

Chapter 10

Predictive Modeling for Classification

10.1 A Brief Overview of Predictive Modeling

Descriptive models, as described in Chapter 9, simply summarize data in convenient ways or in ways which one hopes will lead to increased understanding of the way things work. In contrast, predictive models have the specific aim of allowing one to predict the unknown value of a variable of interest given known values of other variables. Examples include providing a diagnosis for a medical patient on the basis of a set of test results, estimating the probability that a customer will buy product A given a list of other products they have purchased, or predicting the value of the Dow Jones index 6 months from now given current and past values of the index.

In Chapter 6 we discussed many of the basic functional forms of models which can be used for prediction. In this chapter and the next (Chapter 11) we examine such models in more detail, and look at some of the specific aspects of the criteria and algorithms which permit such models to be fitted to the data.

Predictive modeling can be thought of as learning a mapping from an input set of vector measurements $\mathbf{x}$ to a scalar output y (one can learn mappings to vector outputs, but the scalar case is much more common in practice). In predictive modeling the training data D_{train} consists of *pairs* of measurements, a vector $\mathbf{x}(i)$ with a corresponding “target” value $y(i)$. Thus, the goal of predictive modeling is to estimate (from the training data) a mapping or a function $y = f(\mathbf{x}; \theta)$ which can predict a value y given an input vector of measured values $\mathbf{x}$ and a set of estimated parameters θ for the model f . Recall that f is the functional form of the model structure (Chapter 6), the θ 's are the unknown parameters within f whose values we will determine by minimizing a suitable score function on the data (Chapter 7), and the process of searching for the best θ values is the basis for the actual data mining algorithm (Chapter 8).

Defining a predictive modeling algorithm involves choosing three things: a particular model structure (or a family of model structures), a score function, and an optimization strategy for finding the best parameters and model within the model family.

In data mining problems, since we typically know very little about the functional form of $f(\mathbf{x}; \theta)$ ahead of time, there may be attractions in adopting fairly flexible functional forms

or models for f . On the other hand, as discussed in Chapter 6, simpler models have the advantage of often being more stable and more interpretable, as well as often providing the functional components for more complex model structures. For predictive modeling, the score function is usually relatively straightforward to define, typically a function of the *difference* between the prediction of the model $y(i) = f(\mathbf{x}(i); \theta)$ and the true value $y(i)$, i.e.,

$$S(\theta) = \sum_{D_{\text{train}}} d(y(i), \hat{y}(i)) \quad (10.1)$$

$$= \sum_{D_{\text{train}}} d(y(i), f(\mathbf{x}(i); \theta)) \quad (10.2)$$

where the sum is taken over the tuples $(\mathbf{x}(i), y(i))$ in the training data set D_{train} and the function d defines a scalar distance such as squared-error for real-valued y or an indicator function for categorical y (recall Chapter 7 for further discussion in this context). The actual heart of the data mining algorithm then involves minimizing E as a function of θ ; the details of this are determined by both the nature of the distance function and the functional form of $f(\mathbf{x}, \theta)$ which jointly determine how E depends on θ (again recall the discussion in Chapter 8).

If we are constructing multiple models for prediction we will need to get an honest estimate of how well each model will predict on new out-of-sample data, if we are to fairly assess each model. In this case we can redefine the score function $S(\theta)$ above so that it is estimated on a validation data set, or via cross-validation, rather than on the training data directly (as discussed in Chapter 7 on score functions).

We defer detailed discussion of data access and management issues to Chapter 12 since, to some extent, we can decouple the *concepts* behind the algorithms for predictive modeling from the details of how we access the data when we actually want to implement a predictive algorithm in practice.

There are two important distinct kinds of tasks in predictive modeling depending on whether Y is categorical or real-valued. For categorical Y the task is called *classification* (or *supervised classification* to distinguish it from problems concerned with defining the classes in the first instance, such as cluster analysis) and for real-valued y the task is called *regression*. Classification problems are the focus of this chapter and regression problems are the focus of the next chapter. Although one can legitimately discuss both forms of modeling in the same general context (they share many of the same mathematical and statistical underpinnings), in the interests of organizational style we have assigned classification and regression each to their own chapter. However it is important for the reader to be aware that many of the model structures we discuss in each chapter has a “twin” in terms of being applicable to the other task. For example, we discuss tree structures in the classification chapter but they can also be used for regression. Similarly we discuss neural networks under regression, but they can also be used for classification.

In these two chapters we cover many of the more commonly used approaches to classification and regression problems; that is, the more commonly used tuples of model structures, score functions, and optimization techniques. The natural taxonomy of these algorithms tends to be closely aligned with the model structures being used for prediction (e.g., tree

structures, linear models, polynomials, etc), leading to a division of the chapters largely into subsections according to different model structures. Even though specific combinations of models, score functions, and optimization strategies have become very popular “standard” data mining algorithms the reader should nonetheless keep in mind the general reductionist philosophy of data mining algorithms which we espoused in Chapter 5; for a particular data mining problem one should always be aware of the option of tailoring the model, the score function, or the optimization strategy for the specific application at hand rather than just using an “off-the-shelf” technique.

10.2 Introduction to Classification Modeling

We have already discussed classification modeling in Chapter 6,XXX: here we briefly review some of the basic concepts once again. In classification we wish to learn a mapping from a vector of measurements $\mathbf{x}$ to a categorical variable Y . The variable to be predicted is typically called the *class variable* (for obvious reasons) and for convenience of notation we will use the variable C taking values in the set $\{c_1, \dots, c_m\}$ to denote this class variable for the rest of this chapter (instead of using Y). The observed or measured variables $X_1, \dots, X_p$ are variously referred to as the features, attributes, explanatory variables, input variables, etc; the more generic term “input variable” will be used throughout this chapter. We will refer to $\mathbf{x}$ as a p -dimensional vector (that is, we take it to be comprised of p variables), where each component can be real-valued, ordinal, categorical, and so forth. $x_j(i)$ is the j th component of the i th input vector, where $1 \leq i \leq n$, $1 \leq j \leq p$. In our introductory discussion we will implicitly assume that we are using the so-called “0-1” loss function (see Chapter 7) where a correct prediction incurs a loss of zero and an incorrect class prediction incurs a loss of 1 irrespective of what the true class and the predicted class values are.

We begin by discussing below two different but related general views of classification: the decision boundary (or discriminative) viewpoint, and the probabilistic viewpoint.

10.2.1 Discriminative Classification and Decision Boundaries

In the discriminative framework a classification model $f(\mathbf{x}, \theta)$ takes as input the measurements in the vector $\mathbf{x}$ and produces as output a symbol from the set $\{c_1, \dots, c_m\}$. Consider the nature of the mapping function f for a simple problem with just two real-valued input variables X_1 and X_2 . The mapping in effect produces a piecewise constant surface over the (X_1, X_2) plane, i.e., only in certain regions does the surface take the value c_1 . The union of all such regions where a c_1 is predicted is known as the *decision region* for class c_1 , i.e., if an input $\mathbf{x}(i)$ falls in this region its class will be predicted as c_1 (and the complement of this region is the decision region for all other classes).

Knowing where these decision regions are located in the (X_1, X_2) plane is equivalent to knowing where the *decision boundaries* are between the regions. Thus, we can think of the problem of learning a classification function f as being equivalent to learning decision boundaries between the classes. In this context, we can begin to think of the mathematical forms we can use to describe decision boundaries, e.g., straight lines or planes (linear boundaries),

curved boundaries such as low-order polynomials, and other more exotic functions.

In most real classification problems the classes are not perfectly separable in the $\mathbf{x}$ space. That is, it is possible for members of more than one class to occur at some (perhaps all) values of $\mathbf{x}$ —though the probability that members of each class occur at any given $\mathbf{x}$ will be different. (It is the fact that these probabilities differ which permits one to make a classification. Broadly speaking, one will assign a point at $\mathbf{x}$ to the most probable class.) The fact that the classes “overlap” leads to another way of looking at classification problems. Instead of focussing on decision boundaries and decision surfaces, one can seek a function $f(\mathbf{x}, \theta)$ which maximizes some measure of separation between the classes. Such functions are termed *discriminant functions*. Indeed, the earliest formal approach to classification, *Fisher’s linear discriminant analysis method* (Fisher, 1936), was based on precisely this idea: it sought that linear combination of the variables in $\mathbf{x}$ which maximally discriminated between the (two) classes.

10.2.2 Probabilistic Models for Classification

Let $p(c_k)$ be the probability that a randomly chosen object or individual i comes from class c_k . Then $\sum_k p(c_k) = 1$, assuming that the classes are mutually exclusive and exhaustive (MEF). This may not always be the case, e.g., a person may have more than one disease (classes are not mutually exclusive). In this case, we should model the problem as set of multiple two-class classification problems, i.e., “disease 1 or not,” “disease 2 or not,” etc. Or there may be a disease that is not in our classification model (the set of classes is not exhaustive), in which case we could add an extra class c_{k+1} to the model to account for “all other diseases.” Despite these potential practical complications, unless stated otherwise we will use the MEF assumption throughout this chapter since it is widely applicable in practice and provides the essential basis for probabilistic classification.

As an example, imagine that the classes are male and female and that $p(c_k)$, $k = 1, 2$, represents the probability that at conception a person receives the appropriate chromosomes to develop as male or female. The $p(c_k)$ ’s are thus the probabilities that individual i belongs to class c_k if we have no other information (no measurements $\mathbf{x}(i)$) at all about them. The $p(c_k)$ ’s are sometime referred to as the class “prior probabilities,” since they represent the probabilities of class membership *before* observing the vector $\mathbf{x}$. Note that estimating the p_k ’s from data is often relatively easy: if a random sample of the entire population has been drawn, the maximum likelihood estimate of $p(c_k)$ is just the frequency with which c_k occurs in the training data set. Of course, if other sampling schemes have been adopted, things may be more complicated. For example, in some medical situations it is common to sample equal numbers from each class deliberately, so that the priors have to be estimated by some other means.

Objects or individuals belonging to class i are assumed to have measurement vectors $\mathbf{x}$ distributed according to some distribution or density function $p(\mathbf{x}(i)|c_k, \theta_k)$, where the θ_k ’s are unknown parameters governing the characteristics of class c_k . For example, for multivariate real-valued data, the model structure for the $\mathbf{x}$ ’s for each class might be multivariate Gaussian, and the parameters θ_k would represent the mean (location) and variance (scale) characteristics for each class. If the means are far enough apart, and the variances small

enough, we can hope that the classes are relatively *well-separated* in the input space permitting classification with an error-rate near 0. The interesting general problem here from a data mining perspective of course is that neither the functional form nor the parameters of the distributions of the $\mathbf{x}$ s are known a priori.

Via Bayes theorem we have

$$p(c_k|\mathbf{x}) = \frac{p(\mathbf{x}|c_k, \theta)p(c_k)}{\sum_{i=1}^m p(\mathbf{x}|c_i, \theta)p(c_i)}, \quad 1 \leq k \leq m. \quad (10.3)$$

Note that if we were to know the posterior class probabilities precisely then we can make optimal predictions given a measurement vector $\mathbf{x}$. For example, for the case when all errors incur equal cost one should predict the class value c_k which has the highest posterior probability (is most likely given the data). It is worth noting that this scheme is optimal in the sense that no other prediction method can do better. Of course the difficulty is that in practice one does not know the $p(c_k|\mathbf{x})$ functions, or equivalently, the terms $p(\mathbf{x}|c_k, \theta_k)$ and $p(c_k)$.

The posterior probabilities $p(c_k|\mathbf{x}, \theta_k)$ implicitly carve up the input space $\mathbf{x}$ into m decision regions with corresponding decision boundaries. For example, with two classes ($m = 2$) the decision boundaries will be located along the iso-contours where $p(c_1|\mathbf{x}, \theta_1) = p(c_2|\mathbf{x}, \theta_2)$.

In many real problems (perhaps most) the optimal classification scheme will have a non-zero error rate. This arises from the overlap of the distributions $p(\mathbf{x}|c_k, \theta_k)$, mentioned above. Overlap means that the maximum class probability $p(c_k|\mathbf{x}) < 1$. Thus, there is a non-zero probability $1 - p(c_k|\mathbf{x})$ of data arising from the other (less-likely) classes at $\mathbf{x}$, even though the optimal decision at $\mathbf{x}$ is to choose c_k . Extending this argument over the whole space, and averaging with respect to $\mathbf{x}$ (or summing over discrete-valued variables), the *Bayes Error Rate* is defined as

$$p_B^* = \int (1 - p(c^*|\mathbf{x}))p(\mathbf{x})d\mathbf{x} \quad (10.4)$$

where $p(c^*|\mathbf{x}) = \max_k p(c_k|\mathbf{x})$. This is the minimum possible error rate. No other classifier can achieve a lower expected error rate on unseen new data. In practical terms, the Bayes error is a lower-bound on the best possible classifier for the problem.

Display 10.1

Figure 10.1 shows a simple artificial example with a 1-dimensional variable X (the horizontal axis) and two classes. The upper two plots show how the data are distributed within class 1 and class 2 respectively. Each has a uniform distribution over a different range of X ; class c_1 tends to have lower x values than class c_2 . There is a region along the x axis (between values x_1 and x_2) where both class populations overlap.

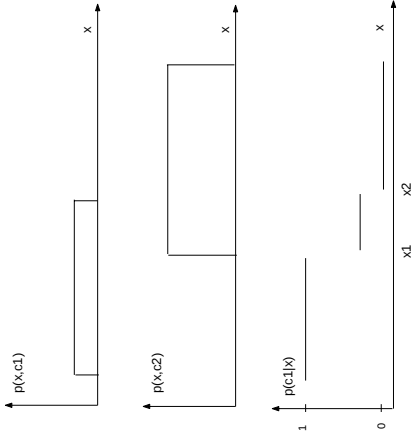
The bottom plot shows the posterior class probability for class c_1 , $p(c_1|x)$ as calculated via Bayes rule given the class distributions shown in the upper 2 plots. For values of $x \leq x_1$, the probability is 1 (since only class c_1 can produce data in that region) and for values of $x \geq x_2$ the probability is 0 (since only class c_2 can produce data in that region). The region of overlap (between x_1 and x_2) has a posterior probability of about $1/3$ for class c_1 (by Bayes rule) since class c_2 is roughly twice as likely as class c_1 in this region. Thus, class c_2 is the Bayes-optimal decision for any $x \geq x_1$ (noting

that in the regions where $p(x, c_1)$ or $p(x, c_2)$ are both zero, the posterior probability is undefined). However, note that between x_1 and x_2 there is some fundamental ambiguity about which class may be present given an x value in this region, i.e., although c_2 is the more likely class there is a $1/3$ chance of c_1 occurring. In fact, since there is a $1/3$ chance of making an incorrect decision in this region, and let us guess from visual inspection that there is a 20% chance of an x value falling in this region, this leads to a rough estimate of a Bayes error rate of about 6% for this particular problem.

Note in Display 10.2.2 that if we had access to a second variable (say Z), and if along the Z axis the classes have no overlap, then by adding this second measurement to the model we would have a zero Bayes error rate, and in principle could try to learn a perfect classifier. This situation occurs in many practical problems, namely, that while one's initial set of variables may have a relatively high Bayes error rate (i.e., the variables are not powerful enough to perfectly discriminate between the classes), one can always in principle reduce the error rate by making more measurements.

This prompts the question of why one not always add lots of measurements in a classification problem, until the error rate is sufficiently low. The answer lies in the the bias-variance principle discussed in Chapter 6. While the Bayes error-rate can only stay the same or decrease if we add more variables to the model, in fact we do not know the optimal classifier or the Bayes error rate. We have to estimate a classification rule from a finite set of training data. Increasing the number of variables for a fixed number of training points means that the training data are representing the underlying distributions less and less accurately. The Bayes error rate may be decreasing, but we have a poorer approximation to it. At some point, as the number of variables increases, the paucity of our approximation overwhelms the reduction in Bayes error rate and the rules begin to deteriorate. The solution is to choose our variables with care: we need variables which when taken together, separate the classes well. Finding appropriate variables (or a small number of features—combinations of variables) is the key to effective classification. This is perhaps especially marked for complex and potentially very high dimensional data such as images, where it is generally acknowledged that finding the appropriate features can have a much greater impact on classification accuracy than the variability which may arise by choosing different classification models.

Figure 10.1: A simple example illustrating posterior class probabilities for a 2-class 1-dimensional classification problem.



While the above framework provides insight from a theoretical viewpoint, it does not provide a prescriptive framework for classification modeling, i.e., it does not tell us specifically how to construct classifiers unless we happen to know precisely the functional form of $p(\mathbf{x}|c_k)$ (which is rare in practice). Nonetheless, there are three general approaches which are suggested by the Bayesian framework and we list them here:

1. **The Discriminative Approach:** Here we try to model the decision boundaries directly, i.e., a direct mapping from inputs $\mathbf{x}$ to one of m class label $c_1, \dots, c_m$. No direct attempt is made to model either the class-conditional or posterior class probabilities. Examples of this approach include perceptrons (Section 10.3) and the more general support vector machines (Section 10.9).

2. The Regression Approach: The posterior class probabilities $p(c_k|\mathbf{x})$ are modeled explicitly, and for prediction the maximum of these probabilities (possibly weighted by a cost function) is chosen. The most widely used technique in this category is known as logistic regression which we discuss in Section 10.7. Note that decision trees (e.g., CART from Chapter 5) can be considered under either the discriminative approach (if the tree only provides the predicted class at each leaf) or the regression approach (if in addition the tree provides the posterior class probability distribution at each leaf).

3. The Class-conditional or “Generative” Approach: The class-conditional distributions $p(\mathbf{x}|c_k, \theta_k)$ are modeled explicitly, and along with estimates of $p(c_k)$ are inverted via Bayes rule (Equation 10.3) to arrive at $p(c_k|\mathbf{x})$ for each class c_k , a maximum is picked (possibly weighted by costs), and so forth, as in the regression approach. We can refer to this a “generative” model in the sense that we are specifying (via $p(\mathbf{x}|c_k, \theta_k)$) precisely how the data are *generated* for each class. Classifiers using this approach are also sometimes referred to as “Bayesian” classifiers because of the use of Bayes theorem. Note, however, that they are not necessarily Bayesian in the sense Bayesian parameter estimation as discussed in Chapter 4—in principle one can use maximum likelihood, MAP, or Bayesian estimation within this general framework.

Note that both the discriminative and regression approaches focus on the differences between the classes (or, more formally, the focus is on the probabilities of class membership conditional on the values of $\mathbf{x}$), whereas the class-conditional/generative approach focuses on the distributions of $\mathbf{x}$ for the classes. Methods which focus on the class-conditional membership distributions are sometimes referred to as *diagnostic* methods, while methods which focus on the distribution of the $\mathbf{x}$ values are termed *sampling* methods. Note also that the class-conditional/generative approach is related to the regression approach in that the former ultimately produces posterior class probabilities, but calculates them in a very specific manner (i.e., via Bayes rule), whereas the regression approach is unconstrained in terms of how the posterior probabilities are modeled. Similarly, both the regression and class-conditional/generative approaches implicitly contain decision boundaries, i.e., in “decision mode” they map inputs $\mathbf{x}$ to one of m classes; however, each does so within a probabilistic framework, while the “true” discriminative classifier is not constrained to do so.

We will discuss examples of each of these approaches in the sections which follow. Which type of classifier works best in practice will depend on the nature of the problem at hand. For some applications (such as in medical diagnosis) it may be quite useful for the classifier to generate posterior class probabilities rather than just class labels. Methods based on the class-conditional distributions also have the advantage of providing a full description for each class (for example this provides a natural way to detect outliers, e.g., future inputs $\mathbf{x}$ which do not appear to belong to any of the known classes). However, as discussed in Chapter 9, it may be quite difficult (if not impossible) to accurately estimate density functions (such as $p(\mathbf{x}|c_k, \theta_k)$) in high-dimensions. In such situations the discriminative classifier may work better. In general, methods based on the class-conditional distributions will require fitting the most parameters (and thus will lead to the most complex modeling), the regression approach will require fewer, and the discriminative model fewest of all. Intuitively this makes sense, since the optimal discriminative model contains only a subset of the information of the

optimal regression model (the boundaries, rather than the full class probability surfaces), and the optimal regression model contains less information than optimal class-conditional distribution model.

10.3 The Perceptron

One of the earliest examples of an automatic computer-based classification rule was the perceptron. The perceptron is an example of a discriminative rule, in that it focuses directly on learning the decision boundary surface. The perceptron model was originally motivated as a very simple artificial neural network model for the “accumulate and fire” threshold behavior of real neurons in our brain—in Chapter 11 on regression models we will discuss more general and recent neural network models.

In its simplest form, the perceptron model (for two classes) is just a linear combination of the measurements in $\mathbf{x}$. Thus, define $h(\mathbf{x}) = \sum w_j x_j$, where the $w_j, 1 \leq j \leq p$ are the weights (parameters) of the model. One usually adds an additional input with constant value 1 to allow for an additional trainable offset term in the operation of the model. Classification is achieved by comparing $h(\mathbf{x})$ with a threshold, which we shall here take to be zero for simplicity. If all class 1 points have $h(\mathbf{x}) > 0$ and all class 2 points have $h(\mathbf{x}) < 0$ then we have perfect separation between the classes. We can try to achieve this by seeking a set of weights such that the above conditions are satisfied for all the points in the training set, i.e., the score function is the number of misclassification errors on the training data for a given set of weights $w_1, \dots, w_{p+1}$. Things are simplified if we transform the measurements of our class 2 points, replacing all the x_j by $-x_j$. Now we simply need a set of weights for which $h(\mathbf{x}) > 0$ for all the training set points.

The weights w_j are estimated by examining the training points sequentially. We start with an initial set of weights and classify the first training set point. If this is correctly classified, the weights remain unaltered. If it is incorrectly classified, so that $h(\mathbf{x}) < 0$, the weights are updated, so that $h(\mathbf{x})$ is increased. This is easily achieved by adding a multiple of the misclassified vector to the weights. That is, the updating rule is $\mathbf{w} = \mathbf{w} + \lambda \mathbf{x}_i$. Here λ is a small constant. This is repeated for all the data points, cycling through the training set several times if necessary. It is possible to prove that if the two classes are perfectly separable by a linear decision surface, then this algorithm will find a separating surface, provided a sufficiently small value of λ is chosen. The updating algorithm is reminiscent of the gradient descent techniques discussed in Chapter 8, although it is actually not calculating a gradient here but instead gradually reducing the error rate score function.

Of course, other algorithms are possible, and others are, indeed, more attractive if the two classes are not perfectly linearly separable (as is often the case in practice). In such cases, the misclassification error rate is rather difficult to deal with analytically (since it is not a smooth function of the weights) and the squared error score function is often used instead, i.e.,

$$S(\mathbf{w}) = \sum_{i=1}^n \left(\sum_j w_j x_{ji} - y(i) \right)^2. \quad (10.5)$$

Since this is a quadratic error function it has a single global minimum as a function of the

weight vector $\mathbf{w}$ and is relatively straightforward to minimize (either by a local gradient descent rule as in Chapter 8, or more directly in closed-form using linear algebra).

10.4 Linear Discriminants

The linear discriminant approach to classification is based on the simple but useful concept of searching for the linear combination of the variables which best separates the classes. Again, it can be regarded as an example of a discriminative approach, since it does not explicitly estimate either the posterior probabilities of class membership or the class-conditional distributions. Fisher (1936) presents one of the earliest treatments of linear discriminant analysis (for the two class case). Let $\hat{\mathbf{C}}$ be the pooled sample covariance matrix defined as

$$\hat{\mathbf{C}} = \frac{1}{n_1 + n_2} (n_1 \hat{\mathbf{C}}_1 + n_2 \hat{\mathbf{C}}_2) \quad (10.6)$$

where n_i is the number of training data points per class, and $\hat{\mathbf{C}}_i$ are the $p \times p$ sample (estimated) covariance matrices for each class, $1 \leq i \leq 2$ (as defined in Chapter 2). To capture the notion of separability along any p -dimensional vector $\mathbf{w}$ he defined a scalar score function as follows:

$$S(\mathbf{w}) = \frac{\mathbf{w}^T \hat{\mu}_1 - \mathbf{w}^T \hat{\mu}_2}{\mathbf{w}^T \hat{\mathbf{C}} \mathbf{w}}. \quad (10.7)$$

The top term is the difference in projected means for each class, which we wish to maximize. The denominator is the estimated pooled variance of the projected data along direction $\mathbf{w}$ and takes into account the fact that the different variables x_j can have different individual variances as well having covariance with each other.

Given the score function $S(\mathbf{w})$, the problem is to determine the direction $\mathbf{w}$ which maximizes this expression. Because this is a linear problem quadratic in the components of $\mathbf{w}$, there exists a closed form solution for the maximizing $\mathbf{w}$, namely:

$$\mathbf{w}_{lda} = \hat{\mathbf{C}}^{-1} (\hat{\mu}_1 - \hat{\mu}_2). \quad (10.8)$$

A new point is classified by projecting it onto the maximally separating direction, and classifying it $\mathbf{x}$ to class 1 if

$$\mathbf{w}_{lda}^T \left(\mathbf{x} - \frac{1}{2} (\hat{\mu}_1 + \hat{\mu}_2) \right) > \log \frac{p(c_1)}{p(c_2)} \quad (10.9)$$

where $p(c_1)$ and $p(c_2)$ are the respective class probabilities.

If one assumes that the distributions within each class have a multivariate normal distribution with a common covariance matrix, then the above method yields the optimal classification rule as in Equation 10.3 (and, indeed, it is optimal whenever the two classes have ellipsoidal distributions with equal quadratic forms). Note, however, that since $\mathbf{w}_{lda}$ was determined without assuming normality, the linear discriminant methodology can often provide a useful classifier even when normality does not hold. Note also that if one approaches the linear discriminant analysis method from the perspective of assumed forms for

the underlying distributions, the method might be more appropriately viewed as based on the class-conditional distribution approach, rather than the discriminative approach.

A variety of extensions to Fisher's original linear discriminant model (described above) model have been developed. *Canonical discriminant functions* generate $m - 1$ different decision boundaries (assuming $m - 1 < p$) to handle the case where the number of classes $m > 2$. *Quadratic discriminant functions* lead to quadratic decision boundaries in the input space when the assumption that the covariance matrices are equal is relaxed. *Regularized discriminant analysis* shrinks the quadratic method towards a simpler form.

The computational complexity of the linear discriminant model scales as $O(p^2n)$ for $m = 2$ classes and $O(m p^2 n)$ for $m > 2$. Here we are assuming that $n \gg \{p, m\}$ so that the main cost is in estimating the class covariance matrices $\hat{\mathbf{C}}_i$, $1 \leq i \leq m$. All of these matrices can be found with at most two linear scans of the database (one to get the means and one to generate the $O(p^2)$ covariance matrix terms). The method thus scales well to large numbers of observations, but is not very useful for large number of variables, as the dependence on the number p of variables is quadratic.

10.5 Tree models

The basic principle of tree models is to recursively partition the space spanned by the input variables to maximize a score of class purity—meaning (roughly, depending on the particular score chosen) that the majority of points in each cell of the partition belong to one class.

Thus, for example, with three input variables, x , y , and z , one might split x at a value t_{1x} , so that the input space is divided into two domains. Each of these domains is then itself split into two, perhaps again at some threshold on x or perhaps at some threshold on y or z . This process is repeated as many times as necessary (see below), with each branch point defining a node of a tree. To predict the class value for a new case with known values of input variables, one works down the tree, at each node choosing the appropriate branch by comparing the new case with the threshold value of the variable for that node.

Tree models have been around for a very long time, although formal methods of building them are a relatively recent innovation. Before the development of such methods they were constructed on the basis of prior human understanding of the underlying processes and phenomena generating the data.

Tree models have many attractive properties. They are easy to understand and explain. They can handle mixed variables (continuous and discrete, for example) with ease since, in their simplest form, trees partition the space using binary tests (thresholds on real variables and subset membership tests on categorical variables). They can predict the class value for a new case very quickly. They are also very flexible, so that they can provide a powerful predictive tool. Having said that, their essentially sequential nature, which is reflected in the way they are constructed, can sometimes lead to suboptimal partitions of the space of input variables.

The basic strategy for building tree models is simplicity itself. One simply recursively splits the cells of the space of input variables. To split a given cell (equivalently, to choose the variable and threshold on which to split the node) one simply searches over all variables and

all possible thresholds to find that which leads to greatest improvement in a specified score function. The score is assessed on the basis of the training data set elements. If the aim is to predict to which one of two classes an object belongs, one chooses the variable and threshold which leads to the greatest average improvement to the local score (averaged across the two child nodes). Splitting a node cannot lead to a deterioration in the score function on the training data. Interestingly enough, for classification, it turns out that using classification error directly is not a useful score function for selecting variables to split on. Other more indirect measures such as entropy have been found to be much more useful. Note that, for ordered variables, a binary split simply corresponds to a single threshold on the variable values. For categorical variables, a split corresponds to partitioning the variable values into two subsets of values.

Display 10.2 *The entropy criterion for a particular real-valued threshold test T (where T stands for a threshold test $X_j > t$ on one of the variables) is defined as the average entropy after the test is performed:*

$$H(C|T) = p(T=0)H(C|T=0) + p(T=1)H(C|T=1) \quad (10.10)$$

where the conditional entropy $H(C|T=1)$ is defined as $-\sum_{c_k} p(c_k|T=1) \log_2 p(c_k|T=1)$. The average entropy above is then the uncertainty from each branch ($T=1$ or $T=0$) averaged over the probability of going down each branch. Since we are trying to split the data into subsets where as many of the data points belong to one class or the other, this is directly equivalent to minimizing the entropy in each branch. In practice, we search among all variables (and all tests or thresholds on each variable) for the single test T which results in minimum average entropy after the binary split.

In principle, this splitting procedure can be continued until each leaf node contained a single training data point—or, in the case when some training data points have identical vectors of input variables (which can happen if the input variables are categorical) continuing until each leaf node contains only training data points with identical input variable values. However, this can lead to severe overfitting. Better trees (in the sense that they lead to better predictions on new data drawn from the same distributions) can typically be obtained by not going to such an extreme (i.e., by constructing smaller more parsimonious trees).

Early work sought to achieve this by stopping the growing process before the extreme had been reached (this is analogous to avoiding overfitting in neural networks by terminating the convergence procedure, as we will discuss in the next Chapter). However, this approach suffers from a consequence of the sequential nature of the procedure. It is possible that the best improvement which can be made at the next step is only very small, so that growth stops, while the step *after* this could lead to substantial improvement in performance. The ‘poor’ step was necessary to set things up so that the next step could take advantage of it. There is nothing specific to trees about this, of course. It is a general disadvantage of sequential methods; precisely the same applies to the stepwise regression search algorithms to be discussed in Chapter 11—which is why more sophisticated methods involving stepping forwards and backwards have been developed. For tree methods similar algorithms have evolved.

A common strategy is to build a large tree—to continue splitting until some termination criterion has been reached in each leaf (e.g. the points in a node all belong to one class or all have the same $\mathbf{x}$ vector)—and then to prune it back. At each step the two leaf nodes are merged which lead to least reduction in predictive performance on the training set. Alternatively, measures such as minimum description length or cross-validation (e.g., the CART algorithm of Chapter 5) are used to trade off goodness of fit to the training data against model complexity. Two other strategies for avoiding the problem of overfitting the training set are fairly widely used. The first is averaging the predictions obtained by the leaves and the nodes leading to the leaves. The second, which has attracted much attention recently, is to base predictions on the averages of several trees, each one constructed by slightly perturbing the data in some way. Such *model averaging methods* are, in fact, generally suitable for all predictive modeling situations.

Often the most common class value among the training data points at a given leaf node (the majority class) is declared as the predicted label for any data points which arrive at this leaf. In effect the region in the input space defined by the branch leading to this node is assigned the label of the most likely class in the region. Sometimes useful information is contained in the overall probability distribution of the classes in the training data at a given leaf. Note that for any particular class, the tree model produces probabilities which are in effect piecewise-constant in the input space and, thus, small changes in the value of an input variable could send a data point down different branches (into a different leaf or region) with dramatically different class probabilities.

When seeking the next best split while building a large tree prior to pruning, the algorithm searches through all variables and all possible splits on those variables. For real-valued variables the number of possible positions for splits is typically taken to be $n' - 1$ (i.e., one less than the number of data points n' at each node), each possible position being located halfway between two data points (putting them half-way between is not necessarily optimal, but has the virtue of simplicity). The computational complexity of finding the best splits among p real-valued variables will typically scale as $O(pn' \log n')$ if carried out in a direct manner. The $n' \log n'$ term results from having to sort the variable values at the node in order to calculate the score function: for any threshold one needs to know how many points are above and below that threshold. For many score functions one can show that the optimal threshold for ordered variables must be located between two values of the variable that have different class labels. This fact can be used to speed up the search, particularly for large numbers of data points. In addition, various bookkeeping efficiencies can be taken advantage of to avoid resorting as one proceeds from node to node. For categorical-valued variables, some form of combinatorial search must be conducted to find the best subset of variable-values for defining a split.

From a database viewpoint, tree-growing can be an expensive procedure. If the number of data points at a node exceeds the capacity of main memory, then the function must operate with a cache of data in main memory and the rest in secondary memory. A brute-force implementation will result in linear scans of the database for each node in the tree, resulting in a potentially very slow algorithm. Thus, when using tree algorithms with data which exceeds the capacity of main memory, one typically either uses clever tree algorithms whose data management strategy is tailored to try to minimize secondary memory access or one

resorts to working with a random sample which can fit in main memory.

Consider a node in the decision tree. It represents some subset D_{node} of the data. The decision on which variable to use at a node of the tree can be made on the basis of knowing for each variable x and each value of x occurring in D_{node} , how many of those values are associated with each value of the class variable y . Call this table $T(node, x)$. Given the data set D_{node} , we can construct the table $T(node, x)$ for all variables x in one pass through the data set. On the basis of these tables, we can decide the optimal splitting criterion. Once that has been decided, the data set can be split into two parts in one pass through the data.

Consider then one level of the tree, i.e., all the nodes in the tree that are on the same level. If the tables $T(node, x)$ for all nodes on this level and all variables x fit into main memory, then we can actually process the whole level in two passes through the data. Thus in this case, the decision tree can be built in a number of passes that is proportional to the height of the decision tree. This is sufficient for most applications.

One disadvantage of the basic form of tree is that it is *monothetic*. Each node is split on just one variable. Sometimes, in real problems, the class variable changes most rapidly with a combination of input variables. For example, in a classification problem involving two input variables, it might be that one class is characterized by having low values on both variables while the other has high values on both variables. The decision surface for such a problem would lie diagonally in the input variable space. Standard methods would try to achieve this by multiple splits, ending up with a staircase-like approximation to this diagonal decision surface. The optimum, of course, would be achieved by using a threshold defined on a linear combination of the input variables - and some extensions to tree methods do just this, permitting linear combinations of the raw input variables to be included in the set of possible variables to be split. Of course, this complicates the search process required when building the tree.

10.6 Nearest neighbor methods

At their basic level, nearest neighbor methods are very straightforward: to classify a new object, with input vector $\mathbf{x}$, one simply examines the k closest training data set points to $\mathbf{x}$ and assigns the object to the class which has the majority of points amongst these k , “Close” here, is defined in terms of the p -dimensional input space. Thus one is seeking those objects in the training data which are most similar to the new object, in terms of the input variables, and then classifying the new object into the most heavily represented class amongst these most similar objects.

In theoretical terms, one is taking a small volume of the space of variables, centred at $\mathbf{x}$, and with radius the distance to the k th nearest neighbor. Then the maximum likelihood estimators of the probability that a point in this small volume belongs to each class are given by the proportion of training points in this volume which belong to each class. The k -nearest neighbor method assigns a new point to the class which has the largest estimated probability. It can be seen from this that nearest neighbor methods are essentially what we have termed “regression” methods—they directly estimate the posterior probabilities of class membership.

Of course, this simple outline leaves a lot unsaid. In particular, we must choose a value for k and a metric through which to define ‘close’. The most basic form takes $k = 1$. This makes a rather unstable classifier, and the predictions can be made more consistent by increasing k . However, increasing k means that the training data points now being included are not necessarily very close to the object to be classified. This means that the “small volume” may not be small at all. Since the estimates are estimates of the average probability of belonging to each class in this volume, this may deviate substantially from the value at any particular point within the volume - and this deviation is likely to be larger the larger is the volume. This means that the predicted probability may be biased from the true probability at the point in question. We are back at the ubiquitous issue of the bias/variance trade-off. There is theoretical work on the best choice of k , but since this will depend on the particular structure of the data set, as well as general issues, the best strategy for choosing k seems to be a data-adaptive one: try various values, plotting the performance criterion (the misclassification rate, for example) against k , to find the best. In following this approach, the evaluation must be on a data set independent of the training data (or else the usual problem of over-optimistic results ensues). However, for smaller data sets it would be unwise to reduce the size of the training data set too much by splitting off too large a test set, since the best value of k clearly depends on the number of points in the training data set. A leaving-one-out cross-validated score function is often a useful strategy to follow, particularly for small data sets.

Many applications of nearest neighbor methods adopt a Euclidean metric: if $\mathbf{x}$ is the input vector for the point to be classified, and $\mathbf{y}$ is the input vector for a training set point, then the distance between them is $\sum_j (x_j - y_j)^2$. As discussed in Chapter 2, the problem with this is that it does not provide an explicit measure of the relative importance of the different input variables. One could seek to overcome this by using $\sum_j w_j (x_j - y_j)^2$, where the w_j are weights. This seems more complicated than the Euclidean metric, but the appearance that the Euclidean metric makes of not requiring a choice of weights is illusory. This is easily seen simply by changing the units of measurement of one of the variables before calculating the Euclidean metric. (An exception to this is when all variables are measured in the same units — as, for example, with situations where the same variable is measured on several different occasions—so-called *repeated measures* data.)

In the two class case, an optimal metric would be one defined in terms of the contours of probability of belonging to class 1 (say): $P(1|\mathbf{x})$. Training data points on the same contour as $\mathbf{x}$ have the same probability of belonging to class 1 as does a point at $\mathbf{x}$, so no bias is introduced by including them in the k nearest neighbors. This is true no matter how far from $\mathbf{x}$ they are, provided they are on the contour. In contrast, points close to $\mathbf{x}$ but not on the contour of $P(1|\mathbf{x})$ through $\mathbf{x}$ will have different probabilities of belonging to class 1, so including them amongst the k will tend to introduce bias. Of course, all this is all very well, but we do not know the positions of the contours. If we did, we would not need to undertake the exercise at all. What this means is that, in practice, one estimates approximate contours and bases the metrics on these. Both global (e.g. estimating the classes by multivariate normal distributions) and local (e.g. iterative application of nearest neighbor methods) have been used for finding approximate contours.

Nearest neighbor methods are closely related to the kernel methods for density estimation

which we discussed in Chapter 6. The basic kernel method defines a cell by a fixed bandwidth and calculates the proportion of points within this cell which belong to each class. This means that the denominator in the proportion is a random variable. The basic nearest neighbor method fixes the proportion (at k/n) and lets the “bandwidth” be a random variable. More sophisticated extensions of both methods (for example, smoothly decaying kernel functions, differential weights on the nearest neighbor points according to their distance from $\mathbf{x}$, or choice of bandwidth which varies according to $\mathbf{x}$) often lead to methods which are barely distinguishable in practice.

The nearest neighbor method has several attractive properties. It is very easy to program. Its classification accuracy can be very good, comparing favorably with alternative more exotic methods such as neural networks. It permits very easy application of the *reject option*, in which a decision is deferred if one is not sufficiently confident about the predicted class. Extension to multiple classes is straightforward (though the best choice of metric is not so clear here). Handling missing values (in the vector for the object to be classified) is simplicity itself: one simply works in the subspace of those variables which are present.

From a theoretical perspective, the nearest neighbor method is a valuable tool: as the design sample size increases, so the bias of the estimated probability will decrease. If one can contrive to increase k at a suitable rate (so that the variance of the estimates also decreases), the misclassification rate of a nearest neighbor rule will converge to a value related to the Bayes error rate. For example, the asymptotic nearest neighbor misclassification rate is bounded above by twice the Bayes error rate.

High dimensional applications cause problems for all methods. Essentially such problems have to be overcome by adopting a classification rule which is not so flexible that it overfits the data, given the large opportunity for overfitting provided by the many variables. Parametric models, of superficially restricted form (such as linear methods) often do well in such circumstances. Nearest neighbor methods often do not do well. With large numbers of variables (and not correspondingly large numbers of training data cases) the nearest k points are often quite far in real terms. This means that fairly gross smoothing is induced, smoothing which is not related to the objectives (as with the linear method, which aims to increase between class separability). The consequence is that nearest neighbor methods can perform poorly in problems with many variables.

Also theoretical analyses suggest potential problems for nearest neighbor methods in high dimensions. Under some distributional conditions the ratio of the distance to the closest point and the distance to the most distant point approaches 1 as the number of dimensions grows. Thus the concept of the nearest neighbour becomes more or less meaningless. However, the distributional assumptions needed for this result are relatively strong, and other more realistic assumptions imply that the notion of nearest neighbour is indeed well-defined.

A potential drawback of nearest neighbor methods is that they do not build a model, relying instead on retaining all of the training data set points. If the training data set is large, searching through them to find the k nearest can be a time-consuming process. Methods have been developed for accelerating this search. For example, branch and bound methods can be applied: if it is already known that at least k points lie within a distance d of the point to be classified, then a design set point is not worth considering if it lies within a distance d of a point already known to be further than $2d$ from the point to be classified. This involves

preprocessing the training data set. Other preprocessing methods discard certain training data points. For example, *condensed nearest neighbor* and *reduced nearest neighbor* methods selectively discard set points so that those remaining still correctly classify all other training data points. The *edited nearest neighbor* method discards isolated points from one class which are in dense regions of another class, so smoothing out the empirical decision surface.

An alternative method for scaling up nearest neighbor methods for large data sets in high dimensions is to use clustering to obtain a grouping of the data. The data points are stored on disk according to their membership in clusters. When finding the nearest point for input point $\mathbf{x}$, the clusters nearest to $\mathbf{x}$ are located and search confined to those clusters. With high probability, under fairly broad assumptions, this method produces the true nearest neighbor.

10.7 Logistic Discriminant Analysis

For the two class case, one of the most widely used basic methods of classification based on the regression perspective is *logistic discriminant analysis*. Given a data point $\mathbf{x}$, the estimated probability that it belongs to class 1 is

$$p(1|\mathbf{x}) = \frac{1}{1 + \exp(\beta\mathbf{x})}.$$

Since the probabilities of belonging to the two classes sum to one, by subtraction, the probability of belonging to class 2 is

$$p(1|\mathbf{x}) = \frac{\exp(\beta\mathbf{x})}{1 + \exp(\beta\mathbf{x})}.$$

By inverting this relationship, it is easy to see that the logarithm of the *odds ratio* is a linear function of the x_j . That is,

$$\log \frac{p(\mathbf{x}|1)}{p(\mathbf{x}|2)} = \beta\mathbf{x}.$$

This approach to modeling the posterior probabilities has several attractive properties. For example, if the distributions are multivariate normal with equal covariance matrices then it is the optimal solution. Furthermore, it is also optimal with discrete $\mathbf{x}$ variables if the distributions can be modeled by loglinear models with the same interaction terms. These two optimality properties can combine, to yield an attractive model for mixed variable (i.e., discrete and continuous) types.

The reader may recall that Fisher's linear discriminant analysis method is also optimal for the case of multivariate normal classes with equal covariance matrices. If the data are known to be sampled from such distributions, then Fisher's method is more efficient. This is because it makes explicit use of this information, by modelling the covariance matrix, whereas the logistic method sidesteps this. On the other hand, the more general validity of the logistic method (no real data is ever exactly multivariate normally distributed) means that this is generally preferred to linear discriminant analysis nowadays.

The word ‘nowadays’ at the end of the previous paragraph arises because of the algorithms required to compute the parameters of the two models. As we noted above, the mathematical simplicity of the linear discriminant analysis model means that an explicit solution can be found. This is not the case for logistic discriminant analysis and an iterative estimation procedure must be adopted. The most common such algorithm is a maximum likelihood approach, based on using the likelihood as the score function. This is described in Chapter 11, in the more general context of *generalized linear models*.

10.8 The Naive Bayes Model

In principle, methods based on the class-conditional distributions in which the variables are all discrete are straightforward: one simply estimates the probabilities that an object from each class will fall in each cell of the cross-classification of the discrete variables, and then uses Bayes theorem to produce a classification as described above.

In practice, however, this is often very difficult to implement because of the sheer number of probabilities which must be estimated as discussed in Chapter 9, i.e., $O(k^p)$ for p k -valued variables. For example, with $p = 30$ and binary variables ($k = 2$) we would need to estimate on the order of $2^{30} \approx 10^9$ probabilities. Assuming (as a rule of thumb) that we should have at least 10 data points for every parameter we estimate (where here the parameters in our model are the probabilities specifying the joint distribution), we would need on the order of 10^{10} data points to accurately estimate the required joint distribution. For m classes ($m > 2$) we would need m times this number. Clearly as p grows the situation becomes impractical.

Recall from Chapters 6 and 9 that we can always simplify any joint distribution by making appropriate independence assumptions, essentially approximating a full table of k^p probabilities by products of much smaller tables. At an extreme, we can assume that all the variables are *conditionally independent*, given the classes. That is, that

$$p(\mathbf{x}|\mathbf{c}_i) = p(x_1, \dots, x_p|\mathbf{c}_i) = \prod_{j=1}^p p(x_j|\mathbf{c}_i), \quad 1 \leq i \leq m \quad (10.11)$$

This is sometimes referred to as the *Naive Bayes* or *first-order Bayes* assumption. The approximation allows us to approximate the full conditional distribution requiring $O(k^p)$ probabilities with a product of univariate distributions, requiring in total $O(kp)$ probabilities per class. Thus, the conditional independence model is linear in the number of variables p rather than being exponential. To use the model for classification we simply use the product form for the class-conditional distributions, yielding the *Naive Bayes classifier*.

It is important to realize that the reduction in the number of parameters by using the Naive Bayes model above comes at a cost, namely that we are making a very strong independence assumption. In some cases the conditional independence assumption may be quite reasonable. For example, if the x_j are medical symptoms, and the c_i are different diseases, then it may be reasonable to assume that given that a person has disease c_i that then the probability of any one symptom depends only on the disease c_i and not on the occurrence of any other symptom. In other words we are modeling how symptoms appear given the disease as having no interactions.

However, in many practical cases the conditional independence assumption may not be very realistic. For example, let x_1 and x_2 be measures of annual income and savings total respectively for a group of people, and let c_i represent their creditworthiness, this being divided into two classes—good and bad. Then, even within each class we might expect to observe a dependence between x_1 and x_2 , because it is likely that people who earn more also save more. Assuming that two variables are independent means, in effect, that will treat them as providing two distinct pieces of information, which we can see is not the case in this example.

Although the independence assumption may not be a very realistic model of the probabilities involved, it may still permit relatively accurate classification performance. There are various reasons for this, including the fact that relatively few parameters are estimated means that the variance of the estimates is small; although the resulting probability estimates may be biased, since we are not interested in their absolute values, but only in their relative order this may not matter; often a variable selection process has already been undertaken, in which one of each pair of highly correlated variables has been discarded; the decision surface from the naive Bayes classifier may coincide with that of the optimal classifier.

Apart from the fact that its performance is often surprisingly good, there is another reason for the popularity of this particularly simple form of classifier. Using Bayes theorem, our estimate of the probability that a point with measurement vector $\mathbf{x}$ will belong to the i th class is

$$p(c_i|\mathbf{x}) \propto p(\mathbf{x}|\mathbf{c}_i)p(c_i) \quad (10.12)$$

$$= p(c_i) \prod_{j=1}^p p(x_j|\mathbf{c}_i) \quad (10.13)$$

by conditional independence. Now let us take the log-odds ratio. After some straightforward manipulation we get

$$\log \frac{p(c_1|\mathbf{x})}{p(c_2|\mathbf{x})} = \log \frac{p(c_1)}{p(c_2)} + \sum \log \frac{p(x_j|\mathbf{c}_1)}{p(x_j|\mathbf{c}_2)} \quad (10.14)$$

Thus the log odds that a case belongs to class 1 is given by a simple sum of contributions from the priors and separate contributions from each of the variables. Despite the simplicity of this form, the decisions surfaces can be quite complicated, and are certainly not constrained to be linear, as are the surfaces corresponding to simple weighted sums of the raw variables. The simplicity, parsimony, and interpretability of the naive Bayes model has led to its widespread popularity, particularly in the machine learning literature.

The naive Bayes model can easily be generalized in many different directions. If our measurements x_j are real-valued we can still make the conditional independence assumption, where now we have products of univariate density estimated instead of just distributions. For any real-valued x_j we can estimate $f(x_j|\mathbf{c}_i)$ using any of our favorite density estimation techniques, e.g., parametric models such as a Gaussian density, more flexible models such as a mixture, or a non-parametric estimate such as a kernel density function.

Equally well we can generalize the the model by including *some but not all dependencies* beyond first-order. One can imagine searching for higher-order dependencies to allow for selected “significant” pairwise dependencies in the model (such as $p(x_j, x_k|c_i)$), and then triples, and so forth. In doing so we are in fact building a general graphical model (or belief network) for the conditional distribution $p(\mathbf{x}|c_i)$ (recall Chapter 9, Section ???). However, the conventional wisdom in practice is that such additions to the model often provide only limited improvements in *classification performance* on many data sets, once again underscoring the difference between building accurate density estimators and building good classifiers.

Finally we comment on the computational complexity of the naive Bayes classifier. Since we are just using (in effect) additive models based on simple functions of univariate densities, the complexity scales roughly as pm times the complexity of the estimation for each individual univariate class-dependent densities. For discrete-valued variables, the sufficient statistics are simple counts of the number of data points in each bin, so we can construct a naive Bayes classifier with just a single pass through the data. A single scan is also sufficient for parametric univariate density models of real-valued variables (we just need to collect the sufficient statistics, such as the mean and the variance for the Gaussian). For more complex density models, such as mixtures models, we may need multiple scans to build the model because of the iterative nature of fitting such density functions (as discussed in Chapter 9).

10.9 Other methods

A huge number of predictive methods have been developed in recent years. Many of these have been powerful and flexible methods, in response to the exciting possibilities offered by modern computing power. We have outlined some of these above, showing how they are related. But other methods also exist - in just one chapter of one book it is not feasible to do justice to all of them. Furthermore, development and invention have not finished. Exciting work continues even as we write. Examples of methods which we have not had space to cover are:

- Mixture models and radial basis function approaches approximate each class-conditional distribution by a mixture of simpler distributions (e.g. Gaussians). Even the use of just a few component distributions can lead to a function which is surprisingly effective in modeling the class-conditional distributions
- Feed-forward neural networks (as discussed in Chapter 5 under the backpropagation algorithm) are a generalization of perceptrons. Sometimes they are called *multilayer perceptrons*. They consist of linear combinations of nonlinear transformations of linear combinations of the raw variables. Sometimes there are more levels of linear combinations and nonlinear transformations. The nonlinearity of the transformations (often logistic transformations) permits highly flexible decision surface shapes, so that such models have been very effective for some classification problems. However, their fundamental nonlinearity means that estimation is not straightforward and iterative techniques (such as hill-climbing) must be used. The slowness of the estimation means that such methods may not be relevant with large data sets.

- Projection pursuit methods provide an alternative class to neural networks. They can be shown, mathematically, to be just as powerful, but they have the advantage that the estimation is more straightforward. They again consist of linear combinations of nonlinear transformations of linear combinations of the raw variables. However, whereas neural networks fix the transformations, in projection pursuit they are data driven.
- Just as neural networks emerged from early work on the perceptron, so also did support vector machines. The early perceptron work assumed that the classes were perfectly separable, and then sought a suitable separating hyperplane. The best generalization performance was obtained when the hyperplane was as far from all of the data points as possible. Support vector machines generalize this to more complex surfaces by extending the measurement space, so that it includes transformations of (combinations of) the raw variables. A linear decision surface which perfectly separates the data in this enhanced space is equivalent to a nonlinear decision surface which perfectly separates the data in the original raw measurement space. Practical experience with such methods is still being gained, but estimation can be slow.

It will be seen from the above, that often a flexible model is fitted, which is then smoothed in some way to avoid overfitting (or the two processes occur simultaneously), and hence strike a suitable compromise between bias and variance. This is manifest in weight decay in fitting neural networks, in regularization in discriminant analysis, in the “flatness” of support vector machines, and so on. A rather different strategy, which has proven highly effective in predictive modeling, is to estimate several (or many) models and average their predictions. We mentioned this in the context of tree classifiers, above. This approach clearly has conceptual similarities to Bayesian approaches, which explicitly regard the parameters of a model as being randomly drawn from some distribution, so that a prediction is based on averaging over the values in this distribution. Whereas model averaging has its natural origins in statistics, the similar approach of majority voting amongst classifiers has its natural origins in machine learning. Yet other ways of combining classifiers are also possible: for example, one can regard the output of classifiers as inputs to a higher level classifier. In principle, any type of predictive classification model can be used at each stage. Of course, parameter estimation will generally not be easy.

A question which obviously arises with the model averaging strategy is how to weight the different contributions to the average — how much weight should each individual classifier be accorded? The simplest strategy is to use equal weights, but there seems obvious that there may be advantages to permitting the use of different weights (not least because equal weights are a special case of this more general model). Various strategies have been suggested for finding the weights, including letting them depend on the predictive performance of the individual model and on the relative complexity of the model. The method of *boosting* can also be viewed as a model averaging method. Here a succession of models is built, each one being trained on a data set in which points misclassified by the previous model are given more weight. This has obvious similarities to the basic error correction strategy used in early perceptron algorithms. There is evidence suggesting that this can be a highly effective strategy.

10.10 Evaluating and Comparing Classifiers

This chapter has discussed predictive classification models — models for predicting the likely class membership of a new object, based on a series of measurements on that object. It will be clear that there are many different methods available, so a perfectly reasonable question is which one should one use. Unfortunately, there is no general answer to this question. Choice must depend on features of the problem, the data, and the objectives. One can be aware of the properties of the different methods, and this can help in making a choice, but theoretical properties are not always an effective guide to practical performance (the effectiveness of the independence Bayes model illustrates this). Of course, differences in expected and observed performance serve as a stimulus for further theoretical work, so leading to deeper understanding.

If practical results sometimes confound the state of current understanding, we must resort to practical comparison of performance to guide our choice of method. There has been a huge amount of work on the assessment and evaluation of classification rules. Much of this work has provided an initial test-bed for enhanced understanding in other areas of model building. This section provides a brief introduction to assessing the performance of classification models.

In the above, we have referred to the error rate or misclassification rate of classification models — the proportion of future objects that the rule is likely to incorrectly classify. We defined the Bayes error rate as optimal error rate — that error rate which would result if one's model was based on the true distribution functions of the problem. In practice, of course, these functions must be estimated (or the alternative discriminative or regression approaches used, and their parameters estimated), so that the model is likely to depart from the optimal. In this case, the model has a true or *actual* error rate (which can be no smaller than the Bayes error rate). The true error rate is sometimes called the *conditional* error rate, because it is conditioned on the given training data set.

The true error rate must be estimated. One obvious way to do this would be to reclassify the training data and see what proportion were misclassified. This is the *apparent* or *resubstitution* error rate. Unfortunately, this is likely to underestimate the future proportion misclassified. This is because the predictive model has been built so that it does well, in some sense, on the training data. (It would be perverse, to say the least, deliberately to choose a model which did poorly on the training data!) Since the training data is merely a sample from the distributions in question, it will not perfectly reflect these distributions. This means that our model may well reflect part of the data specific aspects of the training data. Thus, if these data are reclassified, a higher proportion will be correctly classified than would be the case for future data points.

Many ways have been proposed to overcome this difficulty. One straightforward possibility is to estimate future error rate by calculating the proportion misclassified in a new sample — a *test set*. This is perfectly fine — apart from the fact that, if a test set is available, one might more fruitfully use it to make a larger design set — this will permit a more accurate predictive classification model to be constructed. It seems wasteful, deliberately, to ignore part of the data when constructing the model.

Bearing this in mind, various *cross validation* approaches have been suggested, in which

some small portion (say, one tenth) of the data are left out when the rule is constructed, and then the rule is evaluated on the part which was left out. This can be repeated, with different parts of the data being omitted. Important methods based on this principle are

- the *leaving-one-out* method, in which only 1 point is left out at each stage, but each point in turn is left out, so that one ends up with a test set of size equal to that of the entire training set, but where each single point test set is independent of the model it is tested on. Other methods use larger fractions of the data for the test sets (e.g. one-tenth of the entire data set) but these are more biased than the leaving-one-out method as estimates of the future performance of the model based on the entire data set.
- *bootstrap* methods, of which there are a large number. These model the relationship between the unknown true distributions and the sample by the relationship between the sample and a *subsample* of the same size drawn, with replacement, from the sample. In one method, this relationship is used to correct the bias of the resubstitution error rate. Some highly sophisticated variants of bootstrap methods have been developed, and they are the most effective methods known to date.
- *jackknife* methods are also based on leaving one training set element out at a time, but the principle underlying them is entirely different. They were developed before bootstrap methods, but are mathematically equivalent to an approximation to the bootstrap. Essentially they use performance based on a reduced training set, compared to that of the entire training set, to remove the bias in the estimation of the entire training set.

There are many other methods of error rate estimation. The area has been the subject of several review papers — see the further reading section for details.

Error rate simply regards the misclassification of any object as equally serious. However, this is often (some argue, almost always) unrealistic. Often, certain kinds of misclassification are more serious than other kinds. For example, misdiagnosing a patient with a curable but otherwise lethal disease as suffering from some minor illness is more serious than the reverse. In this case, one may want to attach *costs* to the different kinds of misclassification. In place of simple error rate, then, one seeks a model which will minimize overall loss.

These ideas generalize readily enough to the multiple class case. Often it is useful to draw up a *confusion* matrix, a cross-classification of the predicted class against the true class. Each cell of such a matrix can be associated with the cost of making that particular kind of misclassification (or correct classification, in the case of the diagonal of the matrix) so that overall loss can be evaluated.

Unfortunately, costs are often difficult to determine. When this is the case, an alternative strategy is to integrate over all possible values of the ratio of one cost to the other (for the two class case — generalizations are possible for more than two classes). This approach leads to what is known as the *Gini coefficient* of performance. This measure is equivalent to the test statistic used in the Mann-Whitney-Wilcoxon statistical test for comparing two independent samples, and is also equivalent to the area under a *Receiver Operating Characteristic* or

ROC curve (a plot of the estimated proportion of class 1 objects correctly classified as class 1 against the estimated proportion of class 2 objects incorrectly classified as class 1). ROC curves and the areas under them are widely used in some areas of research. They are not without their interpretation problems, however.

Simple performance of classification models is but one aspect of the choice of a method. Another is how well the method matches the data. For example, some methods are better suited to discrete $\mathbf{x}$ variables, and others to continuous $\mathbf{x}$, while others work with either type with equal facility. Missing values, of course, are a potential (and, indeed, ubiquitous) problem with any method. Some methods can handle incomplete data more readily than others. The independence Bayes method, for example, handles such data very easily, whereas Fisher's linear discriminant analysis approach does not. Things are further complicated by the fact that data may be missing for various reasons, and that the reasons can affect the validity of the model built on the incomplete data. The Further Reading section gives references to material discussing such issues.

In general, the assessment of classification models is an important area, and one which has been the subject of a huge amount of study. The Further Reading section below gives pointers to some of this literature.

10.11 Feature Selection for Classification in High Dimensions

An important issue which often confronts data miners in practice is the problem of having too many variables! Simply put, not all variables that are measured are likely to be *necessary* for accurate discrimination and including them in the classification model may in fact lead to a worse model than if they were removed. Consider the simple example of building a system to discriminate between images of male and female faces (a task which humans perform effortlessly and relatively accurately but which is quite challenging for a classification algorithm). The colors of a person's eyes, hair, or skin are hardly likely to be useful in this discriminative context. These are variables which are easy to measure (and indeed are general characteristics of a person's appearance) but carry little information as to the class identity in this particular case.

In most data mining problems it is not so obvious which variables are (or are not) relevant. For example, relating a person's demographic characteristics to their online purchasing behavior may be quite subtle and may not necessarily follow the traditional patterns (consider a hypothetical group of high-income PhD-educated consumers who spend a lot of money on comic books – if they exist a comic-book retailer would like to know!). In data mining we are particularly interested in letting the data speak, which in the context of variable selection means using data-adaptive methods for variable selection (while noting as usual that should useful prior knowledge be available to inform us about which variables are clearly irrelevant to the task, then by all means we should use this information).

We have discussed this problem in a general modeling context back in Chapter 6, where we outlined some general strategies which we briefly review here:

Variable Selection: The idea here is to select the a subset p' of the original p variables.

Of course we don't know in advance what value of p' will work well or which variables should be included, so there is a combinatorially large search space of variable subsets which could be considered. Thus, most approaches rely on some form of heuristic search through the space of variable subsets, often using a greedy approach to add or delete variables one at a time. There are two general approaches here: the first uses a classification model and algorithm which automatically performs variable selection, the classification tree being the best-known example. The second approach is to use the classifier as a "black-box" and to have an external loop (or "wrapper") which systematically adds and subtracts variables to the current subset, each subset being evaluated by how well the classification model performs. Note that the score function (the evaluation) here should not be the in-sample classification error since this typically can only increase as more variables are added (and thus will be maximized when all variables are used in the model). Instead some form of validation set or cross-validation methodology can be used to generate an honest estimate of how each subset will perform when applied to new out-of-sample data for prediction in the real-world. Clearly there are a huge number of possible search strategies one can employ for this general problem (see the discussion in Chapter 8???) and the trade-off between computation and model quality can be quite important. It helps if the classification algorithm (the "inner loop") is computationally efficient (e.g., a naive Bayes classifier), since then it is relatively cheap to explore many different variable combinations.

Variable Transformations: The idea here is to transform the original measurements by some linear or non-linear function via a preprocessing step, typically resulting in a much smaller set of derived variables, and then to build the classifier on this transformed set. Examples of this approach include principal components analysis (where we try to find the directions in the input space which have the highest variance, essentially a data compression technique — see Chapters 3 and 6), projection pursuit (where an algorithm searches for interesting linear projections — Chapters 6 and 11), and related techniques such as factor analysis and independent components analysis. While these techniques can be quite powerful in their own right, they suffer the disadvantage of not necessarily being well-matched to the overall goal of improving classification performance. A case in point is principal component analysis. Figure 10.2 shows an illustrative example where the first principal component direction (the direction in which the data would be projected and potentially used as input to a classifier) is completely orthogonal to the best linear discriminant for the problem, i.e., it is completely in the wrong direction for the classification task! This is not a problem with the principal component methodology per se but simply an illustration of matching an inappropriate technique to the classification task. This is of course a somewhat artificial and pathological example; in practice principal component projections can often be quite useful for classification, but nonetheless it is worth keeping in mind the overall goal, namely to reduce classification error.

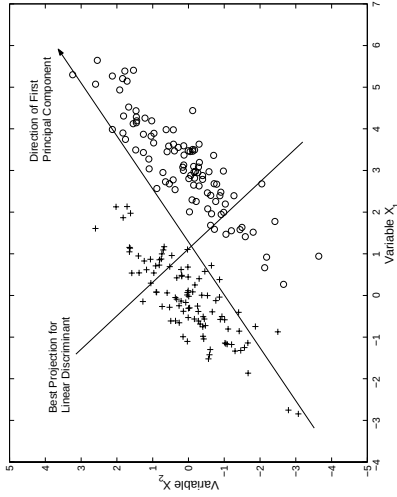


Figure 10.2: An illustration of the potential pitfalls of using principal component analysis as a preprocessor for classification. This is an artificial 2-dimensional classification problem, with data from each class plotted with different symbols. The first principal component direction (which is the first candidate direction on which to project the data if this were actually a high-dimensional problem) is in fact almost completely orthogonal to the best possible linear projection as determined by the linear discriminant technique.

10.12 Further Reading

Fisher's original paper on linear discriminant analysis dates from 1936. Duda and Hart (1973) is a classic text in the statistical pattern recognition literature on classification and related learning problems. Other general reviews of the area are given by Hand (1981), Devijver and Kittler (1982), Fukunaga (1990), McLachlan (1992), Hand (1997), and Webb (1999).

Dasarathy (1991) contains many of the classic papers on nearest neighbor classification from the statistical pattern recognition literature and general descriptions of nearest neighbor methods, including outlines of methods for reducing the size of the retained set, may be found in Hand (1981) and McLachlan (1992). Choice of metric for nearest neighbor methods is discussed in Short and Fukunaga (1981), Fukunaga and Flick (1984), and Myles and Hand (1990). Asymptotic properties of nearest neighbor rules are described in Devroye and Wagner (1982). The related kernel method is discussed in Hand (1981). The problem of meaning of nearest neighbor in high dimensions is considered in Beyer et al (1999) and Bennett et al. (1999), which also discusses the use of clustering for approximate searches.

One of the earliest descriptions of tree-based models is in Morgan and Sunquist (1963). The application of decision trees to classification was popularized in machine learning by Quinlan (1986, 1993). In statistics, the book by Breiman *et al* (1984) describing the CART algorithm (Classification And Regression Trees) was highly influential in the widespread adoption and application of tree models. Chapter 7 of Ripley (1996) contains an extensive overview of the different contributions to the tree-learning literature from statistics, computer science, and engineering. A recent survey article is Murthy (1998). Scalable algorithms for building decision trees are considered in, e.g., Shafer et al. (1996), Gehrke et al (1998), and Rastogi and Shim (1998). The Sprint method of Shafer et al. (1998) operates in a very small amount of main memory but applies only to the CART splitting criterion. The RainForest framework of Gehrke et al. (1998) can be used to scale up a variety of splitting criteria, but its memory usage depends on the sizes of domains of the variables. The method of Rastogi and Shim (1998) interleaves tree building and pruning, thus preventing unnecessary access to the data. A nice survey of scalability issues is Ganti et al. (1999).

Discussions of the independence Bayes method include Russek *et al* (1983), Hilden (1984), Kohavi (1996), Domingos and Pazzani (1997), and Hand and Yu (1999).

Descriptions of support vector machines are given by Vapnik (1995), Burges (1998), and Vapnik (1998).

Techniques for combining classifiers, such as model averaging, are described in Xu *et al* (1992), Ho *et al* (1994), Wolpert (1992), Schaffer (1994), Buntine (1992) and Oliver and Hand (1996).

A detailed review of assessment and evaluation methods for classification algorithms is given in Hand (1997). Reviews of error rate estimation methods in particular are given in Toussaint (1974), Hand (1986), McLachlan (1987), and Schiavo and Hand (1999).

A seminal discussion of missing data, its different types, and how to handle it, is given in Little and Rubin (1987).

References

- Bennett, K., Fayyad, U., and Geiger, D. (1999) Density-based indexing for approximate nearest-neighbor queries. *Proceedings of the Fifth ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, San Diego, CA, ACM Press, pp. 233-243.
- Beyer, K., Goldstein, J., Ramakrishnan, R., and Shaft, U. (1999) When is “Nearest Neighbour” meaningful. *Proc. 7th Int. Conf. Data Theory, ICDT’99*, Lecture Notes in Computer Science, LNCS, Number 1540, pp. 217-235, Springer-Verlag.
- Chipman, H., George, E.I. and McCulloch, R.E. (1998) Bayesian CART model search (with discussion). *J. Am. Statist. Assoc.*, **93**, 935-960.
- Bishop C.M. (1995) *Neural Networks for Pattern Recognition*, Oxford: Clarendon Press.
- Breiman L., Friedman J.H., Olshen R.A., and Stone C.J. (1984) *Classification and Regression Trees*, Belmont, California: Wadsworth.
- Buntine W.L. (1992) Learning in classification trees. *Statistics and Computing*, **2**, 63-73.
- Burges C.J.C. (1998) A tutorial on support vector machines for pattern recognition. *Data Mining and Knowledge Discovery*, **2**, 121-167.
- Dasarathy, B. V. (ed.) (1991) *Nearest Neighbor (NN) Norms: NN Pattern Classification Techniques*, Los Alamitos, CA: IEEE Computer Society Press.
- Devijver P.A. and Kittler J. (1982) *Pattern Recognition: a Statistical Approach*, Englewood Cliffs, New Jersey: Prentice-Hall.
- Devroye L.P. and Wagner T.J. (1982) Nearest neighbor methods in discrimination. In *Handbook of Statistics*, (Vol.2), P.R.Krishnaiah and L.N.Kanal (eds.) Amsterdam: North-Holland, 193-197.
- Domíngos P. and Pazzani M. (1997) On the optimality of the simple Bayesian classifier under zero-one loss. *Machine Learning*, **29**, 103-130.
- Duda R. O. and Hart P. E (1973) *Pattern Recognition and Scene Analysis*, New York: Wiley.
- Fisher, R. A. (1936) The use of multiple measurements on taxonomic problems. *Annals of Eugenics*, **7**, 179-188.
- Fukunaga, K. (1990) *Introduction to Statistical Pattern Recognition*, San Diego: Academic Press.
- Ganti, V., Gehrke, J., and Ramakrishnan, R. (1999), Mining Very Large Databases. *IEEE Computer* 32, 38-45.

- Gehrke, J.E., Ramakrishnan, R., and Ganti, V. (1998) RainForest — a framework for fast decision tree construction of large datasets. *Proceedings of the 24th International Conference on Very Large Databases (VLDB’98)*, pp. 416-427.
- Hand D.J. (1981) *Discrimination and Classification*, Chichester: Wiley.
- Hand D.J. (1982) *Kernel Discriminant Analysis*, Chichester: Wiley.
- Hand D.J. (1986) Recent advances in error rate estimation. *Pattern Recognition Letters*, **4**, 335-346.
- Hand D.J. (1997) *Construction and Assessment of Classification Rules*, Chichester: Wiley.
- Hand D.J., Daly F., Lunn A.D., McConway K.J., and Ostrowski E. (eds.) (1994) *A Handbook of Small Data Sets*, London: Chapman and Hall.
- Hand D.J. and Yu K. (1999) Idiot’s Bayes - not so stupid after all? Working paper, Department of Mathematics, Imperial College, London.
- Hilden J. (1984) Statistical diagnosis based on conditional independence does not require it. *Computers in Biology and Medicine*, **14**, 429-435.
- Ho T.K., Hull J.J., Srihari S.N. (1994) Decision combination in multiple classifier systems. *IEEE Transactions on PAMI*, **16**, 66-75.
- Kohavi R. (1996) Scaling up the accuracy of naive-Bayes classifiers: a decision-tree hybrid. *Proceedings of the Second International Conference on Knowledge Discovery and Data Mining*, Portland, Oregon: AAAI Press, 202-207.
- Little R.J.A. and Rubin D.B. (1987) *Statistical Analysis with Missing Data*, New York: Wiley.
- McLachlan G. J. (1987) Error rate estimation in discriminant analysis: recent advances. In A. K. Gupta (Ed.), *Advances in Multivariate Statistical Analysis*, The Netherlands:Reidel, 233-252.
- McLachlan G.J. (1992) *Discriminant Analysis and Statistical Pattern Recognition*, New York: Wiley.
- Morgan, J. N., and Sonquist, J. A. (1963) Problems in the analysis of survey data, and a proposal. *Journal of the American Statistical Association*, 58, 415-434.
- Murthy, S.K. (1998) Automatic construction of decision trees from data: a multi-disciplinary survey. *Data Mining and Knowledge Discovery* 2, 345-389.
- Myles J.P. and Hand D.J. (1990) The multi-class metric problem in nearest neighbour discrimination rules. *Pattern Recognition*, **23**, 1291-1297.

- Oliver J.J. and Hand D.J. (1996) Averaging over decision trees. *Journal of Classification*, **13**, 281-297.
- Quinlan, J. R. (1986) Induction of decision trees. *Machine Learning*, **1**, 81-106.
- Quinlan, J. R. (1993) *C4.5: Programs for Machine Learning*. San Mateo, CA: Morgan Kaufmann.
- Rastogi, R. and Shim, K. (1998) PUBLIC: A decision tree classifier that integrates building and pruning. *Proceedings of the 24th International Conference on Very Large Databases (VLDB'98)*, pp. 405-415.
- Ripley B.D. (1996) *Pattern Recognition and Neural Networks*. Cambridge: Cambridge University Press.
- Russek E., Krommal R.A., and Fisher L.D. (1983) The effect of assuming independence in applying Bayes' theorem to risk estimation and classification in diagnosis. *Computers and Biomedical Research*, **16**, 537-552.
- Schaffer C. (1994) Cross-validation, stacking and bi-level stacking: meta-methods for classification and learning. In *Selecting Models from Data: AI and Statistics IV*, ed. P.Cheeseman
- Schiavo R.A. and Hand D.J. (1999) Ten more years of error rate estimation. Working paper, Department of Mathematics, Imperial College, London.
- Shafer, J.C., Agrawal, R., and Mehta, M. (1996), SPRINT: A scalable parallel classifier for data mining. *Proceedings of the 22nd International Conference on Very Large Databases (VLDB'96)*, Morgan Kaufmann, pp. 544-555.
- Toussaint G. T. (1974) Bibliography on estimation of misclassification. *IEEE Transactions on Information Theory*, **20**, 472-479.
- Vapnik V. (1995) *The Nature of Statistical Learning Theory*, New York: Springer-Verlag.
- Vapnik V. (1998) *Statistical Learning Theory*, Chichester: Wiley.
- Webb A. (1999) *Statistical Pattern Recognition*, London: Arnold.
- Wolpert D.H. (1992) Stacked generalization. *Neural networks*, **5**, 241-259.
- Xu L., Krzyzak A., and Suen C.Y. (1992) Methods of combining multiple classifiers and their applications to handwriting recognition. *IEEE Pattern Analysis and Machine Intelligence*, **22**, 418-435.