

Application of Bayesian Methods to Habitat Selection Modeling of the Northern Spotted Owl in California: New Statistical Methods for Wildlife Research¹

Howard B. Stauffer,² Cynthia J. Zabel,³ and Jeffrey R. Dunk³

Abstract

We compared a set of competing logistic regression habitat selection models for Northern Spotted Owls (*Strix occidentalis caurina*) in California. The habitat selection models were estimated, compared, evaluated, and tested using multiple sample datasets collected on federal forestlands in northern California. We used Bayesian methods in interpreting Akaike weights calculated for the estimated models. This approach combines Akaike weights with prior probabilities to provide posterior probabilities for the set of competing models for each dataset. This process can be iterated with multiple sample datasets to calculate a succession of posterior probabilities that provide revised assessments of the relative credibility of the models. The posterior probabilities also provide weights for model averaging. They can be used to measure the importance of the covariates in the models, and they provide the weights for model averaging of the predictive values and estimates of the coefficients of the covariates, along with error. This approach offers a robust solution to modeling habitat associations, providing a more realistic assessment of error and uncertainty in the results. We illustrate these methods with sample datasets for the Northern Spotted Owl in California.

Key words: AIC, Akaike weights, Bayesian methods, habitat selection modeling, model averaging, Northern Spotted Owls, posterior probabilities, prior probabilities.

Introduction

It is now widely accepted practice among wildlife researchers and managers to use the parsimonious a priori model selection and inference strategy advocated by Burnham and Anderson (1998). With this approach to model selection, a relatively small collection of biologically plausible candidate models is selected for analysis prior to data collection. The models are then fitted using a sample dataset, and Akaike's Information Criterion (AIC) is used to compare the models (Akaike 1973). AIC "measures" the error, or Kullback-Liebler "distance," between the estimated model and the dataset. The model with the lowest AIC is the best fitting among the collection of candidate models. Since $AIC = \text{deviance} + 2K$ where K is the number of parameters estimated in the model, AIC "penalizes" a model for the number of parameters and discourages a model from having too many parameters and over-fitting a sample dataset. For application, it is best to use a corrected Akaike's Information Criterion, AIC_c , particularly for small datasets, because it is more accurate (Burnham and Anderson 1998).

Burnham and Anderson (1998) also recommend the use of Akaike weights, that can be calculated from the AIC (or AIC_c) values. The Akaike weights A_i of all candidate models

$$A_i = e^{-\{(AIC_i - \text{minimum}(AIC_j))/2\}} / \sum_k e^{-\{(AIC_k - \text{minimum}(AIC_j))/2\}}$$

sum to 1 ($\sum_i A_i = 1$) and can be interpreted as estimates of the relative likelihoods of the models being the best fitting, among the collection of candidate models. Burnham and Anderson (1998) demonstrate convincingly, using bootstrapping techniques with multiple examples, that this interpretation of Akaike weights is reliable in most instances. Akaike weights therefore estimate the relative credibility of the models and can be interpreted as the probability of models being the best fitting, among the collection of candidate models.

Herein we report on the observation that Akaike weights also have a very interesting Bayesian statistical interpretation useful to the wildlife managers. Akaike weights can be interpreted as the relative likelihoods of the "parameter space" of models and used to calculate a probability distribution that is the posterior distribution (Hilborn and Mangel 1997). Wildlife managers

¹A version of this paper was presented at the **Third International Partners in Flight Conference, March 20-24, 2002, Asilomar Conference Grounds, California.**

²Mathematics Department, Humboldt State University, Arcata, CA 95521. E-mail: hbs2@humboldt.edu.

³U.S.D.A. Forest Service, Pacific Southwest Research Station, Redwood Sciences Laboratory, 1700 Bayview Drive, Arcata, CA 95521, USA.

could therefore begin with a prior distribution providing a measure of the relative credibility of a collection of candidate models. The initial prior could be “non-informative,” giving equal probability to each model, if prior information is unavailable or unreliable. Managers could then revise this assessment, calculating the posterior distribution for the models, based upon the prior distribution and the Akaike weights calculated from a dataset. Using a second dataset, this process could be repeated, using prior probabilities that are the posteriors from the first dataset analysis, to obtain new posterior probabilities. If multiple datasets are available, this process could continue sequentially, providing posteriors that represent updated estimates of the probabilities of the models best fitting the population, among the collection of models. This process allows for the continual refinement and evaluation of a suite of candidate models as new data become available. This “strength of evidence” approach (Burnham and Anderson 1998) to model selection and refinement incorporates new “information” which is weighed by the evidence up to that point. Wildlife managers may have traditionally been inclined to stop using a model when it performed poorly on a new dataset, and replace it with a new model that performs well on the new dataset, a “continual reinvention,” as opposed to the approach we are advocating, continual refinement or adaptation. We do, however, recognize that if all models continually perform poorly, reinvention may be necessary. Nonetheless, the approach we present is consistent with the application of adaptive resource management. We will begin by reviewing the ideas of Bayesian statistical analysis.

Bayesian Statistical Analysis

Bayesian statistical analysis begins by assuming a prior distribution representing the current understanding, or a measure of the relative credibility, about a parameter (Iversen 1984, Hilborn and Mangel 1997, Congdon 2001, Gill 2002). A sample dataset is then analyzed, assuming a model such as a normal, Poisson, or binomial distribution. A likelihood function for the parameter is calculated from the data and the model. The posterior distribution is the product of the likelihood and the prior, scaled to sum to 1. The idea is the following

prior → likelihood → posterior

where posterior = prior•likelihood / $\sum$ prior•likelihood for a discrete parameter. In the case of a continuous parameter, summation “ $\sum$ ” should be replaced with integration “ $\int$ ”. This process is based upon Bayes Theorem that states, for parameter value B and data D, in the discrete case,

$$\Pr(B|D) = \Pr(D|B) \cdot \Pr(B) / \sum \Pr(D|B) \cdot \Pr(B)$$

Here $\Pr(B|D)$ = conditional probability of B, given D, is the posterior probability of B. $\Pr(D|B)$ = conditional probability of D, given B, is the likelihood of D. $\Pr(B)$ = probability of D is the prior probability of B. Therefore Bayes Theorem can be interpreted as

$$\text{posterior}(B) = \frac{\text{likelihood}(D|B) \cdot \text{prior}(B)}{\sum \text{likelihood}(D|B) \cdot \text{prior}(B)}$$

The denominator is a scaling factor ensuring that the posterior probabilities sum to 1. So posterior probabilities are scaled products of prior probabilities and likelihood values, obtained from the data and model.

A collection of models can be interpreted as a categorical “parameter space” with discrete values equal to each of the models. The models themselves are estimated using frequentist estimates for the parameters. The Akaike weights, the relative likelihoods of the models being the best fitting to the dataset, therefore, can be interpreted as likelihoods for the models. They can hence be multiplied times prior probabilities, and scaled, to provide posterior probabilities for the parameter space of models.

An Example with Binary Data

As an example, consider the binary dataset {1,0,1} of three measurements, representing the presence or absence of a wildlife species at sample sites. We will use a simple binomial model to describe this dataset for purposes of illustration

$$B(x;p,n) = k \cdot p^x \cdot (1-p)^{(n-x)}$$

where x = the number of 1's (= 2 for this dataset), n = the number of samples (= 3 for this dataset), and k = the binomial coefficient = $\binom{n}{x} = \frac{n!}{x!(n-x)!}$ is a constant

(=3 for this dataset). The likelihood function for the parameter p , based upon this dataset $D = \{1,0,1\}$ and the model, is proportional to $p \cdot (1-p) \cdot p = p^2 \cdot (1-p)$ (fig. 1a). The classical frequentist maximum likelihood estimator for the parameter p , based upon this dataset, calculates the maximum value estimate at $\hat{p} = 2/3 = 0.33$ ($=x/n$) for this model likelihood function. Alternatively, Bayesian statistical analysis begins with a prior for p , such as a non-informative prior based upon a flat conjugate beta distribution (fig. 1b)

$$BE(p;\alpha,\beta) = BE(p;1,1) \\ \text{with parameters } \alpha = 1 \text{ and } \beta = 1,$$

and calculates a posterior also given by a beta distribution, but with different parameters (fig. 1c)

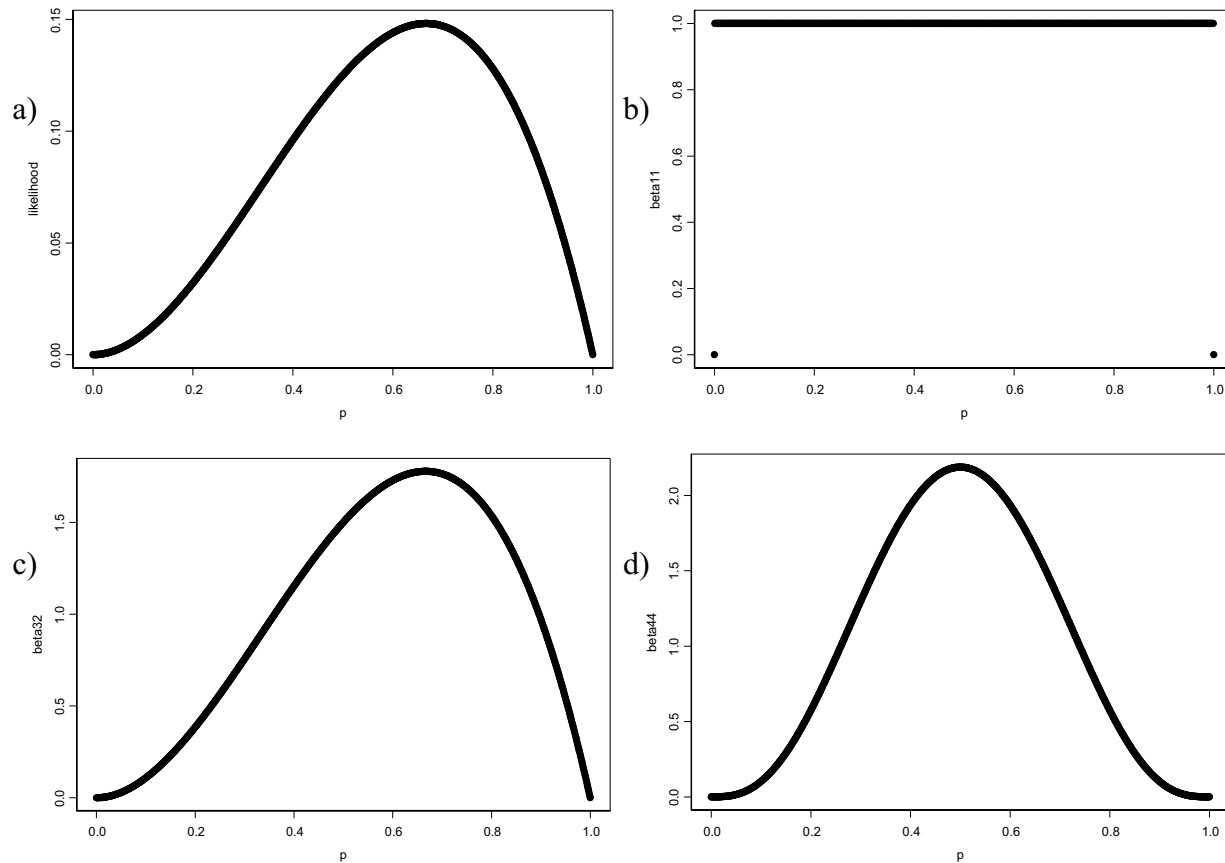


Figure 1— Bayesian statistical analysis for binary datasets, using the binomial model. a) Likelihood $p \cdot (1-p) \cdot p = p^2 - p^3$, for the first dataset $\{1,0,1\}$. b) Non-informative beta prior $BE(p;1,1)$, for the first dataset. c) Beta posterior $BE(p;3,2)$, for the first dataset. d) Beta posterior $BE(p;4,4)$, for the second dataset $\{0,1,0\}$, based upon the $BE(p;3,2)$ prior.

$$BE(p; \alpha+x, \beta+(n-x)) = BE(p; 1+2, 1+(3-2)) \\ = BE(p; 3, 2)$$

Conjugacy theory provides a closed form Bayesian solution for the binomial model $B(x;n,p)$ applied to binary data, with beta prior and posterior distributions, prescribing the addition of x to α and $(n-x)$ to β of the beta prior $BE(p; \alpha, \beta)$ to obtain the beta posterior $BE(p; \alpha + x, \beta + (n-x))$ (Hilborn and Mangel 1997, Carlin and Louis 2000, Gill 2002).

Suppose that we collect a second dataset $\{0,1,0\}$ with $x = 1$, $n = 3$, and likelihood proportional to $(1-p) \cdot p(1-p) = p \cdot (1-p)^2$. The frequentist maximum likelihood estimator for this dataset calculates the maximum value estimate of $\hat{p} = 1/3 = 0.33$. Using the posterior $BE(p;3,2)$ of the first dataset as the prior for this second dataset, we would obtain the new posterior $BE(p;3+1,2+2) = BE(p;4,4)$ (fig.1d). The second posterior would represent our most current assessment about p , based upon the two datasets. Note that the

frequentist maximum likelihood estimate of the combined datasets would be $\hat{p} = 3/6$, occurring at the mode of the second posterior, as we would expect.

A Bayesian Statistical Interpretation of Akaike Weights

Akaike weights estimate the relative likelihoods of the models being the best fitting, among the collection of candidate models. As such, the set of models may be viewed as a “parameter space” with prior probability values assigned to each. The scaled product of the Akaike weights and the priors, scaled to sum to 1, provide a posterior distribution for the parameter space of models, for a given dataset. Furthermore, additional datasets may then be analyzed sequentially, using posteriors obtained from previously analyzed datasets as priors for new datasets, to provide updated new posteriors, based upon all the previously analyzed datasets and the new datasets.

Table 1—Frequentist and empirical Bayesian statistical analysis summaries of 12 habitat selection models for Northern Spotted Owls in California, for RDA, Hayfork, and Mendocino datasets.*a) RDA frequentist analysis.*

Models (200 ha)	AIC _c	Akaike weight	Correct classification (%)	Sensitivity (%)	Specificity (%)
LOGNRE + LOGNRC + NR + NR ² + F + F ²	74.6	0.505505	82.4	87.9	78.1
LOGNR + F + F ²	75.3	0.363419	77.0	93.9	63.4
LOGNRFC + NRF + NRF ² + NRFE	79.4	0.045630	73.0	87.9	61.0
LOGNRE + LOGNRC	79.7	0.041081	73.0	84.9	63.4
LOGNRE	79.7	0.040067	70.3	90.9	53.7
LOGNRFC + NRFE	84.3	0.004037	71.6	90.9	53.1
F + F ²	93.0	0.000052	67.6	84.9	53.7
LOGFEM2	93.0	0.000051	64.9	81.8	51.2
FEME + FEMC	93.3	0.000044	63.5	84.9	46.3
FEM2	93.4	0.000043	63.5	81.8	48.8
LOGFEME + FEMC + FEM + FEM ²	93.5	0.000041	64.9	84.9	48.8
F + F ² + FE	94.2	0.000029	66.2	90.9	46.3

b) RDA empirical Bayesian analysis.

Models	Prior _i	AIC _c	Change in AIC _c	Likelihood	Prior* Likelihood	Posterior _i
FEME + FEMC	0.0833	93.34	18.71	0.00008653	0.00000721	0.00004374
LOGFEME + FEMC + FEM + FEM ²	0.0833	93.46	18.83	0.00008149	0.00000679	0.00004119
LOGNRFC + NRF + NRF ² + NRFE	0.0833	79.44	4.81	0.09026550	0.00751912	0.04562963
LOGNRFC + NRFE	0.0833	84.29	9.66	0.00798652	0.00066528	0.00403722
LOGNRE + LOGNRC	0.0833	79.65	5.02	0.08126824	0.00676964	0.04108147
LOGNRE	0.0833	79.70	5.07	0.07926172	0.00660250	0.04006717
F + F ²	0.0833	92.98	18.35	0.00010360	0.00000863	0.00005237
F + F ² + FE	0.0833	94.15	19.52	0.00005771	0.00000481	0.00002918
LOGNR + F + F ²	0.0833	75.29	0.66	0.71892373	0.05988635	0.36341928
LOGNRE + LOGNRC + NR + NR ² + F + F ²	0.0833	74.63	0.00	1.00000000	0.08330000	0.50550463
LOGFEM2	0.0833	93.04	18.41	0.00010054	0.00000837	0.00005082
FEM2	0.0833	93.36	18.73	0.00008567	0.00000714	0.00004331
				1.97822125	0.16478583	1.00000000

Table 1—contd.

c) Hayfork empirical Bayesian analysis.

Models	Prior ₂	AICc	Change in AICc	Likelihood	Prior* Likelihood	Posterior ₂
FEME + FEMC	0.00004374	211.89	11.48	0.00321477	0.00000014	0.00001408
LOGFEME + FEMC + FEM + FEM ²	0.00004119	256.40	55.99	0.00000000	0.00000000	0.00000000
LOGNRFC + NRF + NRF ² + NRFE	0.04562963	216.44	16.03	0.00033047	0.00001508	0.00150934
LOGNRFC + NRFE	0.00403722	212.20	11.79	0.00275318	0.00001112	0.00111257
LOGNRE + LOGNRC	0.04108147	205.50	5.09	0.07847305	0.00322379	0.32268309
LOGNRE	0.04006717	232.00	31.59	0.00000014	0.00000001	0.00000055
F + F ²	0.00005237	279.22	78.81	0.00000000	0.00000000	0.00000000
F + F ² + FE	0.00002918	281.95	81.54	0.00000000	0.00000000	0.00000000
LOGNR + F + F ²	0.36341928	208.40	7.99	0.01840745	0.00668962	0.66959342
LOGNRE + LOGNRC + NR + NR ² + F + F ²	0.50550463	264.15	63.74	0.00000000	0.00000000	0.00000000
LOGFEM2	0.00005082	200.41	0.00	1.00000000	0.00005082	0.00508691
FEM2	0.00004331	223.57	23.16	0.00000935	0.00000000	0.00000004
				1.10318840	0.00999057	1.00000000

d) Mendocino empirical Bayesian analysis.

Models	Prior ₃	AICc	Change in AICc	Likelihood	Prior* Likelihood	Posterior ₃
FEME + FEMC	0.00001408	39.97	9.21	0.01000170	0.00000014	0.00000502
LOGFEME + FEMC + FEM ²	0.00000000	45.44	14.68	0.00064905	0.00000000	0.00000000
LOGNRFC + NRF ² + NRFE	0.00150934	32.97	2.21	0.33121088	0.00049991	0.01783906
LOGNRFC + NRFE	0.00111257	30.76	0.00	1.00000000	0.00111257	0.03970156
LOGNRE+LOGNRC	0.32268309	38.72	7.96	0.01868564	0.00602954	0.21516183
LOGNRE	0.00000055	40.22	9.46	0.0082647	0.00000000	0.00000017
F + F ²	0.00000000	44.30	13.54	0.00114769	0.00000000	0.00000000
F + F ² +FE	0.00000000	46.24	15.48	0.00043507	0.00000000	0.00000000
LOGNR + F + F ²	0.66959342	37.88	7.12	0.02843882	0.01904245	0.67952257
LOGNRE + LOGNRC + NR + NR ² + F + F ²	0.00000000	65.45	34.69	0.00000003	0.00000000	0.00000000
LOGFEM2	0.00508691	33.43	2.67	0.26315818	0.00133866	0.04776965
FEM2	0.00000004	35.55	4.79	0.09117268	0.00000000	0.00000013
				1.75372622	0.02802328	1.00000000

Application to Northern Spotted Owl Data in California

We applied these ideas to a collection of 12 habitat selection models (Manly et al. 1995) for Northern Spotted Owls in California. For this purpose, we analyzed three datasets, a dataset that was randomly sampled globally for Northern Spotted Owl presence or absence of nesting pairs and habitat attributes throughout federal forestlands in northern California, the so-called Research, Development, and Assessment (RDA) dataset, and two local datasets that were completely censused at the Hayfork Adaptive Management Area and the six adjacent Late Successional Reserves and at Mendocino National Forest. We used the RDA dataset as the initial developmental dataset, estimating, comparing, and testing the models for goodness-of-fit. The Hayfork and Mendocino datasets were then used as test datasets, to further test for goodness-of-fit, but also to “calibrate” our posteriors for the models.

The results of the Bayesian statistical interpretation are summarized in *table 1*. Twelve logistic regression habitat selection models for Northern Spotted Owls in California were estimated and compared (Zabel et al. 2002). The 12 models were selected using habitat definitions thought to be important to the owl, based upon current research findings (Thomas et al. 1990, Gutierrez et al. 1995, Gutierrez et al. 1998, Meyer et al. 1998, Ward et al. 1998, Franklin et al. 1999, Thome et al. 1999, Franklin et al. 2000) and a preliminary analysis screening process. We compared the 12 models using AIC and Akaike weights (Burnham and Anderson 1998) and tested for goodness-of-fit using correct classification (proportion of data points correctly classified), sensitivity (proportion of occupied data points correctly classified), and specificity (proportion of unoccupied data points correctly classified) (Hosmer and Lemeshow 2000). Definitions of owl habitat included NR = nesting, roosting; F = foraging; C = core area; E = edge; FEM = FEMAT definition; and FEM2 = revised FEMAT definition (USDI 1992, USDA and USDI 1993, 1994a, 1994b; Zabel et al. 2002). These 12 models were the best-fitting based upon lowest AIC_c from among a larger collection of several hundred models that were analyzed and compared. These models examined total amount of habitat, amount of core habitat, and edge length of habitat, using linear, quadratic, and pseudo-threshold (log-transformed) forms of the owl habitat definitions.

Table 1a summarizes the analysis results of the developmental dataset, with models ranked by order of fit. The LOGNRE+LOGNRC+NR+NR²+F+F² model was the best fitting, of the 12 candidate models, based upon its lowest AIC_c. This model had the log of NR edge and core, quadratic NR, and quadratic F as its covariates. The LOGNR+F+F² model was second best fit-

ting. This model had the log of NR and quadratic F as its covariates. These top two models dominated the collection of candidate models, as the best-fitting models, with a combined Akaike weight of 0.8689. The goodness-of-fit results were somewhat similar among models, but tended to be higher for the better fitting models. It was interesting to note that the second leading model, although higher in AIC value, did demonstrate a markedly higher sensitivity than the others, at 93.9 percent. High sensitivity was a particularly important management priority for this application. This observation was to become substantiated with further analyses of the other datasets.

Table 1b summarizes the Bayesian interpretation of this RDA dataset analysis. The priors for the models were assumed equal and non-informative, at 1/12 = 0.0833 each. The posteriors for this first dataset were equal to the Akaike weights. The likelihoods in this table were the un-scaled Akaike weights,

$$e^{-\{(AIC-\text{minimum}(AIC))/2\}}.$$

These un-scaled Akaike weights were then multiplied times the priors (assumed constant), and scaled to sum to 1, to obtain the new posteriors, the scaled Akaike weights.

The Bayesian analysis of the Hayfork dataset is summarized in *table 1c*. It was particularly striking with the results of this second step of the Bayesian statistical interpretation that the new posteriors for the local Hayfork dataset were markedly different from the previous RDA posteriors. The second-ranked model, the LOGNR+F+F² model, was now top-ranked, with a posterior probability of 0.6696. The original top ranked model descended to the bottom of the list. A new model, the LOGNRE+LOGNRC model, emerged second on the list, with a posterior of 0.3227. It originally had a posterior of just 0.0411 with the RDA analysis.

Table 1d summarizes the Bayesian interpretation for the Mendocino dataset. Again the original second-ranked model, the LOGNR+F+F² model, had the highest posterior, at 0.6795. The LOGNRE+LOGNRC model was second again, with a posterior of 0.2152. The original top-ranked model still had an insignificant posterior and was at the bottom of the list. As we progressed through additional test datasets (not shown), the results stabilized at these rankings.

Discussion: Use of Posteriors for Model Averaging

The posterior probabilities provide weights for model averaging, based upon the entire collection of candidate models. The predicted response values, the estimat-

ed importance of covariates, and the absolute estimates of covariate coefficients and error can then be calculated using model averaging.

Predictions of the response can be averaged over all the models, using the weighted averages of the predicted values of an observation for all the models. For observation x , the model average predicted response value $f(x)$ for all the models $f_1(x)$, $f_2(x)$, ..., $f_k(x)$ is given by

$$f(x) = \sum_i p_i \cdot f_i(x)$$

where p_i are the model posterior probabilities. If, for example, the $f_1 = \text{LOGNR} + \text{F} + \text{F}^2$ model had a predicted value of $f_1(x) = 0.20$ with posterior of 0.68 (*table 1d*) and the $f_2 = \text{LOGNRE} + \text{LOGNRC}$ had a predicted value of $f_2(x) = 0.30$ with posterior of 0.22, then the model averaged predicted value, based upon these two dominant models, would be

$$\begin{aligned} f(x) &= (p_1/(p_1+p_2)) \cdot f_1(x) + (p_2/(p_1+p_2)) \cdot f_2(x) \\ &= (0.68/0.90) \cdot (0.20) + (0.22/0.90) \cdot (0.30) \\ &= 0.224, \end{aligned}$$

based upon these two dominant models. In actual application, it would be best to take the model averaged predicted value based upon all models in the collection of models under consideration.

The importance of covariates in the models can also be estimated and compared using model averaging. The importance of each covariate x_i can be estimated by taking the sum of the posterior probabilities of the models that include that covariate

$$\text{importance}(x_i) = \sum_{j \in S} p_j$$

where j is indexed with values throughout the set S , the set of models that contain the x_i covariate. In the example above with just the 2 dominant models, $x_1 = \text{NR}$ and its forms are in both models, whereas $x_2 = \text{F}$ is just in the second model, so the relative importance of NR is 0.90 compared to the importance of F of 0.68.

Similarly, model averaging may be used to obtain an absolute or unconditional estimate $\hat{\beta}_i$ of the coefficient of covariate x_i and its error $\text{se}_{\hat{\beta}_i}$, using the weighted average of the individual model estimates $\hat{\beta}_{ij}$ of the coefficient or their standard errors $\text{se}_{\hat{\beta}_{ij}}$ (here i indexes the covariate x_i and j is indexed with values in S , the set of models with covariate x_i)

$$\hat{\beta}_i = \sum_{j \in S} p_j \cdot \hat{\beta}_{ij}$$

$$\text{se}_{\hat{\beta}_i} = \sum_{j \in S} p_j \cdot \text{se}_{\hat{\beta}_{ij}}$$

Burnham and Anderson (1998: 134-137) provide an improved estimator for the unconditional standard error based upon mean square error. Shrinkage estimators can also be obtained by taking the weighted average over all models, not just those in S containing the covariate x_i . An absolute or unconditional estimate of error for a covariate coefficient estimate typically contains more error than is normally calculated conditionally for an individual model estimate that assumes that model is correct. This is a reflection of the added uncertainty of the estimates introduced by variation among the competing models. More details on model averaging, along with examples, can be found in Burnham and Anderson (1998).

Conclusions

A Bayesian interpretation of Akaike weights as likelihoods that can be combined with priors to obtain posteriors for a parameter space of models allows an iterative evaluation process for model comparison, based upon multiple datasets. Datasets need not be analyzed independently of each other, as with a frequentist approach, but rather sequentially, with the results of data analyses building on each other. The posterior distribution outputs of one data analysis becomes the prior distribution inputs of the next data analysis, in a sequential process that reassesses an understanding about parameters and is more attuned to the cumulative scientific method and adaptive management.

Acknowledgments

We thank the U.S.D.A. Forest Service, Pacific Southwest Research Station, Redwood Sciences Laboratory, and Humboldt State University, for providing funding for this project.

Literature Cited

- Akaike, H. 1973. **Information theory and an extension of the maximum likelihood principle**. In: B. N. Petran and F. Csaki, editors. International symposium on information theory. Second edition. Budapest, Hungary: Akademiai Kiado; 267-281.
- Burnham, K. P. and D. R. Anderson. 1998. **Model selection and inference: a practical information theoretic approach**. New York, NY: Springer-Verlag; 353 p.
- Carlin, B. P. and T. A. Louis. 2000. **Bayes and empirical Bayes methods for data analysis, 2nd edition**. Boca Raton, FL: Chapman and Hall.

- Congdon, P. 2001. **Bayesian statistical modeling**. Chichester, England: John Wiley and Sons.
- Franklin, A. B., K. P. Burnham, G. C. White, R. G. Anthony, E. D. Forsman, C. Schwarz, J. D. Nichols, and J. Hines. 1999. **Range-wide status and trends in northern spotted owl populations**. Fort Collins, Colorado, and Corvallis, Oregon: Colorado Cooperative Fish and Wildlife Research Unit, Department of Fishery and Wildlife Biology, Colorado State University; and Oregon Cooperative Fish and Wildlife Research Unit, Department of Fish and Wildlife, Oregon State University; 71 p.
- Franklin, A. B., D. R. Anderson, R. J. Gutierrez, and K. P. Burnham. 2000. **Climate, habitat quality, and fitness in northern spotted owl populations in northwestern California**. Ecological Monographs 70: 539-590.
- Gill, J. 2002. **Bayesian methods: a social and behavioral sciences approach**. Boca Raton, FL: Chapman and Hall.
- Gutierrez, R. J., A. B. Franklin, and W. S. LaHaye. 1995. **Spotted Owl (*Strix occidentalis*)**. In: A. Poole and F. Gill, editors. The Birds of North America, No. 179. Philadelphia, PA and Washington, DC: The Academy of Natural Sciences, and the American Ornithologists' Union.
- Gutierrez, R. J., J. E. Hunter, G. Chavez-Leon, and J. Price. 1998. **Characteristics of spotted owl habitat in landscapes disturbed by timber harvest in northwestern California**. Journal of Raptor Research 32: 104-110.
- Hilborn, R. and M. Mangel. 1997. **The ecological detective**. Princeton, NJ: Princeton University Press.
- Hosmer, D. W. and S. Lemeshow. 2000. **Applied Logistic Regression, 2nd edition**. New York, NY: John Wiley and Sons.
- Iversen, G. R. 1984. **Bayesian statistical inference**. Newbury Park, CA: Sage Publications.
- Manly, B. F. J., L. L. McDonald, and D. L. Thomas. 1995. **Resource selection by animals**. London, UK: Chapman and Hall.
- Meyer, J. S., L. L. Irwin, and M. S. Boyce. 1998. **Influence of habitat abundance and fragmentation on northern spotted owls in western Oregon**. Wildlife Monographs 139: 1-51.
- Thomas, J. W., E. D. Forsman, J. B. Lint, E. C. Meslow, B. R. Noon, and J. Verner. 1990. **A conservation strategy for the Northern Spotted Owl: a report of the interagency scientific committee to address the conservation of the Northern Spotted Owl**. Portland, OR: U.S. Department of Agriculture and U.S. Department of the Interior.
- Thome, D. M., C. J. Zabel, and L. V. Diller. 1999. **Forest stand characteristics and reproduction of northern spotted owls in managed north-coastal California forests**. Journal of Wildlife Management. 63: 44-59.
- Thompson, S. K. 1992. **Sampling**. New York, NY: Wiley.
- USDI. 1992. **Final Draft Recovery Plan for the Northern Spotted Owl, Volume 2**. Portland, OR: Fish and Wildlife Service, U.S. Department of the Interior; 662 p.
- USDA and USDI. 1993. **Forest ecosystem management: an ecological, economic, and social assessment: report of the Forest Ecosystem Management Assessment Team**. Portland, OR: Forest Service, U.S. Department of Agriculture; Fish and Wildlife Service; U.S. Department of the Interior; National Oceanic and Atmospheric Administration, U.S. Department of Commerce; National Park Service, Bureau of Land Management, U.S. Department of the Interior; and Environmental Protection Agency; Irregular pagination.
- USDA and USDI. 1994a. **Record of decision for amendments to Forest Service and Bureau of Land Management planning documents within the range of the northern spotted owl**. Portland, OR: Forest Service, U.S. Department of Agriculture; and Bureau of Land Management, U.S. Department of the Interior; Irregular pagination.
- USDA and USDI. 1994b. **Final supplemental environmental impact statement on management of habitat for late-successional and old-growth forest related species within the range of the northern spotted owl, 2 Volumes**. Portland, OR: Forest Service, U.S. Department of Agriculture; and Bureau of Land Management, U.S. Department of the Interior; Irregular pagination.
- Ward, J. P., Jr., R. J. Gutierrez, and B. R. Noon. 1998. **Habitat selection by northern spotted owls: the consequences of prey selection and distribution**. Condor 100: 79-92.
- Zabel, C. J., L. R. Roberts, B. S. Mulder, H. B. Stauffer, J. R. Dunk, K. Wolcott, D. M. Solis, Jr., M. Gertsch, B. Woodbridge, A. Wright, G. Goldsmith, and C. Keckler. 2002. **A collaborative approach in adaptive management at a large landscape scale**. In: J. M. Scott, P. J. Heglund, J. Haufler, M. Morrison, M. Raphael, B. Wall, editors. Predicting species occurrences: issues of scale and accuracy. Covello, CA: Island Press; 241-253.