

BRIAN L. PIERCE,
ROEL R. LOPEZ, AND
NOVA J. SILVY

Estimating Animal Abundance

INTRODUCTION

CHAPTERS ON **census** methods in The Wildlife Society's "techniques manual" have exploded from 9 pages in the first manual (Wight 1938) to 48 pages in the sixth manual (Lancia et al. 2005). This expansion is testament to the volume of literature produced over the years on this subject, and it has not subsided since the sixth manual. Indeed, the subject has spawned a voluminous literature over the years, including many in-depth books (Caughley 1977; Seber 1982; Caughley and Sinclair 1994; Sutherland 1996; Krebs 1999; Thompson et al. 1998; Buckland et al. 2001, 2004; Borchers et al. 2002; Williams et al. 2002a) on this subject, leading us to ponder how to properly balance coverage of the subject and our intended audience in a limited number of pages.

This chapter differs from those in previous editions (Lancia et al. 1994, 2005) in that we have designed the chapter for use in an undergraduate wildlife techniques class. Our intent is to provide an overview of the basic and most widely used **population estimation techniques**. As pointed out by Lancia et al. (2005), there are several possible approaches to writing a chapter dealing with population estimation that include (1) supplying a detailed treatment that focuses on statistical models and deriving estimators based on these models, (2) providing details on survey protocol design and actually applying different population estimation techniques, or (3) providing the conceptual basis underlying the various estimation methods. Lancia et al. (2005) chose to do the latter. We have chosen the second approach, recognizing, as noted by Lancia et al. (2005), that such an approach has limitations due to the diversity of real-world circumstances and our inability to provide detailed instructions for all possible situations. As such, we do not present all variations of the basic population estimation procedures, but rather provide citations for the relevant literature and computer software where variations of these estimators can be found. However, we believe that a more concise chapter using simple examples will provide a much needed introduction for students, while providing a reference for wildlife biologists and resource managers.

Here we provide an overview of factors that should be considered before choosing a method to estimate population abundance, the pros and cons of using various methods, relevant literature, and available computer software, so the reader may make informed decisions based on their particular needs. For readers with a more quantitative background, literature citations provide access to more detailed coverage of the topics discussed in this chapter.

DEFINITIONS

As terms are used in this chapter, they are defined in relation to population estimation to help the reader understand the material in the chapter. Definitions are based on Overton and Davis (1969), Caughley (1977), Cochran (1977), White et al. (1982), Verner (1985), Caughley and Sinclair (1994), Sokal and Rohlf (1995), Sutherland (1996), Zar (1996), Thompson et al. (1998), Krebs (1999), and Ott and Longnecker (2008).

Population Definitions

Population: A group of animals of the same species occupying a given area (**study area**) at a given time.

Absolute abundance: The number of individuals.

Relative abundance: The number of individuals in a population at one place and/or time period, relative to the number of individuals in a different place and/or time period.

Population density: The number of individuals per unit area.

Relative density: The density in one place and/or time period, relative to the density in another place and/or time period.

Population trend: The change in numbers of individuals over time.

Census: A total count of an animal population.

Census method: The method (e.g., spotlight count) used to obtain data for an estimate of population abundance.

Population estimate: A numerical approximation of total population size.

Population estimator: A mathematical formula used to compute a population estimate calculated from data collected from a sampled animal population.

Closed population: A sampled population in which births, deaths, emigration, and immigration do not occur during the sampling period.

Open population: A sampled population that is not closed.

Population index: A statistic that is assumed to be related to population size.

Detection probability: The probability that an individual animal in a sampled population is detected. Synonyms include **observability**, **sightability**, **catchability**, **detectability**, and **probability of detection**.

Statistical Definitions

Parameter: An attribute (e.g., percentage of females) of a population. If you know the parameters of the population, you do not need statistics.

Statistic: An attribute (e.g., percentage of females) from a sample taken from the population.

Frequency of occurrence: The observed number of occurrences of an attribute relative to total possible number of occurrences of that attribute (e.g., individual was observed on 4 of 5 spotlight counts).

Accuracy: A measure of bias error, or how close a statistic (e.g., a population estimate) taken from a sample is to the population parameter (e.g., actual abundance).

Bias: The difference between an estimate of population abundance and the true population size. However, without knowledge of the true population size, bias is unknown.

Mean estimate: The average of repeated sample population estimates usually taken over a short time period.

Precision: A measure of the variation in estimates obtained from repeated samples. Precision can be measured by (1) **range** (difference between lowest and highest estimates), (2) **variance** (sum of the squared deviations of each n sample measurements from the mean divided by $n - 1$), (3) **standard deviation** (positive square root of the variance), (4) **standard error** (the sample's standard deviation divided by $\sqrt{n}$). It therefore estimates the standard deviation of sampled means based on the population mean), and (5) **confidence interval** (probability that a given estimate will fall within n standard errors of the mean; e.g., a 95% confidence interval would be ± 2 standard errors).

Central Limit Theorem: A statistical theorem stating that for large sample sizes (> 30), the **sampling distribution** of any statistic (e.g., the distribution of means obtained by repeated sampling of the mean from the same population) will be approximately **normally distributed** (form a symmetrical, bell-shaped frequency histogram). Therefore, we can divide the normal curve for the sampling distribution of means into sections represented by n standard deviations above and below the mean. When this is done, 68.26% of the area lies within ± 1 standard deviation, and approximately 95% lies within ± 2 standard deviations of the mean. Accordingly, a **95% confidence interval** implies a range of values within which 95% of the estimated means would fall. Stated differently, there is a 95% chance the true mean lies within ± 2 standard errors of the estimated mean, provided there is no bias in the estimate.

Overton and Davis (1969), in the third edition of *Wildlife Management Techniques Manual*, provided a pictorial presentation (Fig. 11.1) of the relationship between **precision** and **accuracy** that made them easy to visualize. The bull's eye on the rifle target represents the **true population abundance**. If one were to fire 10 shots from a rifle, the 10 bullet strikes would represent the value of each of the 10 individual **population estimates**. The center of the area circumscribed by these 10 shots would then represent where the rifle is firing, on average, or the overall average estimate of population abundance. The distance from the center of all shots fired to the center of the bull's eye represents bias, or the amount of inaccuracy present during those 10 shots.

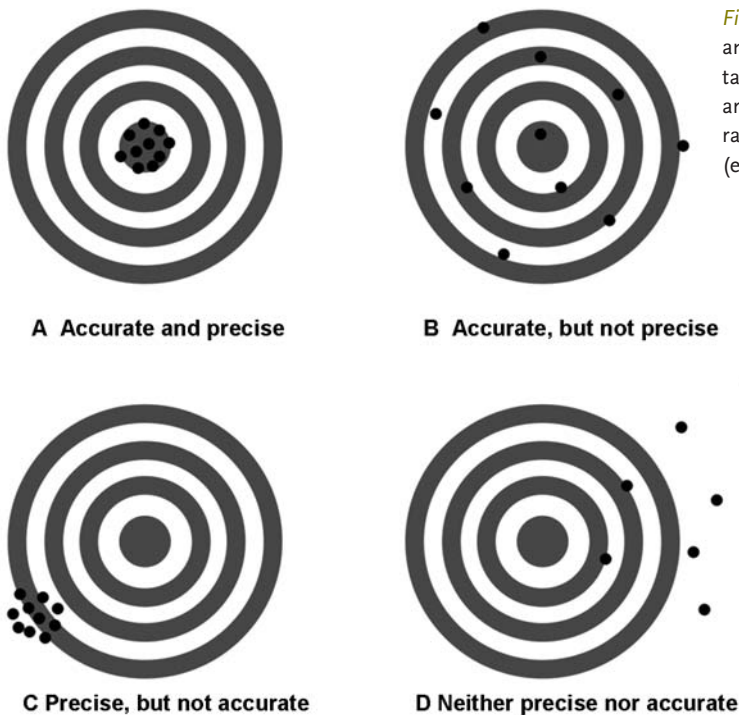


Fig. 11.1. An analogy of precision and accuracy when estimating animal abundance or firing a rifle at a target. Note that in target C, shots are biased to lower left, and in target D, shots are biased to right. When estimating animal abundance, we rarely know in which direction our estimation may be biased (either low or high). Modified from Overton and Davis (1969).

The spread of the bullet strikes would represent **precision** of the population estimates. **Variance** is used to measure precision; the smaller the spread, the smaller the variance and the better the precision will be. For **perfect precision** and **perfect accuracy**, all 10 shots would strike the bull's eye (Fig. 11.1A). However, one can have **poor precision**, but still maintain overall **mean accuracy** if the center of the area circumscribed by the 10 bullet strikes falls on the bull's eye, thereby giving a **mean estimate** equal to the **true population abundance** (Fig. 11.1B). In the same way, one can have **poor accuracy** with **perfect precision** if all bullet strikes hit in 1 spot biased away from the bull's eye (Fig. 11.1C). The worst-case scenario would be to have **poor accuracy** with **poor precision** (Fig. 11.1D). In the real world of population estimation, one does not ever know where the bull's eye lies; therefore, one can only measure **precision** of the estimates.

In practice, population estimates need to be at least precise to be useful. If estimates can be replicated many times in a short time frame, precision can be increased. And, if an estimator or method has good precision, it might be useful as an indicator of population trend, even if it is not accurate. However, if field conditions change (even during the same field season), precision may not increase (Rakestraw et al. 1998). Furthermore, using trend data to manage wildlife populations can be problematic, as the basic assumption when using trend data is that nothing changes over time except population abundance. So, although precision is easy to compute, in real wildlife populations the true population abundance is never known, and therefore accuracy **cannot**

be computed. It can only be implied by the sum of all evidence at hand. As such, if one needs information on population abundance, accuracy is still paramount. Hence the warning precision is no surrogate for accuracy (Lancia et al. 2005).

SURVEY DESIGN

The solution to obtaining a usable estimate of abundance is to choose the right **method** (sampling and/or analysis technique) and to employ proper **survey design or experimental design** (scheme or plan used to obtain samples for abundance or density estimation; see Chapter 1, This Volume). Both must be optimized for the particular circumstance and species to obtain precise (and hopefully accurate) population estimates. Unfortunately, what may work well in some circumstances is useless in others. In addition, there are many combinations of methods and survey designs to choose from, and these can differ by orders of magnitude in their precision and expense. Likewise, there are many opportunities to encounter setbacks and failure. Hence, before any surveying is attempted, the wildlife manager should ask a number of questions:

1. Have I reviewed the relevant literature on the species and/or method?
2. Do I need an estimate of density, or will an index of relative abundance suffice?
3. What methods are available that meet these criteria?
4. What is the extent of the survey area?

5. Are there any limitations on where I can sample?
6. What are the experimental units from which samples will be drawn?
7. How much precision is desired?
8. If comparing areas or time periods, how small a difference must be detected?
9. Given the precision or difference to be detected, how much replication is required?
10. How much replication can I afford?
11. What is the distribution of the species to be surveyed?
12. How will the sample units be distributed?
13. Will sample units be drawn with or without replacement?
14. Do I have the necessary equipment and infrastructure?
15. Do I have sufficient funds to conduct the proposed survey?
16. Is that money better spent on answering another question?
17. Do I have the time required to complete the estimate?
18. Do I have the expertise to collect and analyze the data, or is it available elsewhere?
19. Are there other biologists and biometricians who can provide an independent review?
20. Will I need a pilot study to answer any of the above questions?

Answering the above questions is absolutely necessary, the completion of which should result in a project proposal. Note, the process is iterative and may require several attempts to reach an optimum set of conditions for your particular project. This is typically a good time to contact other biologists and/or biometricians for help, and at the very least to request an independent review of your proposal prior to initiating any work.

Survey Extent

Population estimates typically occur over a defined spatial area, with the estimates representing a specific period of time. As simple as this idea may seem, it is imperative that you define the spatial and temporal extent of the area over which inference is to be made. Answers to these questions will lay the foundation for the statistical analyses ahead and are integral to proper survey design. Integral to this design is an assessment of any nonhabitat and/or nonaccessible areas (private property or dangerous conditions) that may affect species and/or sample distribution.

Experimental Units

Because of the limits of time and costs, a survey of the entire study area of interest is usually not possible. Therefore, an **experimental design** is devised to select a portion of the study area to be sampled (**experimental units**). By definition, experimental units are homogeneous and should be representative of the population or treatment to which in-

ference is to be applied. Experimental units may be time periods, units of space, groups of animals, or an individual animal. It is from **experimental units** that **samples units** are drawn (**replication**). For example, if mice in a cage are given a treatment in diet (e.g., food type A), the cage of animals is the experimental unit, and mice in the cage are sample units. Likewise, if we are comparing abundance among habitat units, the differing habitats are the experimental units, and each survey would be a sample drawn from each of the habitats. In simple surveys, where a population estimate is to be obtained from a single entity with no treatments or controls, there is only one experimental unit.

An **experimental unit** is the smallest entity to which a **treatment** can be randomly assigned (see Chapters 1 and 2, This Volume). If the treatments are manipulative (applied by the experimenter), a **randomization rule** is used to ensure an unbiased assignment of treatments to experimental units. If the treatments are mensurative (categories of time or space; Hurlbert 1984) or organismal (natural categories, such as age class or sex), the randomization rule ensures that experimental units are drawn randomly from each treatment. Thus, proper experimental design helps minimize the effects of uncontrolled variation, allowing you to obtain unbiased estimates of abundance and experimental error (variation among experimental units treated alike).

Sample Units and Sampling Design

Sample units are the entity from which measurements are obtained. Sampling units may be quadrats, transects, or points. Selection of sample units from an experimental unit should be done using a probability sampling scheme, or **sampling design**, where every sample unit has some probability of being selected, and this probability can be accurately determined. Without some type of **randomization rule**, there is no way to avoid discrimination or favoritism in sample unit selection, resulting in **bias** (inaccuracy) and unrepresentative estimates of **variance** (precision) in the estimate of abundance.

Several sampling designs exist to accommodate particular survey conditions (Cochran 1977). The most common sampling design is **simple random sampling**, where sample units are selected randomly to ensure that each sample unit has an equal probability of being selected. You proceed by exhaustively subdividing the experimental unit into sample units, and then you may draw lots, flip a coin, roll dice, or use a random number table to select units to be sampled. During random sampling, sample units may be drawn with replacement (i.e., a sample unit is selected and then placed back into the pool of possible sample units, where it may possibly be drawn again) or without replacement (i.e., sample units may be selected only once). Because sampling without replacement is more precise than sampling with replacement, it is more commonly used in wildlife management (Caughley and Sinclair 1994).

Stratified random sampling is employed when there are implicit differences in sample units that need to be accounted for in the analysis. For instance, differences in habitat quality may produce localized differences in animal density, resulting in increased variance. To reduce variance, the area may be stratified by habitat quality, with sample units selected randomly from each habitat type. For example, a large survey has defined experimental units (areas of homogeneous habitat) by physiognomy (grassland, forest, savanna, desert, etc.). But investigation revealed that controlled burns in each experimental unit created perturbations in the underlying physiognomic matrix (alterations in otherwise homogeneous experimental material). To account for the variability, experimental units are stratified into burned and unburned areas, and sample units are randomly obtained from each stratum.

Systematic sampling (or **systematic random sampling**) is employed to reduce the amount of effort (time or fuel) necessary to navigate among sample units. Systematic sampling typically uses a random start point and then proceeds in an ordered fashion (e.g., a point grid where a sample is collected every 200 m) until the entire area to be covered is sampled. It has the advantage of ensuring thorough coverage of area under investigation, but is susceptible to an array of problems (Cochran 1977), the most pernicious of which is the possible coexistence of an unknown periodic variation in the population being sampled (Krebs 1999). The periodic fluctuation could match the frequency of a systematic sampling design, resulting in a biased estimate with unrepresentative precision (that is unknown to the user).

Several **nonprobabilistic sampling designs** that may be used in error have been described in the literature (Cochran 1977, Krebs 1999), such as accessibility sampling (sampling along trails or roads due to ease of access; later called “convenience sampling” by Anderson [2001]), haphazard sampling (without a plan, as the name implies), or judgmental sampling (selected as “typical” or “representative” on the basis of subjective opinion). Even worse, some sample units may be selected because of the greater opportunity to “see more animals,” despite the obvious bias that will result. Regardless of cause or origin, nonprobabilistic sampling designs are likely to yield biased estimates with levels of precision that are not representative of the area of inference, and they should therefore be avoided.

Sampling Intensity and Statistical Power

Sampling intensity is a concept that encompasses desired **precision**, statistical **power**, and the amount of **variability** among the sample units. Determining the sample size required to achieve study objectives is a central question that must be addressed prior to the initiation of work. If the sample obtained is too big, valuable resources will be wasted obtaining excess precision that produces no change in outcome or conclusions. More catastrophic is a sample size that

is too small, as the information obtained may be incapable of producing useful results, leading to incorrect conclusions. Sampling intensity also is an ethical consideration. Studies with improper sample size exposes subjects (animals or humans) to risks when little (too many samples) or no (too few samples) gain in useful knowledge is possible. Lenth (2001) observed that for such an important and complex issue, there was an alarming paucity of published literature. Fortunately, most popular statistical packages (R [<http://www.r-project.org/>], SPSS [<http://www.spss.com/>], SAS [<http://www.sas.com/>], JMP [<http://www.jmp.com/>], and Statistica [<http://www.statsoft.com/>]) have the tools for sample size determination, and there is a growing number of resources devoted specifically to the task, including books (Armitage and Colton 2005, Chow et al. 2008, Dattalo 2008, van Belle 2008, Julious 2010), standalone software packages (Thomas and Krebs 1997, Lenth 2001, Faul et al. 2007), and several online calculators (Lenth 2001).

There are 5 interrelated components that influence sample size determinations and the conclusions you might reach from a statistical test in a research project. The logic of statistical inference with respect to these components is often difficult to understand and explain (see Chapter 1, This Volume). Here we clarify the 5 components and describe their interrelationships

1. **Significance level:** The significance level is the odds the observed result is due to chance. This concept includes 2 components that define the types of errors possible in statistical tests. **Type I error** is rejecting the null hypothesis when it is true, and the probability of committing this type of error is controlled by the **alpha level** (α) of the test (frequently $\alpha = 0.05$). **Type II error** is failing to reject the null hypothesis when it is false, and the probability of committing this type of error is controlled by the **beta level** (β) of the test (frequently $\beta = 0.05$). The investigator should adjust the levels of alpha and beta according to experimental needs, being mindful of the potential harm that may result from dogmatically applying “typical” or “established” probability levels.
2. **Power:** Power is the odds that you will observe a treatment effect when it occurs. Defined another way, **power** is the probability of rejecting the null hypothesis when it is false, and it is controlled by adjusting beta (i.e., $\text{power} = 1 - \beta$). Increased power results in requisite increases in sample size, due to the relationship between power and beta.
3. **Effect size:** Effect size (d^2) is the difference between treatments (e.g., in number of animals seen) relative to the noise in measurements. **Effect size** expresses the magnitude of difference between 2 sample means and therefore is the logical complement to the *P*-values generated from statistical hypothesis tests. Effect size

and the ability to detect it are indirectly related; the smaller the effect, the more difficult it will be to find, therefore requiring a larger sample size. The term “effect size” is sometimes used synonymously with “standardized difference.” Effect size can be written as

$$d^2 = \frac{\bar{x}_1 - \bar{x}_2}{s},$$

which scales the difference in population means 1 and 2 ($\bar{x}_1 - \bar{x}_2$) by the standard deviation σ (Cohen 1988, van Belle 2008). Although it is useful to think in these terms, one should recognize the dangers of formulating study objectives exclusively in terms of effect size (Lenth 2001, van Belle 2008). For determining sample sizes, it is important to know the anticipated means and variances under the null and alternative hypotheses for the entities being compared.

4. Variation in the response variable: The sample **variance** (s^2) or **standard deviation** (s) are often used to estimate variability in the parameter of interest (e.g., population mean). The standard deviation is calculated as positive square root of the sample variance:

$$s = \sqrt{s^2},$$

where the variance is

$$s^2 = \frac{\sum_{i=1}^n (X_i - \bar{X})^2}{n - 1}.$$

where X_i is 1 data point within a sample and $\bar{X}$ is the mean of all data points within the sample. Similar to the requirement to know the anticipated means for the entities being compared, to accurately determine sample size, we also must estimate the **variance or standard deviation** for the entities being compared. They are typically obtained from either the literature or a pilot study.

5. **Sample size:** Sample size (n) is the number of samples required to obtain the desired precision in an estimate or the desired power in a hypothesis test. Larger sample sizes generally lead to parameter estimates with smaller variances, giving you a greater ability to detect a significant difference. Sample size is typically the variable being solved for in the planning stages, but it can be an input variable when you are attempting to estimate power.

For example, to determine the sample size required for comparing 2 populations with equal variance in a 2-tailed hypothesis (Lehr 1992, van Belle 2008):

$$n = \frac{2(z_1 - \alpha/2 + z_1 - \beta)^2}{\left(\frac{\bar{x}_1 - \bar{x}_2}{s}\right)}.$$

When $\alpha = 0.05$ and $\beta = 0.20$ (typical settings for these parameters in wildlife research), the corresponding critical values from a standardized normal probability table (z -values or z -scores) become 1.96 (z -score for α , the probability of committing a Type I error: $z_1 - \alpha/2$) and 0.84 (z -score for β , the probability of committing a Type II error: $z_1 - \beta$), respectively. The z distribution is a normal or Gaussian distribution with a mean of 0 and a standard deviation of 1. Standardized or z -values then represent deviations from the normalized mean in units of standard deviation. The numerator then simplifies to 15.68. Rounded up to 16 and substituted into the equation, it yields a useful rule of thumb for calculating sample size (Lehr 1992, van Belle 2008):

$$n = \frac{16}{d^2},$$

where

$$d^2 = \frac{\bar{x}_1 - \bar{x}_2}{s},$$

the standardized difference, reflects the difference to be detected between treatment means (effect size) divided by the standard deviation.

It is clear, the ideal experimental design would be one that minimizes the probability of Type I and Type II errors while maximizing power, given the particular experimental constraints of time and resources. Likewise, the above example illustrates the 5 components that are necessary for determining sample size and conducting power analysis, are not independent. The usual objectives of a power analysis are to calculate the sample size (5) required to achieve the desired power (2), given effect size (3) and sample variability (4), at a predetermined level of significance (1). In studies with limited resources, the maximum sample size (5) will be known. In these instances, power analysis then becomes necessary to determine whether sufficient power (2) can be achieved with the known sample size (5), for the desired significance values (1), sample means (3), and sample variances (4). The researcher can then evaluate whether the study is worth pursuing. As indicated above, there are many software packages available for calculating sample size and power (Thomas and Krebs 1997, Lenth 2001, Faul et al. 2007, R Development Core Team 2008). Consult the user's manual of the software package you are using to become familiar with these calculations. The goal is to achieve a balance of components that provides the maximum level of power to detect an effect if one exists, given programmatic, logistical, or financial constraints on the other components.

Proposal Generation and Independent Review

We began this section with a list of questions that should be addressed when developing a survey design. By answering these questions, the researcher should have gained sufficient understanding of the task at hand to finalize the process

with a research proposal. Although many view the writing of a research proposal as an unnecessary formality, we believe that it is an essential part of wildlife management. The steps required to gather the information necessary to write a research proposal forces the investigator(s) to assess the various parameters that will ultimately determine the success, or failure, of a project. The written proposal then represents the investigator's understanding of the problem at hand, as well as the resources and methods believed to provide the solution, given any limitations. As such, the proposal conveys all the information necessary for an independent review. The independent review provides a critique of the survey design, either confirming a sound design or providing the information necessary to improve on the existing knowledge. Therefore, the independent review serves as either the starting point of a new iterative loop through the whole process or the conclusion of the survey design phase.

METHOD CATEGORIES AND CONSIDERATIONS

Animal survey methods have developed over time, building on established knowledge and growing in sophistication. They can be broadly categorized as census methods or estimates derived from sampling, and they are further subdivided by complete or incomplete detection in samples (Fig. 11.2). Early methods focused on complete **census** of a given population. For animals that were elusive or otherwise difficult to census, methods were developed to census animal **indices**. Indices were typically based on cues or other by-products of animal activity (fecal pellets, nests, burrows, tracks, calls, scrapes, etc.) that were believed to be proportional to animal abundance or density. At the same time, methods were developed for obtaining trends or abundance estimates from exploited populations. Later, methods capitalized on existing methodology and attempted to **estimate** abundance by obtaining **complete counts from sample areas**. Finally, because it was impossible to ascertain whether a complete count had been obtained (i.e., to prove a negative: "no animals were missed"), newer methods of estimation were developed utilizing **incomplete counts from sample areas**. It is through this general classification (modified from Lancia et al. 2005) that we introduce the basic methodology of estimating animal density and/or abundance (Fig. 11.2).

Considerations

As noted above, the breadth and depth of the subject of abundance estimation for animal populations spans many methods. The combination of method and survey design then, in turn, dictates how samples may be combined to estimate means and variances. Chapters 1 and 2 (This Volume) should be consulted for more in-depth discussions of experimental design and analysis of data. We intend to pro-

vide a basic overview of methods available for consideration in each category for assessing animal abundance, providing simple examples from historical methods and references for further investigation. We begin by re-emphasizing 2 factors that must be considered due to their impact on precision and accuracy of methods: distribution of the target species relative to the distribution of samples and detection probability.

Species Distribution

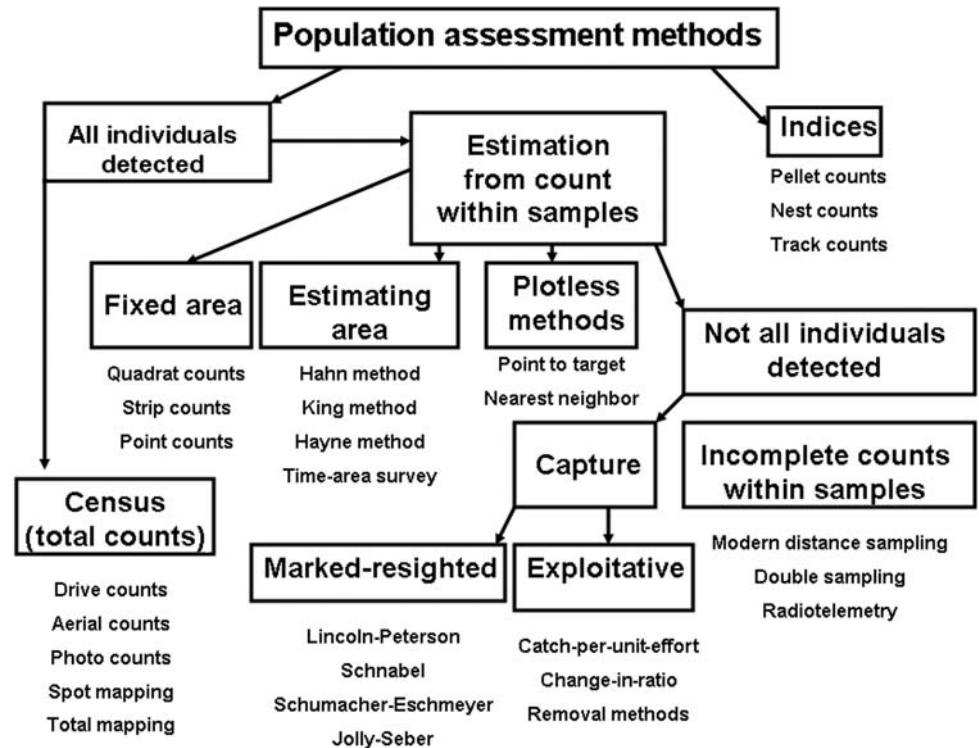
Attempts to manage populations using indices (counts believed to be related to abundance) and complete counts (census) revealed the analytical and practical limitations of these methods. As the size of the area to be surveyed increased, practical limits on available resources were reached, forcing investigators to derive methods for obtaining estimates from samples. Similarly, development of methods for obtaining estimates from samples revealed the importance of sample distribution in relation to species distribution. Resources, and therefore wildlife, are not randomly distributed, which can create bias in estimates of animal abundance. Problems arise when animal distributions are clumped, or when the distribution of samples correlates with the underlying distribution of animals to be sampled. Appropriate survey design is almost always the key component in alleviating this problem, with random sampling or stratified random sampling the most common remedy. Although avoiding the problems resulting from nonrandom distribution of either samples or species is a requisite for obtaining precise and accurate estimates of abundance, defining or describing the underlying distribution of animal abundance is sometimes a necessary objective (Pielou 1974, Cochran 1977, Diggle 1983, Greig-Smith 1983, Ludwig and Reynolds 1988, Ripley 2004). Regardless, we again warn that it is prudent to use probabilistic sampling as the easiest alternative for avoiding unforeseeable problems in obtaining estimates.

Detection Probability

Most animal survey methods do not observe all individuals in the population. Generally, the probability of seeing or trapping an individual animal over a given area is 1. Sampling design and detection probability are major concerns when estimating animal abundance. Usually, one assumes that detection probability is similar across all sampling areas; however, this assumption is not always true, and there may be different detection probabilities for different sampling units. Estimators for these cases take this variation into consideration (Thompson 2002a, Skalski 1994). Lancia et al. (2005) noted that considerable effort in development of abundance estimators has involved ways to estimate detection probability.

Conroy and Nichols (1996), Pollock et al. (2002), and Lancia et al. (2005) noted there are 3 basic approaches used in attempts to deal with variation in detection probability in

Fig. 11.2. The relationship among population estimation procedures.



surveys. The first is to use **standardized methods** when conducting the surveys. All potential sources of variation in detection probability that are under the control of the investigator should be kept constant (methods, effort, observer experience, weather, etc.).

The second involves use of **covariates** in analyses of survey statistics. If exogenous variables, such as weather conditions or observer identity, account for most of the variation in detection probability, models can be developed for estimating change in population size as a function of the relevant exogenous variables (Overton and Davis 1969; Craig et al. 1997; Link and Sauer 1997, 1998, 2002). Lancia et al. (2005) stressed that covariates selected for use in modeling cannot be associated with both detection probability and true abundance. They used the example of vegetation type, because it can affect detection probability, and it also can influence abundance. Therefore, using vegetation type as a covariate affecting detection probability would not be appropriate.

The third approach is to recognize that detection probability is **not constant** over space or time, and that not all exogenous variables can be measured, modeled, or even perceived. This idea leads to implementation of methods that permit direct estimation of detection probability. Of the 3 methods, Lancia et al. (2005) believed this approach was the only one that was scientifically defensible, and they recommended that developers of future surveys and monitoring programs utilize this approach. Despite this recommendation, index use is common in wildlife surveys and

monitoring programs throughout the world (e.g., Thompson et al. 1998).

INDICES

Most **indices** collect **frequency** (number of individual animals or animal sign) information along transects, at quadrats, or points. Examples of index methods include the number of animals seen per kilometer of road, the number of animals present per night at a waterhole, fecal pellets per quadrat, and nest or burrow counts per kilometer of transect. Similarly, a **frequency of occurrence index** only collects presence or absence data. A frequency of occurrence index is based on the proportion of sample units (e.g., scent stations) that contain at least 1 animal or animal sign (Scattergood 1954, Caughley 1977, Seber 1982). However, Seber (1982) noted that a population with a highly clumped distribution will yield a lower frequency of occurrence index (proportion of quadrats with at least 1 animal) than a population of similar density with a more uniform or random distribution.

A **density index** can be defined as any measure that correlates with density (Caughley 1977:12). Indices are used most often because of perceived savings in cost, time, and/or labor. Indices differ from population estimation methods in that only relative abundance or relative density can be derived from the indices. Data are usually presented as deer/km, rabbit pellets/m², or birds/point. Indices can be used to compare animal numbers between treatment and control

areas (e.g., disked with nondisked areas) or to compare the same area over time, based on the assumption that nothing changes except the relative abundance of the animal being studied. The probability of catching, counting, or otherwise detecting an animal in sample units from 2 areas or time periods being compared should be similar. If indices are employed, they should be standardized as to season, time of day, weather conditions, habitat, and observer experience. For multispecies surveys, detection probability will vary with species. Such factors as group status, reproductive cycle, sex and age ratios, and population density also will affect detection probability. For aquatic species, water level, water temperature, and moonlight also affect detection probability.

Data obtained from indices (e.g., relative abundance) are correlated with abundance in some unknown manner. Standardizing methods and using covariates in an analysis can address some sources of variation in index surveys (Lancia et al. 2005). However, as noted by Lancia et al. (2005), other factors that affect detection cannot be handled in these ways or may not even be identified. They recommended caution and skepticism when using and interpreting indices, and they preferred that all indices include an estimate of detection. There are only 2 ways to obtain the relationship of an index to population abundance: (1) estimate detection probability or (2) estimate population abundance and “calibrate” the index (Caughley 1977, Lancia et al. 2005). However, it should be apparent that if an estimate of population abundance can be obtained, there may be no need to do the index survey. We are of the opinion that a calibrated index would only be applicable for the time and place it was done, as conditions typically change over time and for difference areas. Therefore, we will go a step further and recommend that all indices include an estimate of detection probability, and if possible, only population estimation procedures should be used to obtain animal abundance data. Most index surveys can be readily modified to provide information needed in a population estimation procedure. For example, live catch per unit effort (an index) could be easily modified for a mark-recapture population estimation procedure, or deer seen per kilometer could be easily modified for a line-transect population estimation procedure. Regardless, given the advances in sampling methodology, there are relatively few circumstances where index could not be adequately replaced by a more quantitative method.

CENSUS OR TOTAL COUNTS

In few situations are total counts possible. Total counts may be possible for the number of deer in a small paddock or maybe the number of elephants in a small pasture. But for wild populations, it is seldom possible for wildlife managers to obtain a total count of animals in a give study area. As sample area increases, animals are inevitably missed. If it is possible to obtain a total count, then no descriptive statistics

are needed nor apply. The data obtained are not a sample, but an enumeration of the whole population (i.e., no variability is present, because you counted them all). Total counts are assumed to be accurate and can be used to calibrate (i.e., estimate probability of detection) extensive field surveys (Lancia et al. 2005). Total counts on small areas can be derived from intensive surveys (Tilton et al. 1987), from a known number of marked individuals, or by other ingenious means (Kuvlesky et al. 1989). Several methods (presented below) have been purported to produce accurate population counts in some circumstances. However, we warn investigators there is always a possibility that an unknown number of animals will be missed (e.g., nonsinging male birds in a spot-mapping survey). When this occurs, there are no means to detect bias or assess the precision of the sample. So, although the methods listed below are in this section on total counts, in most cases, data resulting from these methods should be viewed skeptically and probably should be considered indices rather than total counts.

Drive Counts

Drive counts occur over limited areas where “beaters,” or herders, drive animals into an enclosure or past counters to count total animals in the study area. The method works best with large, easy to detect animals, such as deer. Drivers remain in sight of one another at all times (to prevent animals from escaping unseen between observers), spaced along a line, and sweep across an area with well-defined boundaries. In the best case scenario, the area would be surrounded by a high fence or water (Tilton et al. 1987). If not, additional observers are placed along the boundaries to count animals that move in or out of the census area. All observers count only those animals that move past them on their right side (this practice eliminates double counting). The census is the sum of the number of animals moving out of the area or back through the line of drivers, minus any moving into the area ahead of the drivers.

McCullough (1979) compared drive counts with population estimates reconstructed from the age of death of individuals in the population. At low densities, drive counts underestimated the true population, and at high densities, they overestimated the true population. Errors could be as large as 20–30%. Thus, drive counts are probably best viewed as an index of population size.

Aerial Photography

Low altitude photography of flocks of birds (or other groups of animals) is often used as a census technique. The entire assemblage of animals is photographed and later counted to give a complete census. However, it is often difficult to ascertain whether all individuals are visible (e.g., some diving ducks may be under water) to be photographed, and errors in counting undoubtedly are made (Bajzak and Piatt 1990).

Spot Mapping or Territorial Mapping

Spot mapping involves plotting locations of individual birds that are seen or heard on a gridded map during repeated visits to a study area. The technique is most suited to birds that regularly sing or call in exclusive territories. Floaters (i.e., nonterritorial birds) and young of the year are usually not surveyed by this technique. The combined data reveal clusters of locations, assumed to represent centers of activity for individual territories during the breeding season. The total number of clusters in the study area equals the number of clusters completely inside the area plus the sum of fractional parts of clusters on the boundaries. Total number of birds is estimated by multiplying the number of clusters by mean number of birds per cluster, which is normally 2, assuming that birds breed in pairs.

Assumptions of the method (Verner 1985, Bibby et al. 2000) are: (1) populations are constant, and birds remain in exclusive spaces or territories during the sampling period; (2) birds in territories produce cues frequently enough to permit repeated location on successive observational visits; (3) estimated proportions of territories along boundaries are accurate; (4) the estimated mean number of birds represented by each cluster is accurate; and (5) observers are skilled, record data accurately, and are consistent. Verner and Milne (1990) provided evidence that spot mapping results should not be considered to be complete counts, as results can vary among observers (Best 1975, O'Conner and Marchant 1981) and map analysts. At best, spot mapping yields an index.

Total Mapping of Bird Territories

This approach is similar to spot mapping, except breeding birds are first trapped and color banded, prior to surveys to delineate territories. This practice facilitates the identification of individuals. Verner (1985:266) believed that, when thoroughly executed, total mapping was probably the most accurate method of estimating population density of breeding birds. He also believed the method should be used as a standard for evaluating the accuracy of other methods of estimating bird density. However, Bibby et al. (2000:42) noted this method only estimates the population of relatively conspicuous birds holding territories, not **floaters** (i.e., nonterritorial birds) or transients. Assumptions are the same as for territorial mapping (described above).

COUNTS ON SAMPLE PLOTS (FIXED AREA)

It may be possible to obtain complete counts of animals on sample units of limited area within some larger area of interest. The sample units must be suitably sized relative to the organism being considered, to ensure that a complete count is obtained. The area being counted is fixed in terms of length and width prior to the start of the survey. Because all individuals are counted, there is no variation associated

with the density or number of animals seen on the sample plots (unless counts on each sample plot are replicated). Instead, only geographic (plot to plot) variation is a concern. The mean density from all sample plots is then extrapolated to the entire study area, giving an estimate of average density and/or population abundance for the area of inference. This basic sampling method has been modified to use sample units of various shape (quadrats, strips, plots, etc.) and size, depending on circumstance and target species. We refer the reader to Caughley and Sinclair (1994) for their excellent illustration of the advantages and disadvantages of sampling with replacement versus sampling without replacement, and transects (long, narrow rectangles) versus quadrats (squares). Here we focus on simple estimates derived from sampling units of equal size. We provide examples for strip and point counts. We again warn investigators of the possibility that an unknown number of animals may be missed in some or all sample units, resulting in negatively biased estimates of population size and/or density.

Strip Counts

This method is the one of the most commonly used to measure density. The counting unit is a strip or transect, which is merely a long, narrow rectangle of fixed area. Transects are randomly placed across the grain of the topography and landscape. Transect lines can be traversed on foot or horseback, by truck or boat, or in a helicopter or airplane. The classic strip census uses a preset distance (0.5-strip width) on each side of the transect line, and then only those animals within this predefined distance are counted. Animals observed outside this distance are not counted, and it is assumed that all animals in the strip are counted with certainty. If these assumptions are valid, the population abundance can be estimated using any of the simple population estimators (Cochran 1977, Krebs 1998, Caughley and Sinclair 1994) for samples of equal area, samples of unequal area, or sampling proportional to size. Here we illustrate the calculations of density and abundance from strip counts of equal area, sampled with and without replacement. Density is calculated as the ratio of the sum of counts to the sum of strip areas (see below for variable definitions):

$$D = \frac{\sum x_i}{\sum a}$$

The density obtained on the sample strips is then multiplied by the size of the study area (area of inference) to obtain populations size:

$$N = DA.$$

By combining the 2 formulas, we obtain the simple strip abundance estimator:

$$N = \frac{A \sum x_i}{2Lwn_s}$$

The variance is obtained from the strip counts using

$$s_x^2 = \frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n_s}}{n_s - 1}.$$

The strip count standard error is then

$$SE_{\bar{x}} = \sqrt{\frac{s_x^2}{n_s}}.$$

The variance of the population estimate when sampling with replacement (SWR) is

$$s_N^2 = \frac{(n_t)^2}{n_s} s_x^2,$$

where n_t is the total number of samples possible on the area of inference (calculated by A/a). The variance of the population estimate when sampling without replacement (SWOR) is

$$s_N^2 = \frac{(n_t)^2}{n_s} s_x^2 \left(1 - \frac{n_s}{n_t}\right).$$

Here the added term (finite population correction) reduces the variance of SWOR relative to SWR by 1 minus the proportion of the area sampled (i.e., the number samples taken over the total number of samples possible on the area of inference). The standard error of the population estimate is then

$$SE_N = \sqrt{SE_{\bar{x}}^2}.$$

Finally, we obtain the 95% confidence interval from

$$95\%CI = N \pm t_{a,df}(SE_N),$$

where N	= population abundance
D	= density of animals in strips
A	= area of inference (study area)
a	= area of each strip ($L \times 2w$)
x_i	= number of animals seen on transect i
w	= preset 0.5-strip width (sample area on each side of transect line)
L	= length of transect
n_s	= number of samples (strips)
n_t	= total possible samples (A/a) in study area
s^2	= sample variance
t	= Student's t for the desired alpha (α) and degrees of freedom ($df = n - 1$)
SE	= standard error
$\bar{x}$	= mean number of animals seen on all transects
$95\%CI$	= 95% confidence interval

Example: We wish to estimate the number of grouse on a 2-km² study area. We utilize 5 counting strips, each 100 m in length with a preset sighting distance of 10 m (0.5-strip width). We divide the study area into strips and select 5 to survey using a random number table. We count each strip, flushing a total of 15 grouse ($x_i = 4, 3, 3, 2$, and 3). The total possible number of samples of this size is 1,000 ($n_t = A/a$). There-

fore, the estimated population abundance would be calculated as follows:

$$N = \frac{(2 \text{ km}^2)(15)}{(2)(0.1 \text{ km})(0.01 \text{ km})(5)} = 3,000.$$

The strip count variance ($s^2 = 0.50$) is then used to obtain the strip count standard error ($SE_{\bar{x}} = 0.7071$). The variance of the population estimate when SWR is then

$$s_N^2 = \frac{(1,000)^2}{5} (0.50) = 100,000,$$

and the standard error of the population estimate when SWR is

$$SE_N = \sqrt{100,000} = 316.23.$$

We can then calculate the 95% confidence intervals when SWR as

$$95\%CI = (\pm 2.776)(316.23) = \pm 877.85.$$

The population estimate, $\pm 95\%CI$ when SWR, is $3,000 \pm 878$ grouse. If we had obtained the counts by SWOR, the population estimate would remain the same, but the variance of the population estimate would change:

$$s_N^2 = \frac{(1,000)^2}{5} (0.50) \left(1 - \frac{5}{1,000}\right) = 99,500.$$

The standard error of the population estimate would become

$$SE_N = \sqrt{99,500} = 315.44,$$

and the resulting 95% confidence interval would be

$$95\%CI = (\pm 2.776)(315.44) = \pm 875.66.$$

The population estimate, $\pm 95\%CI$ when SWOR, is $3,000 \pm 876$ grouse. The increased precision reveals the additional information obtained from n unique samples using SWOR over the possible redundant information contained in repeated samples gathered using SWR. Regardless, we would report the population estimate as $N \pm SE$ (e.g., $3,000 \pm 315.44$ grouse), which would allow other investigators to derive confidence intervals of their choice from the data.

Point Counts

Point counts are typically used to estimate bird density. An observer proceeds to a sample point and might, or might not, allow a rest period of specified duration for equilibration of bird activity (Reynolds et al. 1980). The observer then detects (by both sight and sound) birds for a specified count period within a preset distance (radius) from the point. Although it is generally assumed that all birds are detected within the sample radius, this assumption is typically false unless the preset radius is quite small or the target species is quite conspicuous. Therefore, unless complete counts are certain, point counts should be considered as an index to

relative density. If the assumption is reasonable, then estimation proceeds similar to strip transect counts (described above), differing only in the form of the equation for the simple population estimate:

$$N = \frac{A \sum x_i}{n \pi r^2},$$

where N = population abundance

A = area of study area

x_i = number of birds seen within a fixed radius r of point i

n = total points sampled

π = pi (ratio of the circumference of a circle to its diameter)

r = preset radial distance

Example: A survey consisting of 10 random points, each with a fixed radius of 50 m, is conducted on a 2-km² study area. Surveyors count 50 birds. The estimated population abundance would be calculated as follows:

$$N = \frac{(2 \text{ km}^2)(50)}{(10)(3.1416)(0.050 \text{ km})^2} = 1,273.24.$$

The sample variance (s_x^2), population variance (s_n^2), and population standard error (SE_n) are calculated using the strip count equations for SWOR. We then obtain a point count variance of 0.6667, a population variance of 4,238.13, and population standard error of 65.1. Our calculated 95%CI is then ± 147 birds. Therefore, the population estimate ($\pm 95\%CI$) for the study area is approximately $1,910 \pm 107$ birds. We would report the population estimate $N \pm SE$ (e.g., $1,273 \pm 65$ birds), which would allow other investigators to derive confidence intervals of their choice from the data.

Sample Units of Unequal Area

Samples units of unequal area require an average density to be calculated from all units sampled, as indicated in the discussion of strip counts (above). The average density (D) is then extrapolated to the survey area using $N = DA$. However, the formulas for SWR and SWOR differ for samples of unequal area (Krebs 1998):

$$_{SWR} s_N^2 = \frac{(n_t)^2}{n_s(n_s - 1)} [\sum x_i^2 + D^2 \sum a_i^2 - 2D \sum (x_i a_i)]$$

$$_{SWOR} s_N^2 = \frac{n_t(n_t - n_s)}{n_s(n_s - 1)} [\sum x_i^2 + D^2 \sum a_i^2 - 2D \sum (x_i a_i)],$$

where x_i = count from sample i

a_i = area of sample i

n_s = number of samples taken

n_t = total number of samples in study area

D = average density from the samples

Example: We wish to estimate the number of grouse on a 2-km² study area. From a total of 784 possible transects, we

selected 10 counting strips without replacement. Each strip had a different length, but each was surveyed with a preset sighting distance of 10 m (0.5-strip width) on each side of the centerline. We counted each strip, flushing a total of 50 grouse, with the counts (x) and area (a) of each strip recorded. There are 784 possible transects on the study area. The estimated population abundance would be calculated as follows:

$$D = \frac{50}{0.0255 \text{ km}^2} = 1,960.8$$

and

$$N = (1,960.8)(2) = 3,922.$$

The variance of the population estimate (SWOR) would be calculated as

$$_{SWOR} s_N^2 = \frac{784(784 - 10)}{10(10 - 1)} [256 + (1,960.8)^2(0.00006681) - (2)(1,960.8)(0.1305)] = 7,403.4.$$

Using the equations for strip counts, we obtain a population standard error (SE_N) of 86.0. Our calculated 95%CI is then ± 229 birds. Therefore, the population estimate ($\pm 95\%CI$) for the study area is approximately $3,922 \pm 107$ birds. Again, we would report the population estimate $N \pm SE$ (e.g., $3,922 \pm 86$), which would allow other investigators to derive confidence intervals of their choice from the data.

Sampling with Probability Proportional to Size

Large study areas are seldom homogeneous with respect to resources, species density, or detectability. When this variability occurs, stratification is used to divide the area into units of similar composition. As a result, the units to be sampled are often of unequal size. In these circumstances, one may employ sampling with probability proportional to size (PPS), where the probability of a sample being selected is proportional to the size of the various units being sampled. The PPS method may be used with equal or unequal sized sampling units. Although the PPS method is unbiased and ideally suited for sampling irregular experimental units of differing size, it is limited by design to SWR. Thus, Caughley and Sinclair (1994:202) recommend the method be limited to circumstances where sampling intensity is 15%.

Sampling using the PPS method requires the density to be calculated for each sample, with the average density and variance of the density estimates (s_D^2 ; calculations are the same as sample variance for strip counts above, except they use the density for each sample rather than the count for each sample) to be calculated from all units sampled. The average density (D) is extrapolated to the survey area using $N = DA$. However, the formula for calculating the variance of the total population differs for PPS estimates (Krebs 1998):

$$_{PPS} s_N^2 = \frac{(A)^2}{n_s} s_D^2,$$

where A = total study area size
 n_s = number of samples selected
 D = average density from the samples
 s_D^2 = variance of sample densities

Example: We wish to estimate the number of grouse on a 2-km² study area consisting of 3 vegetation types. We selected 10 samples using sampling PPS. Each strip was 100 m in length with a preset sighting distance of 10 m (0.5-strip width) on each side of the centerline. We counted each strip, flushing a total of 50 grouse ($x_i = 4, 5, 6, 6, 4, 5, 5, 5, 4, 6$). Densities for each sample were calculated (d_i in birds/km² = 2,000, 2,500, 3,000, 3,000, 2,000, 2,500, 2,500, 2,500, 2,000, 3,000), yielding an average density of 2,500 grouse/km², with a variance (s^2) of 166,667. The estimated population abundance ($N = DA$) was 5,000 birds. The variance of the total population was calculated as follows:

$$_{PPS}s_N^2 = \frac{(2)^2}{10} 166,667 = 66,667.$$

The standard error of the population estimate (SE_N) was 258, with a 95%CI of ± 687 birds. We would report the population estimate as $N \pm SE$ (e.g., 5,000 $\pm$ 258 birds), which would allow other investigators to derive confidence intervals of their choice from the data.

COUNTS ON SAMPLE PLOTS (ESTIMATING AREA)

Considerable attention was given to conducting sample counts prior to 1980. In particular, methodology began to center on methods that would allow an accurate estimate of sample area to be obtained from counts without preset strip widths. The thoughts of the day, summarized by Eberhardt (1968), stated that precision was proportional to the square root of the number of animals seen, and therefore efforts should be focused on methods that would allow all sightings to be used. Sightings were expensive to obtain, particularly when many were discarded for being outside the sample frame. The basic solution had several forms, but each attempted to determine the sample area congruent to the area over which counts were obtained.

The King method (Leopold 1933, Buckland et al. 2001) used the average radial distance to all observed animals to estimate the strip width used in the calculations of animal abundance. Kelker (1945) used perpendicular distances to generate a histogram, and from the histogram subjectively determined the strip width over which all animals were likely detected. Hayne (1949a) developed the first widely recognized line-transect density estimator with a solid mathematical foundation (Buckland et al. 2001), based on the sighting distances and angles to flushed birds. Hahn (1949) used visibility measurements, periodically taken perpendicular to the transect line, to estimate the area over which deer were

counted. Density estimates were then based on all detected animals, using average visibility as the estimate of strip width. Robinette et al. (1974) compared the accuracy of these and 6 other early line-transect methods, noting that only the King and Kelker methods showed promise.

We group these methods together based on use of sighting distances to estimate sample area. We refer to this type of distance sampling as **traditional distance sampling**. As **modern distance sampling** has superseded these methods, we provide only the estimators and no examples.

Hahn Method

The Hahn (1949) method is still commonly used to estimate population density. It is very similar to the strip method example provided above, differing only in the use of distances to estimate the strip width defining the sample area. Transects are randomly placed across the grain of the topography and landscape, and they can be traversed on foot, on horseback, or by vehicle. Estimates of maximum visibility are made periodically (e.g., every 200 m) on both sides of the transect, with maximum visibility defined as the maximum distance an observer could see a target animal perpendicular to the transect at each point. The Hahn estimate of population abundance is calculated as

$$N = \frac{A \sum x_i}{2Lv},$$

where N = population abundance
 A = area of study area
 x_i = number of animals seen on transect i
 v = the 0.5-strip width determined by average visibility measurements
 L = total length of all transects

King Method

The King method (Leopold 1933, Buckland et al. 2001) used the average radial distance from all observed animals to estimate the 0.5-strip width to be applied in the calculations of density or abundance. Thus, it is similar to the Hahn method:

$$N = \frac{A \sum x_i}{2L\bar{r}},$$

where N = population abundance
 A = area of study area
 x_i = number of animals seen on transect i
 $\bar{r}$ = the 0.5-strip width determined by average sighting radius
 L = total length of all transects

Hayne Method

The Hayne (1949a) method was commonly used to estimate population density of flushing birds. The method assumed there was a fixed flushing radius for each bird species and habitat. When an observer walking a transect came within

that radial distance, the bird would flush and be spotted. Further, the method assumed the sine of the angle for each observation came from a uniform random distribution ranging from 0 to 1, with an average angle of 32.7° (Hayne 1949a). Later investigators (Robinette et al. 1974, Burnham et al. 1980) determined the mean sighting angles were generally around 40° , with Burnham and Anderson (1976) providing a correction factor for the original Hayne method. The Hayne estimator of density, from Krebs (1998), is

$$D_H = \frac{n}{2L} \left(\frac{1}{n} \sum \frac{1}{r_i} \right)$$

and

$$N = DA,$$

where N = population abundance

D_H = population density

A = area of inference

n = number of animals seen on each transect

r_i = sighting distance to animal i

L = length of transect

The variance associated with this density estimate is calculated as

$$s_{D_H}^2 = D_H \left[\frac{s_n^2}{n^2} + \frac{\sum (1/r_i - R)^2}{R^2 n(n-1)} \right],$$

where D_H = population density

n = number of animals seen

s_n^2 = variance of n

r_i = sighting distance to animal i

R = mean of the reciprocals of sighting distances r_i

Time-Area Squirrel Survey

Time-area surveys are a common method used to census tree squirrels (Goodrum 1940:8). They are a point-based example of using distances to estimate the effective sample area of the counts. Sample points are chosen at random, and counters are stationed at each point (base of a tree nearest to the point) before sunrise. Starting at sunrise, counters wearing camouflaged clothes remain quiet and relatively motionless while counting all squirrels that come into view for 30 minutes. The counter determines the distance to each squirrel when first detected using a laser rangefinder. The average distance to all squirrels detected is then used to compute the area over which the squirrels were counted. Under field conditions, the proportion of a circle observed by each counter will vary from point to point. As such, each observer uses a compass to estimate the portion of a circle under surveillance during the count (e.g., 0.75 or 75% sample effort). This estimate is then factored into the estimation equation (mean area observed by each surveyor). Population size is estimated using

$$N = \frac{A \sum x}{n \Delta \pi r},$$

where N = population abundance

A = area of study area

$\sum x_i$ = number of squirrels seen at point i

n = total points sampled

Δ = average effort in terms of portion of circle observed

π = pi (ratio of the circumference of a circle to its diameter)

r = average radial distance to all detections

The simple strip estimator of variance, standard error, and 95%CI can be used with this method.

COUNTS ON SAMPLE PLOTS (PLOTLESS METHODS)

Although methods of fixed area counts were common in both plant and animal sampling, they suffer from boundary effects, where a decision must be made to determine whether to include each target observed on a plot boundary in the sample, and they are time consuming. Plant biologist developed several "plotless" methods to estimate density and abundance that alleviate these problems and are relatively easy to apply, so long as the target species (e.g., bird nests) remains in place or can be measured before they move (Cottam and Curtis 1956). They have sometimes been referred to as distance methods, because they utilize either point-to-target or target-to-nearest-neighbor distances to estimate density and/or the spatial pattern of the target species.

Two general considerations should be weighed when considering use of plotless methods. The first is the execution of the random sampling design often proposed for this method. Random sampling is great in theory, and reviews well in proposals, but it is difficult and time consuming to achieve in the field. There also is an uncanny proportion of "random" points that do not occur in the thick brush, in the deeper portion of the marsh, on the ant bed, or other "random, but inconvenient" places in the field. Further, Pielou (1977) demonstrated that using random points to select random individuals is biased toward isolated individuals. In some circumstances, systematic random sampling is a good compromise, as the starting points are randomly placed, and they provide broad coverage of the area. Regardless, if you utilize random sampling, then establish a map and/or Global Positioning System (GPS)-based navigation system, allow extra time for navigating to the random points, and develop the willpower to place the points objectively where they fall. The second consideration is the distribution of the target species. Although most methods work well when the target species is randomly or uniformly distributed, many have problems when the target species is clumped or severely clumped (Legendre et al. 2004), and this drawback is especially pronounced for the plotless methods (Engeman et al. 1994).

Point-to-Target and Target-to-Nearest-Neighbor Methods

Byth and Ripley (1980) recommend 2 plotless sampling methods for measuring density and an excellent sampling design procedure for obtaining data from both methods simultaneously:

1. Determine sample size (n) for the density estimate.
2. Set out $2n$ points using a systematic random or other probabilistic sampling design.
3. Randomly select half of the $2n$ points, proceed to those points, and measure the distance from the point to the nearest target species (point-to-target or PTT).
4. On the remaining half of the $2n$ points, lay out a circle of radius sufficient to enclose (on average) the 5 nearest targets. Number these individuals and select n at random. From the randomly selected individuals, measure the distance to the nearest target species (target-to-nearest-neighbor or TNN).

The PTT density is estimated by

$$D_{PTT} = \frac{n}{\pi \sum x_i^2},$$

where D = density

n = number of samples

x_i = distance from point i to nearest target

The TNN density is estimated by

$$D_{TNN} = \frac{n}{\pi \sum x_j^2},$$

where D = density

n = number of samples

x_j = distance from target j to nearest neighbor

The variance for both estimates is calculated from the reciprocal of the density,

$$y = \frac{1}{D},$$

with the variance of y calculated as

$$s_y^2 = \frac{y^2}{n}.$$

The standard error of y is then

$$SE_y = \sqrt{\frac{s_y^2}{n}},$$

where D = density from either the PTT or TNN estimator

n = number of samples

y = reciprocal of the density estimate (D)

Example: We wish to estimate the number of active nest on a 2-km² study area during the breeding season. We used a map to delineate 20 systematic samples and randomly selected

10 for PTT measurements, reserving the other 10 for TNN measurements. At the PTT locations, we obtained the distances ($x_i = 0, 10, 1, 10, 11, 15, 7, 12, 10, 9$), with the sum of squared distances (x^2) equal to 921. At the TNN locations, we obtained the distances ($x_i = 15, 7, 3, 12, 9, 15, 5, 11, 1, 7$), with the sum of squared distances (x^2) equal to 929. As the calculations are the same for each estimator, we illustrate the density estimate from the PTT measurements:

$$D_{PTT} = \frac{10}{(3.14159)(921 \text{ m}^2)} = 0.0035.$$

So we estimate 0.0035 nests/m² or 34.56 nests/ha. The variance of the PTT estimate is

$$s_y^2 = \frac{\left(\frac{1}{0.003456}\right)^2}{10} = 8.371.8.$$

The standard error of the population estimate (SE_y) is

$$SE_y = \sqrt{\frac{8,371.8}{10}} = 28.934.$$

Therefore, the 95%CI for y is

$$95\%CI_y = \pm(2.262)(28.934) = \pm 65.45.$$

The upper and lower bounds on 95%CI are calculated as:

$$\frac{1}{0.003456} + 65.45 = 289.35 + 65.45 = 354.8$$

and

$$\frac{1}{0.003456} - 65.45 = 289.35 - 65.45 = 223.9.$$

We take the reciprocal of the results and multiply by 10,000 to convert to nests per hectare, so

$$\left(\frac{1}{354.8}\right)(10,000) = 28.18$$

and

$$\left(\frac{1}{223.9}\right)(10,000) = 44.66.$$

Therefore we have a mean of 34.56 nests/ha with 95%CI of 28–45 nests/ha.

Point-Quarter Method

The point-quarter method is a classic for sampling vegetation that dates back to the first land surveys in the United States. Surveyors would locate and describe the 4 trees nearest to each corner of a section (1 square mile) of land. The method was used by Cottam and Curtis (1956) for estimating forest species and continues to be used today. The method has application to animal density estimates as long as the target species (e.g., bird nests) remains in place or can be measured before they move. Using this technique, selected points from a sampling design are located in the field,

and the area around the point is precisely divided into 4 (90°) quadrants (either perpendicular to the transect for point-transect sampling, or by compass bearing for random points). The distance from the point to the nearest target within each quadrant is measured, so that 4 distances are obtained at each point. The population density is then calculated as (Pollard 1971, Krebs 1998)

$$D_{PQ} = \frac{4(4n - 1)}{\pi \sum (x_{ij}^2)},$$

where D = point-quarter estimate of density
 n = number of points sampled
 x_{ij} = distances from point i to the nearest target in quadrant j

Variance of the density estimate is

$$s_{PQ}^2 = \frac{D_{PQ}}{4n - 2}.$$

The standard error of the density estimate is

$$SE_{PQ} = \sqrt{\frac{s_{PQ}^2}{4n}}.$$

The 95%CI can be obtained by

$$95\%CL_{PQ} = \left(\frac{\sqrt{16n - 1} \pm 1.96}{\sqrt{\pi \sum (x_i^2)}} \right)^2.$$

Example: We wish to estimate the number of active nests on a 2-km² study area during the breeding season. We use a map to delineate a point transect through a patch of forest, with 5 points spaced at 100 m. At the 5 locations, we obtain the distances ($x_i = 0, 10, 1, 10, 11, 15, 7, 12, 10, 9, 15, 7, 3, 12, 9, 15, 5, 11, 1, 7$), with the sum of squared distances (x_i^2) equal to 1,850. The density estimate is

$$D_{PQ} = \frac{(4)[(4)(5) - 1]}{(3.1416)(1850)} = 0.0131$$

with a variance of the density estimate equal to

$$s_{PQ}^2 = \frac{0.01308}{(4)(5) - 2} = 0.000727.$$

The standard error of the density estimate is then

$$SE_{PQ} = \sqrt{\frac{0.00072647}{(4)(5)}} = 0.00603,$$

and the lower and upper bounds on the 95%CI are

$$95\%LCL_{PQ} = \left(\frac{\sqrt{(16)(5) - 1} - 1.96}{\sqrt{(3.1416)(1,850)}} \right)^2 = 0.00826$$

and

$$95\%LCL_{PQ} = \left(\frac{\sqrt{(16)(5) - 1} + 1.96}{\sqrt{(3.1416)(1,850)}} \right)^2 = 0.02025.$$

The above units are in nests per square meter. We multiply by 10,000 to get nests per hectare, so we have a mean of 131 nests/ha, with 95%CI of 83–202 nests/ha.

COUNTS ON SAMPLE PLOTS (DETECTION PROBABILITY)

The preceding methods for estimating population size either reduced the survey area to ensure complete detection or attempted to correct the survey area to allow for unbounded counts with incomplete detection. The strategy was to either standardize or estimate the survey parameters necessary to obtain accurate estimates without direct evaluation of detection probability. The methods that follow use the opposite strategy: to estimate detection probability directly or collect ancillary data necessary to develop models for predicting detection probability.

Double Sampling

Double sampling (Jolly 1969a, b; Eberhardt and Simmons 1987; Pollock and Kendall 1987; Estes and Jameson 1988; Prenzlow and Lovvorn 1996; Anthony et al. 1999; Bart and Earnst 2002) is a modified form of sampling based on ratio estimation, where a large number of samples are obtained using a rapid method, such as point counts, followed by the surveying of a random subsample of those same plots using an intensive method that determines actual density. In the subsampled area, the densities obtained from the intensive method are used to estimate the proportion of animals seen using the rapid method. The relative probability of detection derived from the ratio of the rapid-method results to actual density is then used to correct estimates obtained from the rapid method over the remaining surveyed region.

The estimate of the proportion of animals seen (β) is the ratio of the mean counts (or density estimate) from the rapid method (y) to the mean count (or density estimate) from the intensive method (x):

$$\beta = \frac{y}{x}.$$

We can then use this estimate of the proportion of animals (β) on the subsamples to correct the population estimate (N) using the rapid method on the larger set of samples:

$$N = \frac{A \sum y}{na\beta},$$

where A = area of the study area (area of inference)

$\sum y$ = sum of counts or density estimates from the rapid method

n = the number of rapid-method samples

a = the area of each rapid-method sample

β = the relative proportion of animals (rapid method verses intensive method)

Jolly (1969a, b) and Pollock and Kendall (1987) presented an estimator for the variance of this estimator.

The **assumptions** of double sampling are the intensive method is accurate and reflects the actual density of the subsamples. Inaccuracy in the intensive method will result in multiplicative bias in the population estimate. For instance, lack of complete detection using the intensive method will create negative bias in the “corrected” population estimator. Similarly, the timing of the counts should coincide, and ideally would be simultaneous, so that both methods sample the same population. Differences in timing will increase variability, and perhaps bias, in the final population estimate.

Double Observer Sampling

Multiple observer methods were developed initially for aerial transect surveys (Caughley 1974, Magnusson et al. 1978, Cook and Jacobson 1979, Grier et al. 1981, Caughley and Grice 1982, Pollock and Kendall 1987, Graham and Bell 1989), but more recently they have been applied to ground point count surveys (e.g., Nichols et al. 2000). These methods can be divided into groups based on use of independent or dependent observers.

Independent Observers

Aerial or surface (ship, car, etc.) transects may be conducted with 2 observers, each collecting observations independently. The animal locations can be annotated on maps by each observer, or precise offset locations (x , y , and time) can be obtained using survey equipment (total station or GPS and offset laser rangefinder), allowing maps to be created post-survey. The mapped data are assigned to categories based on the type of detection: those seen by observer 1, those seen by observer 2, and those seen by both observers as in the equation below. Caughley (1974) demonstrated that data of this sort can be analyzed using the Lincoln–Petersen estimator (see Marked–Resight Methods later in this chapter) to estimate population size in the surveyed area (Grier et al. 1981, Caughley and Grice 1982, Pollock and Kendall 1987):

$$N = \frac{n_1 n_2}{m},$$

where N = population size in the area of inference

n_1 = total number of animals seen by observer 1

n_2 = total number of animals seen by observer 2

m = total number of animals seen by both observers

The method has several assumptions that will affect precision and accuracy:

1. Observations must be independent.
2. Category assignments must be accurate.
3. Targets must have equal detectability.

The **assumption** of independence of sightings between the observers is a strict requirement that may be difficult to

achieve. For example, the independence assumption will be violated if the activity of one observer, such as speaking into a tape recorder or writing on a map, alerts the other observer to an animal's presence. Likewise, if separate surveys are conducted (e.g., ground and aerial), different observers should be used to ensure independence. Further, all animals must have equal detection probabilities, but these probabilities can differ between the 2 observers. If some animals differ in detectability (e.g., if males are more conspicuous than females), the resulting heterogeneity will produce negative bias in the Lincoln–Petersen estimator. However, Magnusson et al. (1978) noted the assumption of equal detection probabilities is not critical. Observation locations from each observer must be precise and unambiguous, or categorical assignments will be inaccurate. Similarly, because animal movement may contribute to this problem, surveys of mobile animals should be conducted simultaneously, so that each observer views the same sample population. Immobile targets (nests, middens, lodges, etc.) pose no such problem, and therefore, separate surveys may be made so long as the sample frame remains the same. Chapman (1951) provided a modified estimator with less bias:

$$N = \frac{(n_1 + 1)(n_2 + 1)}{m + 1} - 1,$$

and the variance of N was provided by Seber (1982):

$$s_N^2 = \frac{(n_1 + 1)(n_2 + 1)(n_1 - m)(n_2 - m)}{(m + 1)^2(m + 2)}.$$

The method also has been used to estimate bird abundance from fixed-radius point counts, using 2 independent observers at each point. The point method requires there be no undetected movement into or out of the fixed radius, and that each observation must be accurately assigned as either inside or outside the fixed radius.

Dependent Observers

Another double observer approach involves 2 observers working in tandem. One is designated as the primary observer, the other as the secondary observer. The primary observer detects animals and reveals all sightings to the secondary observer. The secondary observer then records any additional sightings independently. Animal locations can be annotated on maps by each observer, or precise offset locations (x , y , and time) can be obtained using survey equipment (total station or GPS and offset laser rangefinder), allowing maps to be created post-survey. The mapped data are assigned to categories based on the type of detection: those seen by observer 1 and those additional animals seen by observer 2. Assuming equal detection probabilities for the 2 observers, we can obtain estimation of population size under the 2-sample removal model (Seber 1982, Pollock and Kendall 1987):

$$N = \frac{n_1^2}{n_1 - n_2}.$$

The variance of N is estimated as

$$s_N^2 = \frac{n_1^2 n_2^2 (n_1 + n_2)}{(n_1 - n_2)^4}.$$

The probability of an animal being detected is

$$p = 1 - \left(\frac{n_2}{n_1 + 1} \right)$$

and the variance of the detection probability is

$$s_p^2 = \frac{n_2(n_1 + n_2)}{n_1^3},$$

where N = population size in the area of inference

n_1 = total number of animals seen by observer 1

n_2 = total number of animals seen by observer 2

p = probability of an animal being detected

As with the independent observer approach, heterogeneous detection probabilities will produce negatively biased estimates of population size. Pollock and Kendall (1987) noted this method does not use the number of animals seen by both observers, and it assumes both observers have equal sighting probabilities. Therefore, it may not be as useful as the independent double observer method using the Lincoln–Petersen estimator.

Cook and Jacobson (1979) developed a similar dependent double-observer approach for transect surveys, but it has the 2 observers switch roles midway through the survey to overcome the possible difference in detectability between the observers. This method assumes that swapping roles does not alter the detection probability of the observers, all other assumptions being the same as above. Nichols et al. (2000) suggested applying the Cook and Jacobson (1979) method to estimate bird abundance from fixed-radius point counts, noting the model (DOBSERV; Nichols et al. 2000) permits estimation of observer-specific detection probabilities and bird abundance.

The advantage of the dependent double-observer approach occurs when there are practical or logistical reasons prohibiting the use of the independent double-observer method. The disadvantage is the dependent approach is less efficient than the independent approach, because capture–recapture methods are more efficient than removal methods (Seber 1982:324, Pollock and Kendall 1987:505). Therefore, we agree with Pollock and Kendall (1987) the independent approach using the Lincoln–Petersen estimator is more precise, simpler to understand, and allows the 2 observers to have different sighting probabilities.

Generalizations using program MARK (White and Burnham 1999) or DOBSERV (<http://www.mbr-pwrc.usgs.gov/software.html>) give the researcher the option to fit generalized Lincoln–Petersen models that allow for detection probability to depend on covariates, such as species, wind speed, and distance. MARK and DOBSERV use Akaike's Information Criterion (AIC; Burnham and Anderson 1998, 2002) to

pick the most parsimonious model that explains the data adequately.

Marked Sample

We can use marked animals in a population to estimate detection probabilities. Using this method, some marked animals are released into the population and are therefore available for detection at the time of the survey. Marked and unmarked animals are counted during the survey, and the probability of detection for the marked animals can be estimated as

$$\beta = \frac{m}{n_1}.$$

By rearranging the terms, we get the Lincoln–Petersen estimator:

$$N = \frac{n_2}{\beta} = \frac{n_1 n_2}{m},$$

where N = total population size in the surveyed area

n_1 = number of marked animals present in the area at the time of the survey

n_2 = number of animals (both marked and unmarked) seen during the survey

m = number of marked animals seen during the survey

β = probability of detection

In practice, we recommend use the bias-adjusted modification of this estimator provided by Chapman (1951):

$$N = \frac{(n_1 + 1)(n_2 + 1)}{m + 1} - 1.$$

Although this approach is straightforward, the practical aspects require careful consideration. Marked and unmarked individuals must have the same probability of being detected. The mark must be conspicuous, so that no marked animals are erroneously or inadvertently recorded as unmarked. But the mark must not be so obvious that it draws attention to marked animals, making them more visible than unmarked animals. Further, it is necessary to determine how many marked animals are actually present for observation during the survey. The number present, and therefore available for observation, is frequently not equal to the number released. Radiotelemetry is a commonly used approach, as it can be used to determine the number of radiomarked animals in the surveyed area at the time of the survey (e.g., Packard et al. 1985, Samuel et al. 1987) and to verify whether each animal seen is marked. However, it is not necessary to have individually identifiable animals, and batch marks (e.g., collars with no alphanumeric identification code) will suffice, so long as the number of marked animals available for detection can be determined prior to the survey. Similarly, any marked animals seen during the survey that were not known to be present prior to the survey are not included in n_1 . They are treated as unmarked in the survey data and included in n_2 , but not in m .

Program NOREMARK (White 1996; <http://www.cnr.colostate.edu/~gwhite/software.html>) provides multiple estimators to determine the number of animals in the study area (Bartmann et al. 1987, White and Garrott 1990, Neal et al. 1993), a simulation capability for anticipating estimator performance, routines for estimating sample sizes, and a simulation for calculating the relative effort to put into marking versus resighting.

Time of Detection

Farnsworth et al. (2002) were the first to recognize that useful information on detection probabilities were available from the times when birds were detected in point count surveys. Their method was a modification of removal methods that used only the time interval when a bird was first detected to estimate detection probabilities. Similar to the development chronology of double observer methods, more recent work (Allredge et al. 2007a) has extended the approach using a capture–recapture formulation, because capture–recapture methods are generally more efficient than removal methods (Seber 1982:324, Pollock and Kendall 1987: 505). Both approaches capitalize on the common practice of recording data at point counts in temporal intervals, where the number of birds counted is separated into those first observed in the first 3 minutes, those first observed in the next 2 minutes, and those first observed in the final 5 minutes. This procedure was recommended by Ralph et al. (1995) and was originally designed to allow results from 10-minute counts to be comparable with those from studies employing 3- and 5-minute counts.

Using the removal method of Farnsworth et al. (2002), the simplest application of the time of detection approach can be illustrated with just 2 time intervals of equal duration. Suppose that an observer records all birds seen and/or heard in the first 5 minutes and then records any additional birds detected in the second 5 minutes. We can then define x_1 as the number of birds counted in the first time interval and x_2 as the number of new birds (not detected in the first period) detected in the second time interval. The expected values of the random variables x_1 and x_2 are

$$E(x_1) = Np_1$$

$$E(x_2) = N(1 - p_1)p_2,$$

where N = total number of birds within the detection radius of the observer

p_1 = detection probability for an individual bird in the first time period

p_2 = detection probability for an individual bird in the second time period

The term $(1 - p_1)$ is necessary, because all birds first detected in the second interval must, by definition, have been missed in the first time interval. If we assume the detection probability for the 2 intervals is equal (i.e., $p_1 = p_2 = p$), solving the

above equations for p and N produces the moment estimator (Zippin 1958)

$$p = \frac{x_1 - x_2}{x_1}$$

and the population estimator

$$N = \frac{x_1^2}{x_1 - x_2} = \frac{x_1}{p}.$$

The estimators can fail if $x_1 \leq x_2$, which is possible when p is small. We present this 2-sample removal estimator to illustrate the approach with the simplest possible situation.

Example: During the first 5 minutes, we observe 20 birds, and during the second 5 minutes, we observe 5 birds that we did not observe during the first 5 minutes. The probability of detection is then

$$p = (20 - 5)/(20) = 0.75$$

and the population estimate is

$$N = (20)(20)/(20 - 5) = 20/0.75 = 26.67 = 27 \text{ birds.}$$

In practice, we use >2 intervals, because doing so permits relaxation of the assumption of equal detection for different species. For instance, Farnsworth et al. (2002) present a more general model with 3 count intervals of variable length, allowing for differences in detection probabilities among intervals and heterogeneity of detection among individual birds. These differences are taken into account by assuming there are 2 groups of individuals in an unknown proportion, and that all members of the first group are detected in the first time interval.

Allredge et al. (2007a) suggested a more efficient approach using a closed population capture–recapture model with k time intervals to account for more sources of variability in the point count data. Their method was specifically designed to account for variation in detection probabilities associated with the singing rate of birds, by modeling both availability and detection bias. They recommended using ≥ 4 equal intervals to reduce assumptions. For example, the assumption of constant detection rates over time required by the removal model of Farnsworth et al. (2002) is not required in the capture–recapture approach, because it models temporal variation from the full detection history.

Assumptions of the general time-of-detection method (Farnsworth et al. 2002, Allredge et al. 2007b) model are: (1) there is no change in the population of birds within the detection radius during the point count (i.e., the population is closed and birds do not move into or out of the radius), (2) there is no double counting of individuals, (3) all members of group 1 are detected in the first interval, (4) all members of group 2 that have not yet been detected have a constant per minute probability of being detected, and (5) observers accurately assign birds to within or beyond the

radius used for the fixed radius circle. As noted by Alldredge et al. (2007b), these restrictions are not trivial, because movement of individuals and difficulties associated with aural detections may result in violation of all assumptions.

Program CAPTURE can produce maximum likelihood estimates for N , as well as the estimated variance of N (Otis et al. 1978, White et al. 1982), for equal time intervals using the method of Farnsworth et al. (2002). **Program MARK** (White and Burnham 1999) can be used to model detection history over k intervals, constant detection probability for all individuals, time effects on detection probability, difference due to previous detection, and unobservable heterogeneity, following the method of Alldredge et al. (2007a, b).

Modern Distance Sampling

Modern distance sampling is a widely used method for estimating size or density of biological populations. It is a comprehensive approach that encompasses study design, data collection, and statistical analyses (Buckland et al. 2001). Modern distance sampling is based on the observation that detection probabilities decrease with increasing distance from the observer (Burnham and Anderson 1984). Distance data are used to estimate the specific shape of the **detection function** relating detection probability to distance for a particular target species and set of conditions. We can define the detection function $g(x)$ as the conditional probability of detecting an animal, given that it is located at some distance (w) from the line. Although the various analyses can be quite sophisticated, the data collected along line transects or points counts for modern distance sampling methods are the same data one would use for traditional distance sampling methods. When properly applied, distance sampling yields estimates of absolute density and detection probability, meeting the requirements for inference put forth by Rosenstock et al. (2002). The history and development of distance sampling is described by Buckland et al. (2001), and extensions to the basic theory are covered in Buckland et al. (2004) and Thomas et al. (2010). An extensive reference archive, covering methodological development and practical applications of modern distance sampling, is available on the Distance Sampling website (<http://www.ruwpa.st-and.ac.uk/distancesamplingreferences/>).

In distance sampling, counts are assumed to be incomplete. Thus, the proportion of animals present that are actually seen (β) must be estimated, and actual counts must be corrected by these detection probabilities. Perpendicular or radial distance data are used to estimate these detection probabilities. To examine what a detection function looks like, we can plot a **histogram** using the frequency of detections (y axis) grouped into small distance intervals (x axis) from the center of a transect line (distance 0) to the maximum observation distance (w). If our sample size is large, we can approximate the shape of the detection function by drawing a smooth curve through the top of each distance

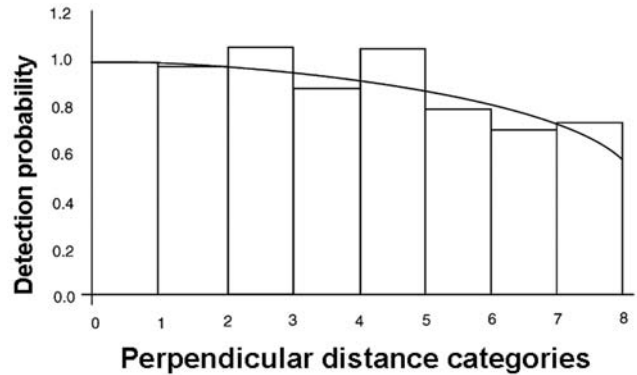


Fig. 11.3. The detection function for the uniform plus one-term detection function for duck nest data. From Anderson and Pospahala (1970).

interval in the histogram (Fig. 11.3). In practice, sample sizes are often too small, and this procedure does not work well.

Survey planning, including sampling design and estimates of sample size, can be performed in program **DISTANCE**. Data collection can be performed using either line transects or points. Analysis of the resulting data typically involves 4 steps: (1) data examination via graphical displays, (2) model fitting using various functions and adjustment terms, (3) model selection using the AIC criteria, and (4) inference under the chosen model. Program **DISTANCE** allows for the fitting of complex detection functions (half normal, uniform, or hazard rate) using a series of adjustment terms (cosine, simple polynomial, or hermite polynomial). Rather than review these models and associated parameter estimators here, we recommend the excellent book by Buckland et al. (2001). A concise overview of distance sampling and program **DISTANCE**, including newly available advanced options, can be found in Thomas et al. (2010). Details concerning the actual use of program **DISTANCE** (Thomas et al. 1998; available at <http://www.ruwpa.st-and.ac.uk/distance>) are contained in the help files provided with the program.

Actual field application of modern distance sampling methods involves many decisions and **considerations** specific to each survey situation. For example, many animals exhibit gregarious behavior and tend to occur in groups. This behavior requires measuring the distance from the line or point to the geometric center of each group and recording the number of animals in present in each cluster. Because groups of animals are easier to detect than individuals, detection bias can occur as a function of group or cluster size. Thus, decisions must be made concerning whether to measure distances to groups or to individuals. The density of groups or clusters along with estimates of cluster size are modeled to improve the precision of estimates of density and population size. Drummer and McDonald (1987) and Otto and Pollock (1990) discussed models for use when detection probability for fixed distance depends on group size.

Another consideration involves grouping of data. Accurate measurement of distances in the field may not be possible; therefore, detections may need to be grouped into distance categories. Even when direct distance measurements are recorded, anomalous patterns may be apparent, such as few objects detected at short distances, heaping of detections at commonly rounded measurements (e.g., 50 m or 100 m), or a relatively large number of detections near the boundary distance. Buckland et al. (2001) recommended truncation of data at distances greater than that at which observations seem likely to be outliers. Further, data may be grouped into a histogram before analysis as a smoothing technique (Buckland et al. 2001). However, exact distance measurements are to be preferred when possible, as they allow the data to be placed into distance intervals during analysis.

Additional problems may arise due to insufficient sample size in terms of observations, transects, or points. The variability between lines and points is an important factor that influences encounter rates (n/k) and detection probability. Failure to obtain a representative sample of the true variability within a population will lead to bias, and too few lines or points will result in lack of precision. The number of lines or points (k) should be 4, and sampling should be probabilistic to adequately represent the area of inference. We also suggest that transect length be selected to provide a minimum of 40 animals detected, and preferably 60–80 (Buckland et al. 2001).

We recommend those planning to conduct a modern distance sampling study consult Buckland et al. (2001) and, if available, published recommendations for specific field situations or species (e.g., Karanth et al. 2002). For instance, Anderson et al. (2001b) used field trials to estimate the abundance of artificial desert tortoise (*Gopherus agassizii*) models to test whether assumptions that underlie distance sampling were met. They found the density estimate of adult tortoise models was relatively unbiased, whereas the estimate for subadult (small) tortoise models was biased low (about 20%). They attributed the bias to failure to detect small tortoises on or near the centerline and presented ideas to better train observers before commencing the survey. And standard distance theory, based on the premise that detection probability is a decreasing function of distance and that nothing else influences detection, can be violated. Breeden et al. (2008) noted the effects of traffic noise on auditory point surveys of urban white-winged doves (*Zenaida asiatica*).

Distances also can be measured to animals (usually land birds) that are counted around a point rather than along a transect. There are advantages and disadvantages associated with use of points rather than line transects. For example, a line transect can yield more data per unit time than can points, particularly when more time is spent traveling between transects or points than actual sampling (Bibby et al. 2000, Rosenstock et al. 2002). Scale also is important, as a

typical transect generally covers more spatial area than a typical set of points; thus, the scale of spatial habitat diversity must be commensurate with the scale of transects or points. The main disadvantage with points, according to Bibby et al. (2000:92), is the area surveyed is proportional to the square of the distance from the observer, whereas in transects the area is proportional to lateral distance from the transect line. Thus, density estimates from point data are more susceptible to errors arising from inaccurate distance measurements or from violation of assumptions about detecting animals.

However, points are often preferred to transects in habitats with a variety of small patches of habitat relative to the home range of an animal (Bibby et al. 2000). Likewise, points can be preferred over transects when vegetation or terrain hinders navigation, or when observer movement signals the animals of observer presence. For instance, Reynolds et al. (1980) noted that observers traveling along line transects, in structurally complex vegetation and rough terrain, tended to watch the path of travel, reducing their ability to detect birds. Consequently, they recommended establishing equally spaced observer stations, positioned along a transect of points that could be located randomly. Similarly, Koenen et al. (2002) used point transects to estimate seasonal density and group size of mule deer (*Odocoileus hemionus*) by gender and age class on the Buenos Aires National Wildlife Refuge in southeastern Arizona. The authors believed their survey design balanced the often conflicting objectives of random placement of transects and detecting animals before they moved. Burnham et al. (1980) and Buckland et al. (1993, 2001) provide details for sampling designs of point transects.

The **assumptions** of distance sampling are: (1) points or transects are located randomly with respect to the distribution of animals; (2) all objects at the center of the point or transect are detected with certainty; (3) objects are detected at their initial location prior to any movement in response to the observer; (4) distances are measured accurately (ungrouped data), or objects are counted in the proper distance category (grouped data); and (5) objects are detected independently. Violation of the second assumption is a critical failure and is probably common when conducting bird surveys. This violation will result in negatively biased estimates of density. Similarly, if animals are attracted to the observer, the data are not likely to indicate a problem, resulting in positive bias in the estimate of density.

REMOVAL METHODS

Removal methods of population estimation are old and have been analyzed by numerous investigators. Yet these methods are attractive, because often someone other than the investigator, such as hunters, can collect the removal data. Thus, the investigator may not have to actually capture and mark animals to develop population estimates based on removals,

which often makes these methods inexpensive to implement in the field.

Catch-per-Unit-Effort

Catch-per-unit-effort (e.g., catch/day) is based on the premise that as more animals are removed from a population, fewer are available to be “caught,” and catch/day **will decline** (Fig. 11.4). Eventually, if all animals are removed, the expected catch will become zero, and the total number of animals removed will equal the initial population size. Because it is generally not desirable (and seldom possible) to remove all individuals in a population, this method involves developing a linear regression of the number of animals removed each day on the cumulative total number of animals removed prior to that day (Leslie and Davis 1939). An advantage of this method is that population estimates can be derived prior to all animals being removed, and they can be used with removals that are a part of routine management activities, such as hunter or fisher harvests. Animals do not have to be physically taken or removed to be “caught.” Animals can be trapped, shot, photographed, or seen. If animals are marked (i.e., live-trapping of small mammals), they would be included in the calculation on the day they were trapped and marked, but they would be ignored on subsequent days if re-trapped.

Assumptions for this method include: (1) sampling units are taken at random; (2) the population is closed (e.g., the removal period is kept as short as possible); (3) all individual animals have an equal probability of being caught; and (4) unit of effort is constant, and all the removals are known. Catch-per-unit-effort estimates are not likely to be accurate or precise unless a large proportion of the population is removed (i.e., large enough to cause a decline in catch-per-unit-effort; Krebs 1998, Bishir and Lancia 1996).

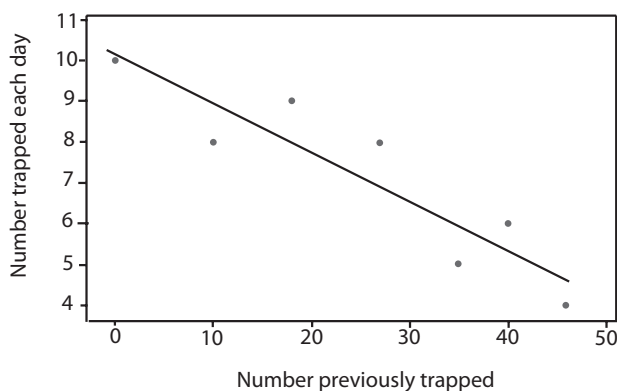


Fig. 11.4. Estimating population abundance by plotting the daily number trapped against the total number previously captured. In this example 10, 8, 9, 8, 5, 6, and 4 mice were trapped on 7 consecutive days. The regression equation for these data is $x = 74.7 - 6.94y$; therefore, when $y = 0$ (all mice are removed), x would equal 75 mice in the population.

The regression equation is not a typical regression, because the catch/day and the cumulative removals depend on the same removals. This lack of independence makes calculation of variances and 95%CI difficult. Bishir and Lancia (1996) have shown that estimates do not follow a normal distribution and, therefore, standard variance equations are not appropriate.

Change-in-Ratio

Kelker (1940) first used this method on selective harvest of male and female deer. Often it is referred to as the **sex-ratio estimator**, because in most cases, sex determines the 2 classes (e.g., male and female deer or pheasants) used with the estimator. However, the method can be used on any 2 classes of animals as long as harvest varies between the classes, such as age classes (e.g., adults and juveniles); species harvested at the same time (e.g., gray [*Sciurus carolinensis*] and fox [*Sciurus niger*] squirrels if one species is selected over the other); if only 1 species is harvested (e.g., deer and cows, where cows are not hunted); or with marked and non-marked populations, where restrictions are placed on harvest of marked animals (e.g., collar-marked deer).

This methods primarily has been used on public hunting areas or on private ranches, where hunts are controlled, and animals taken must come through a check station. In this way, total kill is known. Additional information needed is an estimate of the proportion of each class (e.g., proportion of males and females) in the population just prior to the hunt and an estimate of the proportion of each class after the hunt. **Assumptions** include: (1) the proportion of the classes will change after the hunt due to selective harvest of one class over the other (e.g., more bucks killed than does), (2) observed proportions of the 2 classes are unbiased (a major problem with the estimator to be discussed later), (3) the population is closed, and (4) the number of removals of each class is known. If these assumptions are valid, the population abundance can be estimated using the following equation:

$$N_1 = \frac{[(T)(p_2) - F]}{p_2 - p_1}$$

$$N_2 = N_1 - T,$$

where N_1 = pre-hunt population

T = total kill (all animals harvested regardless of sex class)

F = number of females killed

p_1 = proportion of females in survey before hunt

p_2 = proportion of females in survey after hunt

N_2 = post-hunt population

Example: On a public hunting area, you observed 300 male and 300 female pheasants on a road survey prior to a hunt. During the hunt, 400 male and 25 female pheasants (females are illegal to shoot) are shot and brought through a check sta-

tion. On a road survey immediately after the hunt, you observe 100 male and 300 female pheasants. The estimated population abundance prior to the hunt would be calculated as

$$N_1 = \frac{[(425)(0.75) - 25]}{0.75 - 0.50} = 1,175.$$

Therefore, **post-hunt** population is

$$N_2 = 1,175 - 425 = 750.$$

Comment: Note that a small change in p_1 or p_2 will have a great effect on the estimate. As noted for assumption 2 above, observed proportions of the 2 classes should be unbiased. We believe this is a major problem with the estimator, and we discuss this issue in some detail here. Prior to a hunt, animals have probably not been hunted for at least a year, or in the case of released pheasants, not at all. Therefore, sighting probabilities for male and females may be unbiased; however, once hunting centers on 1 sex class, we believe that sex class will have a lower probability of sighting after the hunt, whereas the nonhunted sex will have a higher probability of being sighted. This bias would exist for pheasants or deer and similarly, for different age classes for which larger animals (either trophy deer or larger deer hunted for meat) are harvested more than are young of the year (i.e., fawns). In addition, the probability of sighting different sex and age classes of deer varies by month even if they are not hunted (Downing et al. 1977), thereby giving a bias between pre- and post-hunt observations. We, therefore, do not recommend this method to estimate population abundance, and we know of few people who currently use it. We have presented the method here only because readers may come across this method in the literature and should be aware of its problems.

MARKED-RESIGHT METHODS

Unlike previous editions of the *Techniques* manual, in which this section was usually titled Capture-Recapture or Capture-Mark-Recapture, we prefer the term marked-resight, because animals do not have to be captured to be marked (e.g., they may have natural marks, including DNA, or may be marked remotely with paint-ball guns, etc.; see Chapter 9, This Volume), nor do they need to be recaptured (e.g., they can be observed; photographed; or DNA fingerprinted from hair, feathers, or feces) to determine whether they are marked. In fact, they do not need to be marked at all (we explain this later). There is **only one assumption** for marked-resight methods: the proportion of marked to nonmarked individuals in a sample is the same as it is in the population. All other purported assumptions are just violations of this assumption. We examine this issue more closely later in this section.

We consider marked-resight methods to be the **gold standard** for conducting population estimates. For if done correctly, we believe they produce more accurate and reli-

able estimates. However, the percentage of marked animals in the population will affect the accuracy of the estimates (Silvy et al. 1977). Silvy et al. (1977) noted that when 50% of the population was marked, more accurate estimates were obtained; however, due to cost of marking animals, they recommended that at least 25% of the population be marked. How does one know when 25% of the animals are marked? When 25% of animals seen on random resight surveys are observed, 25% of the population is marked.

Known Number Alive

Many times when conducting marked-resight studies on small populations (e.g., bobcats [*Lynx rufus*] on small areas), few if any animals are resighted. A common method to estimate abundance is to simply use the number of original captures as an estimate of abundance. **Known-to-be-alive** or **minimum-number-live** estimates are often the most appropriate estimates when conducting these types of studies. These estimates tend to underestimate population density; however, an overestimate of density may lead to inappropriate management action, whereas an underestimate may produce inefficient, but safe management strategies.

Lincoln-Petersen Estimator

A known number of animals in a study is “marked” during a short time period, and then within a few days, a random sample is taken to determine the number marked in the sample. A **rule of thumb** is to use a different method to obtain the sample than was used to mark the animals. For example, do not use a net gun from a helicopter to capture and mark deer and then use a helicopter to obtain the sample, as deer captured and marked may hide from the helicopter, thereby producing a bias in the sample that will cause an overestimation of population abundance. If the assumption given above is valid, an unbiased estimate of population abundance can be obtained using the following equation:

$$N = \frac{Mn}{m},$$

where N = population abundance

M = number marked in study area

n = number of marked and nonmarked animals observed in sample

m = number marked in sample

Example: You mark 100 deer using box traps on a study area, and a week later, you conduct a random road survey and see 50 deer, of which 10 are marked. The estimated population abundance would be calculated as

$$N = \frac{(100)(50)}{10} = 500.$$

Assuming a normal distribution, the 95%CI would be approximately ± 2 standard errors (SE). An estimate for 1 SE can be obtained using

$$SE_N = \sqrt{\frac{(m^2n) + (n-m)}{m^3}}.$$

Therefore,

$$SE_N = \sqrt{\frac{[(100)(100)(50)] + (50-10)}{(10)(10)(10)}} = 141.42$$

$$2 SE = 283;$$

therefore, 95%CI = 217–783 deer. Again, we should replicate the sample several times to obtain a mean estimate, a 95%CI, and conduct a power analysis to determine the sample size needed to detect a desired **effect size**.

Comment: The **only assumption** made in marked–resight methods is the proportion of marked to nonmarked individuals in a sample is the same as it is in the population. If animals lose their marks, then fewer marked animals would be seen than expected, which would cause an overestimation of population abundance. A **rule of thumb** is that any factor (e.g., marked animals leave study area) that causes one to see fewer marked animals than expected will cause an overestimation of population abundance. In contrast, factors (e.g., trap-happy animals) that cause one to see more marked than expected will cause an underestimation of population abundance. There is a premise the Lincoln–Petersen estimator is limited to a closed population. This scenario is best case; however, if the ratio of marked to nonmarked animals leaving a study area is the same as it is in the population, an open population will have **no effect** on the estimate. Similarly, if the same number of nonmarked animals emigrate from and immigrate to the study area, it will have no effect on the estimate. The best way to avoid any problems with population closure is to mark the animals within a short time frame and conduct the resight sample soon thereafter.

Another **misconception** is that animals have to be marked randomly or uniformly throughout the study area. This marking would be ideal; however, if a random resight sample is taken, it does not have to be done, because a random sample should contain the ratio of marked to nonmarked as they are found in the population. To illustrate this idea, we use an extreme example. Say there are 2 identical (e.g., size and vegetation types) islands crossed by a single road with a bridge between them. On the first island, you mark 100 deer, and on the second island you mark none. Later that week you conduct a resight road census over both islands. On the first, you sight 50 deer, of which 10 were marked, and on the second, you sight 50 deer (this result would be expected if the islands were truly identical), of which none were marked. Using the example given above, your estimate for the first island would be 500 deer. Now let us recalculate using all the information from both islands.

Unlike the example above, n now equals 100 (50 seen on each island):

$$N = \frac{(100)(100)}{10} = 1,000$$

The 1,000 deer would be expected if the islands were truly identical. We use this illustration to debunk the idea that marked and nonmarked animals must be evenly mixed. What **must be done** is to obtain a random sample across the study area that will give you a true ratio of marked to nonmarked individuals in the population. For large animals, such as deer, that can more easily be trapped and marked along roads, a resight survey using randomly placed infrared motion-sensitive cameras is ideal, especially if neck collars are used to mark the deer. Also, remember the Lincoln–Petersen estimator does not require that animals be individually marked, making this method ideal when photos may not get a good angle of the marked animal. However, additional information (e.g., movements or survival) can be obtained from animals if they are individually marked, and we recommend that you do so.

At the beginning of this section, we made the comment that no animals need be marked to conduct a marked–resight estimate. In south Texas on large ranches, some landowners stock ranches with known numbers of exotic deer. Given this practice, one could use the number of exotic deer as marked animals and all native deer as nonmarked animals to estimate the number of exotic deer and native deer on the ranch, especially if randomly located infrared cameras were used to resight animals. Subtracting the known number of exotic deer from the estimate would give you an estimate of native deer abundance. The **assumption** is that exotic deer and native deer have the same detection probability. Only one's imagination limits the use of marked–resight methods.

In practice, the major problem we find with marked–resight methods is defining the study area. This is not a problem if working on islands or estimating deer abundance within high fences. But it is a real problem when using live traps in a defined grid to mark–resight small mammals. We recommend using the maximum daily movement (i.e., obtained from maximum distance between traps in the grid in which an individual was trapped on consecutive days) of the mammals in question to define the limits outside the grid. For larger animals (e.g., deer), we also recommend using the maximum daily movement (i.e., obtained from maximum distance between daily sightings of marked animals during resight surveys). This distance is then used to expand an area obtained by including all locations of marked animals within a convex polygon using the minimum number of locations to connect all other locations.

Schnabel Estimator

In situations where animals are continually being marked as resight surveys are conducted, there are several ways to analyze the data for a population estimate. A common way is to

treat each resight survey as a separate data set (i.e., using the total number marked at the time of the survey) to obtain multiple estimates and then calculate a mean estimate of population abundance using the Lincoln–Petersen estimator. Or, because the number of marked animals in the population affects the estimate, one could use only the data obtained from the final survey to obtain an estimate. If the former is used, then one is giving equal weight to each survey and if only the last survey is used because it has a larger sample size, one is not using all data available. To overcome this problem, Schnabel (1938) developed a method (i.e., weighted average) to use all available data without giving each survey an equal weight. The **assumption** for the Schnabel estimator is the same as for the Lincoln–Petersen estimator; namely, the resight sample has the same ratio of marked to nonmarked animals as is found in the population. If the assumption given above is valid, an unbiased estimate of population abundance can be obtained using the following equation:

$$N = \frac{\sum Mn}{\sum m},$$

where N = population abundance

M = number marked in study area

n = number of marked and nonmarked animals
observed in sample

m = number marked in sample

Example: Over a 5-day period, you trapped and marked mice using 100 live traps, with the results shown in Table 11.1. The death of some animals during trapping must be accounted for as noted below. If no animals die, then $A = n$. If animals are found dead in the sample, they must be accounted for (i.e., dead marked animals must then be subtracted from M , and dead nonmarked animals are then not added to M). Using the data from Table 11.1 in the above equation yields

$$N = \frac{1,268}{19} = 66.7.$$

If we had run 4 Lincoln–Petersen estimates for days 2–5, our estimates would be 60 mice for day 2, 57 mice for day 3, 78 mice for day 4, and 68 mice for day 5. If we average these estimates, we get 66 mice with a standard error of 4.70 mice. Assuming a normal distribution, we have a 95%CI (about $\pm 2 SE$) of 57–75 mice. Even though the mean Lincoln–Petersen estimator (66) and Schnabel estimator (67) are similar, the Schnabel estimator gives greater weight to the last days of trapping when a greater number of mice were marked. Silvy et al. (1977) have shown that accuracy of estimates is greater when more animals are marked; therefore, one should use the Schnabel estimator when there are 1 day of resightings.

Schumacher–Eschmeyer Estimator

The Schumacher–Eschmeyer estimator (Schumacher and Eschmeyer 1943) is a variation of the Schnabel estimator, itself a variation of the Lincoln–Petersen estimator. Like the Schnabel estimator, it uses all available data without giving each survey an equal weight. Using the data from the Schnabel example above, 2 additional columns are calculated (Table 11.1). The **assumption** for the Schumacher–Eschmeyer estimator is the same as for the Lincoln–Petersen and Schnabel estimators: the resight sample has the same ratio of marked to nonmarked animals as is found in the population. If the assumption is valid, an unbiased estimate of population abundance can be obtained using the following equation:

$$N = \frac{\sum nM^2}{\sum mM},$$

where N = population abundance

M = number marked in study area

n = number of marked and nonmarked animals
observed in sample

m = number marked in sample

Using data from the Schnabel example above and the last 2 columns of Table 11.1, we obtain

$$N = \frac{42,164}{611} = 69.$$

Table 11.1. Hypothetical example of 5 days of trapping and marking mice with data presented in format suitable for estimation of population abundance using the Schnabel and Schumacher–Eschmeyer estimators^a

Day	Number trapped (n)	Number recaptured (m)	Number alive (A)	Total marked alive prior to date (M)	Mn	nM^2	mM
1	10	0	10	0	0	0	0
2	12	2	11	10	120	1,200	20
3	15	5	15	19	285	5,415	95
4	10	5	9	39	390	15,210	195
5	11	7	11	43	473	20,339	301
Totals		19			1,268	42,164	611

^aNote that only the first 6 columns are needed for the Schnabel estimator, whereas all 8 columns are needed for the Schumacher–Eschmeyer estimator.

Jolly–Seber Estimator

The Jolly–Seber estimator (Jolly 1965, Seber 1965) is used for open populations and estimates population size, survival rates, and births. A marked–resight experiment is conducted, during which, on ≥ 3 successive occasions, animals are marked from the population. The identity of marked individuals is recorded on each occasion, unmarked animals are marked, and all animals are released. An estimate of population size is calculated from the simple relationship that population size is equal to the size of the marked population divided by the proportion of animals marked. Estimates can be obtained for each occasion except the first and last. Calculations for the Jolly–Seber estimator are complicated and are best done with available computer programs; therefore, they are not presented here. Estimates of population size, survival rates, and births can be computed directly by **program JOLLY** (Pollock et al. 1990; <http://www.mbr-pwrc.usgs.gov/software.html>). **Program POPAN-5** (Arnason and Schwarz 1999; <http://www.cs.umanitoba.ca/~popan/>) is based on a different approach to the Jolly–Seber model (Crosbie and Manly 1985, Schwarz and Arnason 1996). It includes estimation of the total number of individuals that are in the population at any time during the study, and from the program computes an estimate of population size (plus survival rate and recruitment). To achieve this estimate, one must make some **assumptions** about the values of parameters at the beginning and end of the study (Schwarz and Arnason 1996).

Assumptions of the Jolly–Seber estimator are: (1) all individuals have equal probability of capture; (2) every marked animal present in the population has the same probability of survival; (3) marks are not lost or overlooked; (4) all samples are instantaneous, and each release is made immediately after the sample; and (5) every animal in the population is equally likely to emigrate, and all emigration from the population is permanent.

COMPUTER SOFTWARE PACKAGES

Several computer software packages are available that can be used to estimate population abundance using the methods described above plus other methods not covered here. Prior to the use of these packages, however, one must be aware that errors may exist in these programs (e.g., early versions of program CAPTURE). If results obtained from a software program appear unrealistic, compare them to results from a different software package. Also, be aware that input errors also can give unrealistic or erroneous results (i.e., “garbage in is garbage out,” and it is not the fault of the software package). Input errors include, but are not limited to (1) data entry, (2) data transfer, (3) column and/or row selection in spreadsheets, and (4) model selection (i.e., assumptions). The best way to test software results is to run a small “known” data set through the soft-

ware program, where the outcome has been previously determined without the use of software programs. If the result obtained is the same or similar, then the data have been entered correctly and the software program is probably working properly.

Program CAPTURE

Program CAPTURE computes tests to select a model from several possible models and then computes estimates of capture probability and population size for **closed** population marked–resight data. However, some models in CAPTURE do not work with small data sets. For those who want to learn more about Program CAPTURE, references are provided in Box 11.1.

Program MARK

Program MARK (<http://warnercnr.colostate.edu/~gwhite/mark/mark.htm>) provides parameter estimates from marked animals when they are re-encountered later. Generalizations using program MARK (White and Burnham 1999) or DOBSERV (<http://www.mbr-pwrc.usgs.gov/software.html>) give the researcher the option to fit generalized Lincoln–Petersen models that allow for probability of detection to depend on covariates, such as species, wind speed, and distance. MARK and DOBSERV use AIC (Burnham and Anderson 1998, 2002) to pick the most parsimonious model that explains the data adequately. Currently, no paper documentation is available for MARK. Electronic documentation can be found at <http://warnercnr.colostate.edu/~gwhite/mark/mark.htm>. This material can be printed if you want hard copy. A reasonably complete description of MARK can be found in White and Burnham (1999). Other references are given in Box 11.1.

Program DISTANCE

Program DISTANCE provides an analysis of **distance sampling data** to estimate density and abundance of a population. Considerably more detail is provided at the web site (<http://www.ruwpa.st-and.ac.uk/distance/>), that includes the software and an electronic manual. The methods used by this program are documented in the references listed in Box 11.1.

SUMMARY

Obtaining precise estimates of animal abundance or density in wild populations is difficult, time consuming, and costly. Most techniques have problems related to estimating the probability of capturing or detecting animals during a survey and to taking insufficient and/or nonrandom samples. When using indices, it is assumed the detection probability is constant, but unknown and that over time nothing changes except population abundance. These assumptions may or may not be true, and we caution against use of indi-

BOX 11.1. REFERENCES FOR POPULATION ESTIMATION COMPUTER PROGRAMS

Program CAPTURE

- Otis, D. L., K. P. Burnham, G. C. White, and D. R. Anderson. 1978. Statistical inference from capture data on closed animal populations. *Wildlife Monographs* 62.
- Rexstad, E., and K. P. Burnham. 1991. Users guide for interactive program CAPTURE. Colorado Cooperative Fish and Wildlife Research Unit, Colorado State University, Fort Collins, USA.
- White, G. C., D. R. Anderson, K. P. Burnham, and D. L. Otis. 1982. Capture–recapture and removal methods for sampling closed populations. LA-8787-NERP, Los Alamos National Laboratory, Los Alamos, New Mexico, USA.

Program MARK

- White, G. C., and K. P. Burnham. 1999. Program MARK: survival estimation from populations of marked animals. *Bird Study (Supplement)* 46:120–138.

Program DISTANCE

- Buckland, S. T., D. R. Anderson, K. P. Burnham, and J. L. Laake. 1993. Distance sampling: estimating abundance of biological populations. Chapman and Hall, New York, New York, USA.

- Buckland, S. T., D. R. Anderson, K. P. Burnham, J. L. Laake, D. L. Borchers, and L. J. Thomas. 2001. Introduction to distance sampling: estimating abundance of biological populations. Oxford University Press, Oxford, England, UK.
- Buckland, S. T., D. R. Anderson, K. P. Burnham, J. L. Laake, D. L. Borchers, and L. J. Thomas, editors. 2004. Advanced distance sampling. Oxford University Press, Oxford, England, UK.
- Laake, J. L., S. T. Buckland, D. R. Anderson, and K. P. Burnham. 1994. DISTANCE user's guide V2.1. Colorado Cooperative Fish and Wildlife Research Unit, Colorado State University, Fort Collins, USA.
- Thomas, L., S. T. Buckland, K. P. Burnham, D. R. Anderson, J. L. Laake, D. L. Borchers, and S. Strindberg. 2002. Distance sampling. Pages 544–552 in A. H. El-Shaarawi and W. W. Piegorsch, editors. *Encyclopedia of environmetrics*. Volume 1. John Wiley and Sons, Chichester, England, UK.
- Thomas, L., J. L. Laake, S. Strindberg, F. F. C. Marques, S. T. Buckland, D. L. Borchers, D. R. Anderson, K. P. Burnham, S. L. Hedley, and J. H. Pollard. 2002. DISTANCE 4.0., Release 1. Research Unit for Wildlife Population Assessment, University of St. Andrews, Scotland, UK. <http://www.ruwpa.st-and.ac.uk/distance/>.

ces unless these assumptions can be verified for the comparisons being made. In the case of population estimation, techniques range from complete counts, where sampling concerns dominate, to incomplete counts, where detection concerns dominate. Examples of population estimation procedures include multiple observer, removal, and capture–resight methods.

Before conducting a survey to estimate population abundance, determine what information is needed, for what purpose the information will be used, how precise an estimate

is needed, and the time and cost required to conduct the survey. The key to deriving population abundance estimates is to select a method that fits the situation. If necessary, techniques can be adapted to meet a particular need. Generally, a biometrician familiar with population estimation literature should be consulted. However, most biometricians consider a method “better” when it has greater precision than another method, but remember that most of these methods have never been tested for accuracy under field conditions. Great precision does not mean great accuracy.