



Exploratory Projection Pursuit

Jerome H. Friedman

Journal of the American Statistical Association, Vol. 82, No. 397. (Mar., 1987), pp. 249-266.

Stable URL:

<http://links.jstor.org/sici?sici=0162-1459%28198703%2982%3A397%3C249%3AEPP%3E2.0.CO%3B2-E>

Journal of the American Statistical Association is currently published by American Statistical Association.

Your use of the JSTOR archive indicates your acceptance of JSTOR's Terms and Conditions of Use, available at <http://www.jstor.org/about/terms.html>. JSTOR's Terms and Conditions of Use provides, in part, that unless you have obtained prior permission, you may not download an entire issue of a journal or multiple copies of articles, and you may use content in the JSTOR archive only for your personal, non-commercial use.

Please contact the publisher regarding any further use of this work. Publisher contact information may be obtained at <http://www.jstor.org/journals/astata.html>.

Each copy of any part of a JSTOR transmission must contain the same copyright notice that appears on the screen or printed page of such transmission.

The JSTOR Archive is a trusted digital repository providing for long-term preservation and access to leading academic journals and scholarly literature from around the world. The Archive is supported by libraries, scholarly societies, publishers, and foundations. It is an initiative of JSTOR, a not-for-profit organization with a mission to help the scholarly community take advantage of advances in technology. For more information regarding JSTOR, please contact support@jstor.org.

Exploratory Projection Pursuit

JEROME H. FRIEDMAN*

A new projection pursuit algorithm for exploring multivariate data is presented that has both statistical and computational advantages over previous methods. A number of practical issues concerning its application are addressed. A connection to multivariate density estimation is established, and its properties are investigated through simulation studies and application to real data. The goal of exploratory projection pursuit is to use the data to find low- (one-, two-, or three-) dimensional projections that provide the most revealing views of the full-dimensional data. With these views the human gift for pattern recognition can be applied to help discover effects that may not have been anticipated in advance. Since linear effects are directly captured by the covariance structure of the variable pairs (which are straightforward to estimate) the emphasis here is on the discovery of nonlinear effects such as clustering or other general nonlinear associations among the variables. Although arbitrary nonlinear effects are impossible to parameterize in full generality, they are easily recognized when presented in a low-dimensional visual representation of the data density. Projection pursuit assigns a numerical index to every projection that is a functional of the projected data density. The intent of this index is to capture the degree of nonlinear structuring present in the projected distribution. The pursuit consists of maximizing this index with respect to the parameters defining the projection. Since it is unlikely that there is only one interesting view of a multivariate data set, this procedure is iterated to find further revealing projections. After each maximizing projection has been found, a transformation is applied to the data that removes the structure present in the solution projection while preserving the multivariate structure that is not captured by it. The projection pursuit algorithm is then applied to these transformed data to find additional views that may yield further insight. This projection pursuit algorithm has potential advantages over other dimensionality reduction methods that are commonly used for data exploration. It focuses directly on the "interestingness" of a projection rather than indirectly through the interpoint distances. This allows it to be unaffected by the scale and (linear) correlational structure of the data, helping it to overcome the "curse of dimensionality" that tends to plague methods based on multidimensional scaling, parametric mapping, cluster analysis, and principal components.

KEY WORDS: Exploratory data analysis; Multivariate analysis; Density estimation; Data mapping; Statistical graphics; Multiparameter numerical optimization.

1. INTRODUCTION

Often, especially during the initial stages, the analysis of a data set is exploratory. One wishes to gain insight and understanding about the nature of the phenomena or system that produced the data without imposing preconceived notions or models. For multivariate data a first set of useful summary statistics is based on the locations and scales of the measurement variables as well as on their correlational structure. Classical multivariate analysis has provided powerful tools for their estimation (see Anderson 1958). If the data closely follow an elliptically symmetric distribution (such as the normal) in the p -dimensional variable space, then these summary statistics usually provide nearly all of the relevant information.

Sometimes, however, important structure in the data is

not adequately captured by the linear associations (correlations) among the variables. Such effects include clustering of the observations into distinct groups and/or concentrations near nonlinear manifolds (which may be made up of parts of linear manifolds). Knowledge of the existence and nature of such effects can sometimes help in understanding the underlying phenomena. In contrast to linear effects, the variety of shapes and other attributes of nonlinearity are immense. It is thus impossible to pre-specify all possibilities in advance and attempt to estimate their corresponding parameters. Powerful approaches can be based on making informative pictorial representations of the data upon which the human gift for pattern recognition can be applied. By viewing (lower-dimensional) representations of the data density, the analyst can often detect striking as well as subtle structure that was impossible to anticipate.

Unfortunately, the human gift for pattern recognition is limited to low dimension. In addition, the technology available to the analyst may place further restrictions on the viewing dimension. Ordinary plotting is limited to two dimensions. Sophisticated uses of motion and color with computer graphics displays can increase the dimensionality for viewing the data to three or perhaps a bit more. If one is to graphically explore multivariate data, it is necessary to find highly revealing lower-dimensional (one-, two-, or three-dimensional) representations.

The most commonly used dimension-reducing transformations are linear projections. This is because they are among the simplest and most interpretable. Moreover, projections are smoothing operations in that structure can be obscured by projection but never enhanced. Any structure seen in a projection is a shadow of an actual (usually sharper) structure in the full dimensionality. In this sense those projections that are the most revealing of the high-dimensional data distribution are those containing the sharpest structure. It is of interest then to pursue such projections.

Friedman and Tukey (1974) presented an algorithm for attempting this goal. The basic idea was to assign a numerical index to every (one- or two-dimensional) projection that characterized the amount of structure present (data density variation) in the projection. This index was then maximized (via numerical optimization) with respect to the parameters defining the projections. They termed this method "projection pursuit." Since all (density) estimation is performed in a univariate (or bivariate) setting this method has the potential of overcoming the "curse of dimensionality" (Bellman 1961) that afflicts such nonlinear methods as parametric mapping (multidimensional scaling) and cluster analysis that are based on interpoint

* Jerome H. Friedman is Professor, Department of Statistics, and Group Leader, Stanford Linear Accelerator Center, Stanford University, Stanford, CA 94305. Helpful discussions with Persi Diaconis, Brad Efron, Iain Johnstone, and Art Owen are gratefully acknowledged. The author thanks John Tukey for reigniting his interest in this subject.

distances. In addition, the projection index can (and often should) be affine invariant (see Huber 1985). Therefore, unlike (projections based on) principal components or factor analysis, it is unaffected—and thus not distracted—by the overall covariance structure of the data, which often has little to do with clustering and other nonlinear effects.

Projection pursuit solutions are seldom unique. Usually data structuring in the full dimensionality will be observable in several lower-dimensional projections and the viewing of each can provide additional insight. It is thus important that a projection pursuit algorithm find as many of these views as possible. Each of these views will (it is hoped) give rise to a substantive local maximum in the projection index. One way to encourage the discovery of several important views is to repeatedly invoke the optimization procedure, each time removing from consideration the solutions previously found. Structure removal is, therefore, an important part of any successful projection pursuit procedure.

This article presents a new projection pursuit algorithm. Its projection index has superior sensitivity and similar robustness properties to the Friedman–Tukey (1974) index, and it is much more rapidly computable. The optimization procedure is faster and far more thorough. Finally, a systematic solution to the structure removal problem is presented.

2. THE PROJECTION INDEX

The projection index forms the heart of a projection pursuit method. It defines the intent of the procedure. Our intent is to discover interesting structured projections of a multivariate data set. This rather vague goal must be translated into a numerical index that is a functional of the projected data distribution. This functional must vary continuously with the parameters defining the projection and have a large value when the (projected) distribution is defined to be “interesting” and a small value otherwise. The notion of interesting will obviously vary with the application (see Huber 1985). As stated in the Introduction, our goal is to discover additional structure not captured by the correlational structure of the data. A way to insure this is to make the projection index invariant to all nonsingular affine transformations in the p -variable data space (see Huber 1985; Jones 1983).

Although the notion of interesting may be difficult to quantify, the converse notion of uninteresting seems more straightforward. Huber (1985) and Jones (1983) gave strong heuristic arguments to the effect that the (projected) normal distribution ought to be considered the least interesting:

1. The multivariate normal density is elliptically symmetric and is totally specified by its linear structure (location and covariances).

2. All projections of a multivariate normal distribution are normal. Therefore, evidence for nonnormality in any projection is evidence against multivariate joint normality. Conversely, if the least normal projection is—not significantly different from—normal, then there is evidence for joint normality of the measurement variables.

3. Even if several linear combinations of variables are (possibly highly) structured (nonnormal), most linear combinations (views) will be distributed approximately normally. Roughly, this is a consequence of the central limit theorem (sums tend to be normally distributed). This notion was made precise by Diaconis and Freedman (1984).

4. For fixed variance, the normal distribution has the least information (Fisher, negative entropy).

Since our goal is the discovery of nonlinear structure, it is principally the elliptical symmetry of the multivariate normal that causes it to be regarded as least interesting for our purposes. Depending on the context, other elliptically symmetric distributions may serve as well or better. But (as exploited below) the computational tractability of the normal makes it an attractive choice among the family of elliptically symmetric distributions.

Following this view, any test statistic for testing normality could serve as the basis for a projection index. Different test statistics have the property of being more (or less) sensitive to different alternative distributions. It is the particular alternatives that are of interest here, since (in the context of projection pursuit) they define the notion of an *interesting* distribution. We must choose a statistic that has preferential power against the (projected) distributions that we are seeking with our projection pursuit algorithm.

The most powerful tests for nonnormality emphasize alternative distributions with heavy tails. Our intent is to seek projected distributions that exhibit clustering (multimodality) or other kinds of nonlinear associations. Such distributions differ from the normal mainly near the center of the distribution, rather than in the tails. Therefore, we seek a projection index (test statistic) that emphasizes departure from normality in the main body of the distribution and gives correspondingly less weight to the tails.

Since the projection index serves as the objective function for a multiparameter optimization, its computational properties are crucial. For a given set of parameter values, its value should be rapidly computable. It should be absolutely continuous so that at least its first derivatives (with respect to the parameters) exist everywhere. These derivatives should also be rapidly computable.

The most computationally attractive projection indexes are based on polynomial moments. No sorting of the projected values is required, and the values of the polynomials as well as their derivatives are rapidly computed by means of recursion relations. Since we are interested in departure from normality, it would be natural to base a projection index on the standardized (absolute) cumulants of the projected distribution (see Huber 1985, p. 445). These are the moments of the Hermite polynomials. Jones (1983) suggested a projection index based on sums of squares of standardized cumulants.

Despite their computational attractiveness, projection indexes based on standardized cumulants are (unfortunately) not useful for our application. This is because they very heavily emphasize departure from normality in the tails of the distribution. For example, a (projected) distribution with only slightly heavier than normal tails re-

ceives a much higher index value than a highly clustered projection.

It is possible to base a projection index on moments having the required statistical properties. Such an index is developed in this section. The main idea is to change scale by first transforming the projected data using the normal cdf and then comparing the transformed distribution with the uniform.

Following Huber (1985), the algorithm will be described first in its abstract version. That is, we imagine it operating on a p -dimensional continuous probability distribution. This makes some of the notation simpler. In our case, the practical version (i.e., applied to data samples) is usually obtained by simply replacing the expected value operation with the corresponding data average. There are other minor differences that are pointed out at the appropriate places in the description. Random variable terminology will be used. Uppercase letters will denote random variables and their lowercase counterparts (usually with subscripts) will denote realized values in the sample.

As a first computational economy, we begin by “sphering” the data (Huber 1981; Jones 1983; Tukey and Tukey 1981). The idea is to perform a linear transformation (rotation, location, and scale change) that removes (incorporates) all of the location, scale, and correlational structure. Let Y be a random variable in R^p . We perform an eigenvalue–eigenvector decomposition of the covariance matrix

$$\Sigma = E[(Y - EY)(Y - EY)^T] = UDU^T,$$

with U an orthonormal and D a diagonal $p \times p$ matrix. We then define new variables $Z = D^{-1/2}U(Y - EY)$. More specifically, let q be the rank of Σ . Then the q components of Z are given by

$$Z_j = (1/\sqrt{D_j}) \sum_{i=1}^p U_{ij}(Y_i - EY_i), \quad 1 \leq j \leq q. \quad (1)$$

The rows of U and D are assumed ordered in descending (nonnegative) values of D_j . By definition $E[Z] = 0$ and $E[ZZ^T] = I$, the identity matrix.

The computational advantage gained by sphering is reflected by the fact that data constraints in the N -dimensional observation space become geometrical constraints in the p -dimensional variable space. First, any linear combination $X = \alpha^T Z = \sum_{i=1}^q \alpha_i Z_i$ has variance $\text{var}(X) = \alpha^T \alpha = \sum_{i=1}^q \alpha_i^2$; thus enforcing the constraint

$$\alpha^T \alpha = 1 \quad (2)$$

insures that all linear combinations have unit variance. Second, two linear combinations based on orthogonal vectors are uncorrelated. That is, $\alpha^T \beta = 0$ implies that $E[(\alpha^T Z)(\beta^T Z)] = 0$.

All operations are performed on the sphered variables Z . (Only at the end do we transform the solutions back to reference the original coordinates Y .) This frees us from having to compute variances in individual projections in order to standardize density estimates during the numerical optimization. Since sphering need only be done once at the beginning, this results in a substantial computational

saving. By definition the Z variables are affine invariant, thus any (orthogonally invariant) projection index based on them will inherit this property.

Although in most exploratory applications two or higher dimensional projection pursuit will likely prove the more informative, we begin by describing our projection index for a one-dimensional projection pursuit. The concepts underlying the two algorithms are nearly identical and the notation is simpler for the one-dimensional case. The extension to two (and higher) dimensions is seen to be straightforward.

In a one-dimensional exploratory projection pursuit we seek a linear combination

$$X = \alpha^T Z \quad (3)$$

such that the probability density $p_\alpha(X)$ is relatively highly structured. As discussed earlier, we regard the (standard) normal as the least structured density, and we are concerned with finding departures that manifest themselves in the main body of the distribution rather than in the tails. To this end we begin by performing a transformation

$$R = 2\Phi(X) - 1, \quad (4)$$

with $\Phi(X)$ being the standard normal cdf

$$\Phi(X) = (1/\sqrt{2\pi}) \int_{-\infty}^X e^{-1/2t^2} dt. \quad (5)$$

Clearly, R takes on values in the interval $-1 \leq R \leq 1$, and if X follows a standard normal distribution then R will be uniformly distributed in this interval. Specifically,

$$p_R(R) = \frac{1}{2} p_\alpha \left[\Phi^{-1} \left(\frac{R+1}{2} \right) \right] / g \left[\Phi^{-1} \left(\frac{R+1}{2} \right) \right],$$

with $g(X)$ being the standard normal density. Thus, a measure of nonuniformity of R corresponds to a measure of nonnormality of X . We take as a measure of nonuniformity the integral-squared distance between the probability density of R , $p_R(R)$, and the uniform probability density, $p_U(R) = \frac{1}{2}$, over the interval $-1 \leq R \leq 1$:

$$\int_{-1}^1 \left[p_R(R) - \frac{1}{2} \right]^2 dR = \int_{-1}^1 p_R^2(R) dR - \frac{1}{2}. \quad (6)$$

Our projection index $I(\alpha)$ is taken to be a moment approximation to (6).

Expanding $p_R(R)$ in Legendre polynomials we have

$$\int_{-1}^1 p_R^2(R) dR - \frac{1}{2} = \int_{-1}^1 \left[\sum_{j=0}^{\infty} a_j P_j(R) \right] p_R(R) dR - \frac{1}{2},$$

where the Legendre polynomials are defined by

$$P_0(R) = 1$$

$$P_1(R) = R$$

$$P_j(R) = [(2j-1)RP_{j-1}(R) - (j-1)P_{j-2}(R)]/j \quad (7)$$

for $j \geq 2$. The coefficients are given by

$$a_j = \frac{2j+1}{2} \int_{-1}^1 P_j(R) p_R(R) dR = \frac{2j+1}{2} E_R[P_j(R)]$$

so that

$$\int_{-1}^1 p_R^2(R) dR - \frac{1}{2} = \sum_{j=1}^{\infty} (2j+1) E_R^2[P_j(R)]/2. \quad (8)$$

For a uniform distribution $U(-1, 1)$, $E[P_j(R)] = 0$ for $j > 0$.

Our projection index is obtained by truncating the sum in (8) at order J ,

$$I(\alpha) = \sum_{j=1}^J (2j+1) E_R^2[P_j(R)]/2. \quad (9)$$

Note that this projection index measures departure from normality even if the J -term expansion is not an accurate approximation to $p_R(R)$, since it achieves its minimum value (0) for X normally distributed (R uniform). Of course, for finite J the (projected) normal is not the unique minimizer of $I(\alpha)$. Any distribution of X that after the transformation (4) results in a distribution $p_R(R)$ with zero values for its first J Legendre polynomial moments will also be regarded as a least interesting distribution.

For a practical version of the algorithm operating on a data sample, the expected values in (9) are estimated by the corresponding sample averages. Substituting from (3) and (4) we have

$$\hat{I}(\alpha) = \frac{1}{2} \sum_{j=1}^J (2j+1) \left[\frac{1}{N} \sum_{i=1}^N P_j(2\Phi(\alpha^T z_i) - 1) \right]^2 \quad (10)$$

as the sample version of our projection index. This is to be maximized with respect to the q components of α under the constraint $\alpha^T \alpha = 1$. Details of the optimization procedure are given in Section 5.

The projection index (10) can be computed fairly rapidly. Fast approximations (to machine accuracy) for the normal integral (5) exist (see Kennedy and Gentle 1980) and are provided as built-in intrinsic functions by many programming language compilers. The Legendre polynomials to order J are quickly obtained via the recursion relation (7).

For efficient optimization it is useful to have derivatives of the objective function. These are easily obtained for our projection index via the chain rule for differentiation. The result is

$$\frac{\partial I}{\partial \alpha_k} = \frac{2}{\sqrt{2\pi}} \sum_{j=1}^J (2j+1) \times E[P_j(R)] E[P'_j(R) e^{-X^2/2} (Z_k - \alpha_k X)], \quad (11)$$

with X given by (3) and R given by (4). The derivatives of each Legendre polynomial with respect to its argument is rapidly obtained by the recursion relation

$$P'_j(R) = 1, \quad P'_j(R) = R P'_{j-1}(R) + j P_{j-1}(R), \quad (12)$$

for $j > 1$. The derivative calculation (11) takes into account the constraint $\alpha^T \alpha = 1$ by keeping the gradient vector $\nabla_{\alpha} I(\alpha)$ orthogonal to the gradient of the constraint function $\nabla_{\alpha} (\alpha^T \alpha) = 2\alpha$. This is the purpose of subtracting $\alpha_k X$ from Z_k in the second expectation (11). The derivatives

of $\hat{I}(\alpha)$ (10) are obtained by applying sample averages in place of the expectation operators.

The projection index for a two-dimensional projection pursuit is developed in direct analogy with the one-dimensional index. We seek two linear combinations

$$X_1 = \alpha^T Z, \quad X_2 = \beta^T Z, \quad (13)$$

such that the joint distribution (probability density) $p_{\alpha\beta}(X_1, X_2)$ is highly structured. Since we are interested in non-linear structure we require the two linear combinations to be uncorrelated, $\text{corr}(X_1, X_2) = 0$. As a consequence of our definition of Z (data sphering) this constraint is equivalent to requiring α to β to be orthogonal, $\alpha^T \beta = 0$. We must also require that the variances in all projections be equal. This is insured by the sphering and constraints $\alpha^T \alpha = \beta^T \beta = 1$.

We regard the bivariate standard normal to be the least structured joint distribution and are interested in departures that are manifest in the main body of the distribution rather than in the tails. We begin by transforming the X_1, X_2 plane to the square $(-1, 1) \times (-1, 1)$ by means of

$$R_1 = 2\Phi(X_1) - 1, \quad R_2 = 2\Phi(X_2) - 1, \quad (14)$$

with $\Phi(X)$ defined by (5). If X_1, X_2 have a joint standard normal distribution, then R_1, R_2 will be uniformly distributed on the square. We take as a measure of nonuniformity the integral-square distance from the uniform

$$\begin{aligned} & \int_{-1}^1 \int_{-1}^1 \left[p_R(R_1, R_2) - \frac{1}{4} \right]^2 dR_1 dR_2 \\ &= \int_{-1}^1 \int_{-1}^1 p_R^2(R_1, R_2) dR_1 dR_2 - \frac{1}{4}. \end{aligned} \quad (15)$$

Our projection index $I(\alpha, \beta)$ is taken as a product moment approximation to (15). Expanding $p_R(R_1, R_2)$ in a product Legendre expansion and proceeding in direct analogy with the development of the one-dimensional index we have

$$\begin{aligned} & \int_{-1}^1 \int_{-1}^1 p_R^2(R_1, R_2) dR - \frac{1}{4} \\ &= \frac{1}{4} \sum_{j=0}^{\infty} \sum_{k=0}^{\infty} (2j+1)(2k+1) E^2[P_j(R_1) P_k(R_2)] - \frac{1}{4}, \end{aligned}$$

with the Legendre polynomials defined by (7). Our bivariate projection index is obtained by truncating the expansion at order J ,

$$\begin{aligned} I(\alpha, \beta) &= \sum_{j=1}^J (2j+1) E^2[P_j(R_1)]/4 \\ &+ \sum_{k=1}^J (2k+1) E^2[P_k(R_2)]/4 \\ &+ \sum_{j=1}^J \sum_{k=1}^{J-j} (2j+1)(2k+1) \\ &\times E^2[P_j(R_1) P_k(R_2)]/4. \end{aligned} \quad (16)$$

As for the univariate index, derivatives of the bivariate

index are easily obtained:

$$\begin{aligned}\frac{\partial I}{\partial \alpha_m} &= (1/\sqrt{2\pi}) \sum_{j=1}^J (2j+1) E[P_j(R_1)] \\ &\quad \times E[P'_j(R_1) e^{-(1/2)X_1^2} (Z_m - \alpha_m X_1 - \beta_m X_2)] \\ &\quad + (1/\sqrt{2\pi}) \sum_{j=1}^J \sum_{k=1}^{J-j} (2j+1)(2k+1) \\ &\quad \times E[P_j(R_1) P_k(R_2)] \\ &\quad \times E[P'_j(R_1) P'_k(R_2) e^{-(1/2)X_2^2} (Z_m - \alpha_m X_1 - \beta_m X_2)], \\ \frac{\partial I}{\partial \beta_n} &= (1/\sqrt{2\pi}) \sum_{k=1}^J (2k+1) E[P_k(R_2)] \\ &\quad \times E[P'_k(R_2) e^{-(1/2)X_2^2} (Z_n - \alpha_n X_1 - \beta_n X_2)] \\ &\quad + (1/\sqrt{2\pi}) \sum_{j=1}^J \sum_{k=1}^{J-j} (2j+1)(2k+1) \\ &\quad \times E[P_j(R_1) P_k(R_2)] \\ &\quad \times E[P_j(R_1) P'_k(R_2) e^{-(1/2)X_2^2} (Z_n - \alpha_n X_1 - \beta_n X_2)].\end{aligned}$$

The quantities X_1 and X_2 are given by (13), and R_1 and R_2 are given by (14). The data version of the index and its derivatives are obtained by substituting sample averages for the expectations. The derivatives account for the constraints ($\alpha^T \alpha = \beta^T \beta = 1$, $\alpha^T \beta = 0$) by keeping the gradient vector simultaneously orthogonal to the gradients of the three constraint functions.

The computation of the bivariate index is analogous to that of the univariate index. For a corresponding order J it is more expensive, since it and its derivatives contain more terms. In addition, the optimization is with respect to twice as many parameters. On the other hand, the bivariate solutions often contain considerably more information concerning the multivariate density.

There is one (user-defined) parameter associated with the (one- or two-dimensional) projection index. It is the order J ((9), (16)) of the polynomial expansion of the (transformed) density. It controls the amount of smoothness imposed on the approximation. Limited experience indicates that the results are insensitive to the value chosen for J over a fairly wide range ($4 \leq J \leq 8$) except for very small sample size. Intuition suggests that the value of J should increase with the sample size, but there are as yet no specific guidelines. The computation increases linearly with increasing J for one-dimensional projection pursuit and quadratically for two-dimensional projection pursuit.

3. STRUCTURE REMOVAL

The purpose of the optimization algorithm (detailed in Sec. 5) is to find a substantive maximum of the projection index. The corresponding (one- or two-dimensional) projection will (we hope) present an informative view of the p -dimensional data density. It is unlikely, however, that there is only one such informative view. Usually, the non-normality of the full p -dimensional data distribution will be manifest in several one- or two-dimensional projec-

tions. Each of these projections can help in the identification and interpretation of the effects. Moreover, there is no reason to believe that the algorithm will find the most informative of these views first. For these reasons it is important that the projection pursuit procedure find as many of these informative views as possible.

A variety of approaches for accomplishing structure removal have been suggested (see Huber 1985, p. 449). The most systematic of these (for one-dimensional projection pursuit) is the recursive approach associated with the projection pursuit density estimation procedure (Friedman, Stuetzle, and Schroeder 1984). After an interesting projection has been found (solution maximizing the projection index), remove the structure that makes the projection interesting (deflate that maximum of the objective function) and then remaximize the projection index. This can be done repeatedly. In the projection pursuit density estimation approach this was implemented using a complex strategy for maintaining and updating an estimate for the p -dimensional probability density and involved Monte Carlo sampling. It is thus computationally quite expensive. We present a simple procedure for structure removal that is computationally much faster and has a straightforward implementation for two-dimensional projections.

By definition of our projection index, a view (projected density) has zero interest if it is standard normal. Therefore, the structure can be removed by applying a transformation that takes the (projected) density to a standard normal distribution. We thus require a transformation of the q variables, the result of which renders a standard normal distribution in the projected subspace, but leaves all orthogonal directions unchanged. We develop such a transformation below.

We first describe the procedure for a one-dimensional projection. Let $X = \alpha^T Z$ be a one-dimensional projection ($\text{var}(X) = 1$) and $F_\alpha(X)$ be its cdf. Then applying the transformation

$$X' = \Phi^{-1}(F_\alpha(X)) \quad (17)$$

to X results in a standard normal distribution for X' . Here Φ^{-1} is the inverse of the standard normal cdf (5).

Let U be an orthonormal ($q \times q$) matrix with α (the projection pursuit solution) as the first row. Then applying the linear transformation $T = UZ$ results in a rotation such that the new first coordinate is $T_1 = \alpha^T Z = X$. Let Θ be a (vector) transformation (with components $\theta_1 \cdots \theta_q$) that takes T_1 to a standard normal distribution and is the identity transformation on $T_2 \cdots T_q$:

$$\theta_1(T_1) = \Phi^{-1}(F_\alpha(T_1))$$

$$\theta_j(T_j) = T_j, \quad 2 \leq j \leq q. \quad (18)$$

Then the transformation

$$Z' = U^T \Theta(UZ) \quad (19)$$

transforms the projection $X = \alpha^T Z$ to a standard normal distribution leaving all orthogonal directions unchanged. We then reapply the maximization procedure with a projection index based on Z' .

By definition, the q -variate distribution of Z' exhibits no structure in the projection $X' = \alpha^T Z'$ (zero value of the projection index). The joint distribution of Z , $p(Z)$, and that of Z' , $p'(Z')$, determine the same conditional density given $\alpha^T Z$:

$$p(\cdot | \alpha^T Z) = p'(\cdot | \alpha^T Z'). \quad (20)$$

In fact,

$$p'(Z') = p(Z') [g(\alpha^T Z') / p_\alpha(\alpha^T Z')], \quad (21)$$

with p_α the (univariate) density of $\alpha^T Z$ and g the standard normal density

$$g(X) = (1/\sqrt{2\pi})e^{-X^2/2}. \quad (22)$$

In this sense the transformation (19) produces a new (vector-valued random) variable Z' whose distribution is as close as possible to that of Z under the constraint that its marginal distribution along α be normal (zero interest). It also produces the closest distribution under this constraint in the sense of the relative entropy distance measure

$$\int \log\left(\frac{p}{p'}\right) p \, dZ = \int \log\left(\frac{p_\alpha}{g}\right) p_\alpha \, d(\alpha^T X) = \min \quad (23)$$

(see Huber 1985).

The data sample version of (17) is easily implemented. One substitutes the empirical distribution $\hat{F}_{\alpha N}(X) (X = \alpha^T Z)$ for the distribution function $F_\alpha(X)$:

$$x'_i = \Phi^{-1}(\hat{F}_{\alpha N}(x_i)) = \Phi^{-1}[(r(x_i)/N) - 1/2N], \quad (24)$$

with $r(x_i)$ being the rank of x_i among the N (projected) observations. This transformation simply replaces each observation by its corresponding normal score in the projection ("Gaussianization").

The process of repeatedly applying projection pursuit on the (structure removed) output of the previous pursuit can be continued until several applications result in finding no additional interesting structure. It should be noted that "Gaussianizing" a solution projection in this way at a particular stage modifies the normality of (nonorthogonal) previous solution projections so that they no longer have exactly zero interest (unless backfitting is employed—see Friedman et al. 1984). Experience indicates, however, that the structure induced in previous solution projections by structure removal at later stages is small.

Structure removal in two dimensions is more difficult. We need a transformation that takes a general bivariate distribution $p_{\alpha\beta}(X_1, X_2)$ to the bivariate standard normal

$$g(X_1, X_2) = (1/2\pi)\exp(-(X_1^2 + X_2^2)/2). \quad (25)$$

There is no difficulty in theory. One can transform one of the margins, say X_1 , to normality via (17) and then transform each conditional orthogonal marginal $p(X_2 | X_1)$ to normality [again via (17)]. But this prescription does not lead to a practical algorithm for application to (bivariate) data. A practical algorithm can be based on the observation that all projections of a normal distribution are normal. The idea is to repeatedly Gaussianize rotated (about

the origin) projections of the solution plane until they stop becoming more normal.

Let

$$X'_1 = X_1 \cos \gamma + X_2 \sin \gamma$$

$$X'_2 = X_2 \cos \gamma - X_1 \sin \gamma \quad (26)$$

be a rotation about the origin through angle γ . The distributions of X'_1 and X'_2 are then each transformed to normality via (17). This process is repeated (on the previously transformed distributions) for several values of γ (0, $\pi/4$, $\pi/8$, $3\pi/8$). This entire process is then repeated until the distributions stop becoming more normal. Any convenient index of (non) normality can be used.

During the first few iterations the nonnormality decreases rapidly in a monotonic fashion as the planar distribution approaches joint normality. After approximate normality has been achieved, the value of the nonnormality index tends to oscillate with small amplitude on successive iterations, sometimes decreasing a small amount on the average. Convergence is defined when approximate stability has been achieved. Note that with finite samples absolute stability is impossible to achieve. Typically, the procedure takes from 5 to 15 complete iterations to converge. It produces bivariate data distributions that are quite close to normal.

In analogy with the univariate case, let U be an orthonormal ($q \times q$) matrix with α and β (the linear combinations determining the solution plane) as the first two rows. The linear transformation $T = UZ$ performs a rotation aligning the first two new coordinates with α and β . Let Θ be a transformation that takes the joint distribution of T_1 and T_2 to standard normal (as described earlier) and is the identity transform on $T_3 \cdots T_q$. Then the transformation $Z' = U^T \Theta(UZ)$ transforms the solution (α, β) plane to bivariate standard normal leaving all orthogonal directions unchanged. Thus the joint distribution of Z and Z' determines the same conditional distribution given $\alpha^T Z$ and $\beta^T Z$,

$$p'(Z') = p(Z') \frac{g(\alpha^T Z') g(\beta^T Z')}{p_{\alpha\beta}(\alpha^T Z', \beta^T Z')}, \quad (27)$$

with $g(X)$ the standard normal and the denominator the joint distribution of $\alpha^T Z$ and $\beta^T Z$.

Friedman and Tukey (1974) suggested two rudimentary forms of structure removal. One was to restrict later solutions to be orthogonal (with respect to the original coordinates and their scales) to previous solutions. This is clearly supplanted by the structure removal technique outlined here. There is no reason to expect good views of the data to be orthogonal with respect to any prespecified metric. The second suggested method was applicable when the structure in the solution projection took the form of clustering. The idea was to isolate the clusters into separate subsamples and apply projection pursuit to each such isolate individually. This could be iterated if clustering were found in subsequent solutions. This second structure removal technique (when applicable) can be viewed as com-

plementary to the method outlined here. If there is clustering and it is largely hierarchical, then the isolation technique can provide a straightforward means for interpreting this kind of structure.

4. DENSITY ESTIMATION

Exploratory projection pursuit as described in the preceding two sections can be incorporated into a multivariate density estimation procedure. Its properties (for one-dimensional projection pursuit) are similar to the projection pursuit density estimation procedure of Friedman et al. (1984) and the projection pursuit density approximation techniques of Huber (1985). But its computational aspects are considerably more attractive.

The projection pursuit strategy outlined in the preceding two sections (distribution version) consists of finding the least normal projection $p_{\alpha_1}(\alpha_1^T Z)$ of the probability density $p(Z)$ by maximizing a measure of nonnormality (9). The procedure is then repeated on the density

$$p_1(Z) = p(Z)g(\alpha_1^T Z)/p_{\alpha_1}(\alpha_1^T Z)$$

[see (21)], obtaining a second solution $\alpha_2^T Z$. The distribution is again modified,

$$\begin{aligned} p_2(Z) &= p_1(Z)g(\alpha_2^T Z)/p_{\alpha_2}^{(1)}(\alpha_2^T Z) \\ &= p(Z) \frac{g(\alpha_1^T Z)g(\alpha_2^T Z)}{p_{\alpha_1}(\alpha_1^T Z)p_{\alpha_2}^{(1)}(\alpha_2^T Z)}, \end{aligned}$$

and so on. Here $p_{\alpha_2}^{(1)}(\alpha_2^T Z)$ is the univariate marginal density of $\alpha_2^T Z$ under the joint density $p_1(Z)$. At the K th iteration one has

$$p_K(Z) = p(Z) \prod_{k=1}^K \frac{g(\alpha_k^T Z)}{p_{\alpha_k}^{(k-1)}(\alpha_k^T Z)}. \quad (28)$$

The quantity in the denominator, $p_{\alpha_k}^{(k-1)}(\alpha_k^T Z)$, is the marginal density of $\alpha_k^T Z$ under the joint distribution $p_{k-1}(Z)$, with $p_0(Z) = p(Z)$.

At some point in the iterative process the projection pursuit algorithm cannot find a projection that deviates substantially from normality. This indicates that $p_K(Z)$ is approximately multivariate standard normal. We then take as our density approximation

$$\bar{p}(Z) = g(Z) \prod_{k=1}^K \frac{p_{\alpha_k}^{(k-1)}(\alpha_k^T Z)}{g(\alpha_k^T Z)}, \quad (29)$$

with

$$g(Z) = \frac{1}{(2\pi)^{q/2}} \exp(-Z^T Z/2). \quad (30)$$

The projected univariate densities $p_{\alpha_k}^{(k-1)}$ can be approximated by any appropriate method. One possibility is to use the Legendre polynomial expansion associated with the projection index

$$p_{\alpha_k}^{(k-1)}(\alpha_k^T Z) = g(\alpha_k^T Z) \sum_{j=0}^J (2j+1) E_{k-1}[P_j(R_k)] P_j(R_k)/2,$$

with $R_k = 2\Phi(\alpha_k^T Z) - 1$ and $E_{k-1}(\cdot)$ the expected value under $p_{k-1}(Z)$. Truncation can be used to insure non-

negativity. Substituting this into (29) we have for this (multivariate) density approximation

$$\begin{aligned} \bar{p}(Z) &= g(Z) \prod_{k=1}^K \left[\sum_{j=0}^J (2j+1) \right. \\ &\quad \times \left. E_{k-1}[P_j(R_k)] P_j(R_k) / 2 \right] / 2, \quad (31) \end{aligned}$$

with $E_{k-1}[P_j(R_k)]$ being the expected value of the associated (adjacent) Legendre polynomial under $p_{k-1}(Z)$. A density estimate is obtained by estimating $E_{k-1}[P_j(R_k)]$ by sample averages over the transformed variables

$$Z_{k-1} = U_{k-1}^T \Theta_{k-1}(U_{k-1} Z_{k-2})$$

[see (19)] obtained from the structure removal process during the projection pursuit. Thus

$$\hat{E}_{k-1}[P_j(R_k)] = \frac{1}{N} \sum_{i=1}^N P_j[2\Phi(\alpha_k^T z_{(k-1)i}) - 1]. \quad (32)$$

Here $Z_0 = Z$, the original (sphered) data.

This density approximation/estimate is strongly influenced by the main body of the data and will give poor (usually under-) estimates in the outlying tails. This is a result of the transformation (4), which compresses the tails into small intervals near the extremes of the interval $(-1, 1)$. Long-tailed (compared with the normal) projected (univariate) distributions will result in very sharp spikes in the transformed density $p_R(R)$ at the ends of the interval. These cannot be captured by a low to moderate degree ($4 \leq J \leq 8$) Legendre polynomial expansion that will substantially underestimate them. This is how the projection index (9), (10), and (16) achieves its robustness against long-tailed scatter. In addition, of course, the projection index (by design) will tend not to produce solutions for which the projected density has no other structure but long tails. As a result the density approximation/estimate provided by (31) and (32) will focus on capturing the density variation in the central part of the distribution and will approximate its tails by the tails of the normal. Of course, there is no other method that produces accurate density estimates in the tails of a multivariate distribution either. It is interesting to note the connection between this approach to projection pursuit density approximation and the "analytic" approach proposed by Huber (1985).

It is possible to base a density approximation/estimation procedure on the two-dimensional algorithm in direct analogy with the development of the one-dimensional procedure described previously. But it might not work as well as the one-dimensional algorithm (for this purpose) because of its increased complexity.

5. OPTIMIZATION STRATEGY

Although it is an engineering detail, the technique used for maximizing the (one- and two-dimensional) projection index strongly influences both the *statistical* and the *computational* aspects of the procedure. The statistical power of the method is reflected in its ability (for a given sample size and data dimension) to find substantive maxima of

the projection index. As observed by Switzer (1970) there are "an almost inevitable multiplicity of decidedly sub-optimal local maxima" mostly caused by sampling fluctuations. This can distract a projection pursuit algorithm from finding important views (substantive maxima). These pseudomaxima can be visualized as a high-frequency ripple superimposed on the main variational structure of the objective function (projection index). The amplitude of these ripples increases with increasing dimension and decreasing sample size. The extent to which the optimization procedure can ignore ("step over") these pseudomaxima, and thus avoid being trapped by them, determines to a great extent its statistical power.

On very smooth objective functions the most powerful optimization methods [steepest descent, conjugate gradients, quasi-Newton (see Gill, Murray, and Wright 1981)] involve the use of first derivatives. This is why the ability to rapidly compute derivatives was a design goal for our projection index. These methods very effectively (rapidly and accurately) find the first maximum of an objective function uphill from a starting point. Unfortunately, when applied to a projection index with the ripple phenomenon described previously, this will very likely be a pseudomaximum, unless the starting point is within the domain of attraction of a substantive maximum. The optimization strategy used by Friedman and Tukey (1974) did not employ (exact) derivatives and took fairly large steps in its search for a maximum. This gave it some robustness against pseudomaxima at the expense of considerable computational effort.

We employ a hybrid optimization strategy. It begins with a simple (coarse stepping) optimizer that is designed to very rapidly get close to (within the domain of attraction) of a substantive maximum. A gradient method (quasi-Newton) is then used to quickly converge to the solution.

We begin with the maximization of the one-dimensional index $I(\alpha)$ with respect to the q components of α ($\alpha_1 \cdot \dots \cdot \alpha_q$). As a first step $I(\alpha)$ is maximized over the coordinate axes $\alpha = e_i$ ($1 \leq i \leq q$). Note that since we are working with the sphered data, Z , these axes are in fact the principal component directions when referenced to the original data, Y (1). Let α^* be the resulting maximizing axis. Starting with this direction the following optimization algorithm is performed:

Loop:

$$I_0 = I(\alpha^*)$$

For $i = 1$ to q do:

$$f_+ = I[(1/\sqrt{2})(\alpha^* + e_i)/(1 + \alpha_i^*)^{1/2}]$$

$$f_- = I[(1/\sqrt{2})(\alpha^* - e_i)/(1 - \alpha_i^*)^{1/2}]$$

If $f_+ > f_-$ then $f = f_+$; $s = +1$

else $f = f_-$; $s = -1$

end If

If $f > I(\alpha^*)$ then

$$\alpha^* \leftarrow (1/\sqrt{2})(\alpha^* + s \cdot e_i)/(1 + s \cdot \alpha_i^*)^{1/2}$$

end If

end For

If $I(\alpha^*) = I_0$ then done

end Loop

This search algorithm takes large steps, and thus it cannot be expected to converge to the value of α^* at a maximum of $I(\alpha)$. Because of its coarse stepping, however, it is much less likely than a gradient method to be trapped on pseudomaxima, thereby allowing it to converge in the vicinity of a substantive maximum. This algorithm typically requires from two to four passes over the coordinates (executions of the For loop) to converge.

Starting with the value of α^* obtained upon convergence of (33), a gradient-directed optimization method is then employed to rapidly ascend to a maximum of the projection index $I(\alpha)$. We have employed both steepest-ascent and quasi-Newton methods with comparable results.

The maximization of the two-dimensional index $I(\alpha, \beta)$ is done in an analogous manner. First, it is maximized over the $q(q-1)/2$ pairs of coordinate axes, $\alpha = e_i$, $\beta = e_j$ ($2 \leq i \leq q$, $1 \leq j \leq i$). Then, starting with the best coordinate pair, an algorithm analogous to (33) is executed. In this algorithm the For loop is over $2q$ variables (the q components of α and the q components of β), and the constraint $\alpha^T \beta = 0$ must be maintained in addition to $\alpha^T \alpha = \beta^T \beta = 1$. Finally, after this procedure converges a gradient directed optimization is employed to rapidly find a maximum.

6. REMARKS

6.1 Robustness

The one- and two-dimensional projection indexes (10) and (16) are (by design) quite robust in that they are largely unaffected by extreme outliers. As a consequence the pursuit procedure is thus similarly robust. Structure removal is also clearly unaffected by outliers. The only nonrobust aspect of the procedure is the data sphering. It is based on the sample covariance matrix, which is strongly influenced by extreme outliers. Experience indicates that this projection pursuit procedure does not seem to be severely degraded when based on badly sphered data due to outliers. Nevertheless it seems sensible to use robust sphering when possible.

There are several methods for robust estimation of a covariance matrix (see Devlin, Gnanadesikan, and Kettenring 1981). In fact there are several attractive (but computationally expensive) projection pursuit approaches (Chen and Li 1981). We have implemented a simple multivariate trimming method. It begins by sphering using all of the data. All observations for which $Z^T Z > D$ [see (1)] are deleted, and the remaining data are resphered. Here D is some prespecified threshold conveniently taken to be a high ($\sim 1 - .01/N$) quantile of the chi-squared distribution on q degrees of freedom. This procedure can be iterated until no observations are deleted. Often D is adjusted so that no more than a certain (small) fraction of the data are deleted.

6.2 Preliminary Dimensionality Reduction

The power of the projection pursuit algorithm to find important structure decreases with decreasing sample size and increasing dimension (see the next section). As re-

marked earlier, the covariance structure (linear associations) often does not align with the nonlinear structure (clustering, nonlinear relationships) that we are seeking with our projection pursuit algorithm. A typical exception to this, however, has to do with the existence of a subspace containing only a tiny fraction of the data variation.

Clearly, if a subspace contains no data variation, it cannot contain any structure. In this case the covariance matrix is singular, and the dimension of the search space is reduced to $q < p$ (1), the rank of the covariance matrix. If there exists a $(p - q)$ -dimensional subspace for which the data variation is very small compared with the complement subspace (covariance matrix nearly singular), then this subspace is usually dominated by the noise in the system and contains little data structuring. If this turns out to be the case, then the power of the projection pursuit procedure can be enhanced (and computation reduced) by restricting the pursuit search to the q -dimensional complement space. If not, any structure represented in the $(p - q)$ -dimensional subspace will be ignored. But, in cases in which the data dimension is very high and the sample size is small, there may be no alternative but to restrict the projection pursuit search to the subspace spanned by the largest q principal component axes, where q is determined by the sample size. Moreover, if a high-dimensional projection pursuit is unsuccessful in finding interesting structure, it might be worthwhile to restrict the search dimension (as described earlier) and try again.

6.3 Preliminary Transformations

Sometimes marginal distributions on the original measurement coordinates exhibit considerable (nonnormal) structure. For example, substantial skewness is often associated with quantities that take on only positive values. Since inspection of the coordinate marginals should always be among the first parts of any data analysis, this structure is easily discovered. Often the data analyst would like to know if there is additional structure associated with combinations of the variables. In this situation it makes sense to perform a transformation on highly structured coordinates, to remove the obvious structure, and then apply projection pursuit to the data after these selected transformations. For example, taking logarithms often removes positive skewness [see Mosteller and Tukey (1977, chap. 5) for a catalog of such "first aid" transformations]. Of course, the structure along any marginal can be completely removed (from the point of view of projection pursuit) by replacing the coordinate values with their corresponding normal scores. Such "Gaussianized" variables would then only contribute to data structuring through their (nonlinear) associations with other variables.

Sometimes measurement variables take on only a small number of distinct values. This can be caused by the nature of the variables themselves or by the measuring process. If the number of such values is small (compared with the sample size), the marginal distribution will exhibit many identical values or ties. If the number of distinct values is very small (less than five or so), the marginal distribution

appears highly structured when compared with the normal. Gaussianization of such variables is a possible remedy; however, it is important that observations with the same value be ordered randomly so that associations between variables are not induced by the fact the original ordering of the observations may be associated with the values of some of the measurement variables. Categorical (nominal) variables can be accommodated by introducing corresponding (0/1) dummy variables along with randomly breaking the resulting ties and then Gaussianizing.

6.4 Significance

It is important to know whether a view is indicative of actual structure in the population or whether it is an artifact of sampling fluctuations. One way to assess this is to compare the corresponding solution projection index with values obtained by applying the procedure to Gaussian data. One can repeatedly generate random samples from a Gaussian distribution of the same dimension and cardinality as the data sample. The identical procedure that was applied to the data can then be applied to these Monte Carlo multivariate normal samples. A comparison of the resulting (null) distribution of projection index values to the data sample value gives an indication of its significance.

6.5 Adjusted Data Plots

With the exception of the first, there are two projections of interest associated with each projection pursuit solution. One is the distribution of the data projected onto the solution line or plane. The other is the projection of the transformed data after removal of the structure associated with all previous solutions. In distributional terms, the former projection is the (joint) distribution of $\alpha^T Z$ (and $\beta^T Z$) under the original joint data density $p(Z)$. The latter projection is that distribution under $p_{K-1}(Z)$ (28), or the corresponding two-dimensional analog [see (27)], where K is the iteration number. At the K th iteration, projection pursuit is applied to $p_{K-1}(Z)$ to find additional structure. The K th solution projection index and the resulting projection of the transformed data reflect the additional structure adjusted for (not directly associated with) all previous solutions. We refer to these as *adjusted* data plots. They are the analog of residual plots in regression analysis.

6.6 Interpretation

The output of an exploratory projection pursuit is a collection of views of the multivariate data set. These views are selected to be those that independently best represent the nonlinear aspects of the joint density of the measurement variables as reflected by the data. The nonlinear aspects are emphasized by maximizing a robust affine invariant projection index, whereas the independence is induced by the structure removal process. The data analyst has at his disposal the values of the parameters (variable loadings) that define each solution (line or plane) as well as the projected data density. This information can be used

to try to interpret any nonlinear effects that might be uncovered. A visual representation of the projected data density (histogram, smoothed density estimate, scatter or contour plot) can be inspected to ascertain the nature of the effect (clustering, nonlinear relationship). The (scaled) variable loadings that define the corresponding solution indicate the relative strength that each (corresponding) variable contributes to the observed effect.

In a two-dimensional projection pursuit the visual impact of the projected data density is insensitive to a particular orientation (within the plane) of the orthogonal axes used to define the solution plane. A rigid rotation of the projected data about the origin of the solution plane provides the same picture of the nonlinear effects. Therefore, it makes sense to orient the defining axes so that the resulting variable loadings are most easily interpreted. This usually means maximizing the variance (or some other dispersion measure) of the (normed) variable loadings so as to give large loadings to as few variables as possible. Note that this "varimax" rotation is performed on the defining axes in the sphered variable representation Z , whereas the criterion to be maximized is the variance of the corresponding (normed) coefficients in the original data variables Y [see (1)]. Experience indicates that a most useful varimax rotation is one that maximizes the variance of the loadings associated with *one* of the defining axes (e.g., the vertical axis).

Often projection pursuit solutions give rise to small loadings on several variables. If all (original) variables Y_j ($1 \leq j \leq p$) have similar scales, then those with small loadings have correspondingly less effect in defining the solution. For interpretational purposes it is often important to know whether these variables have any importance to the observed structure. This is most easily determined by (manually) setting the small coefficients to 0 (in perhaps a reverse stagewise manner) and then reprojecting the data.

6.7 Three- and Higher-Dimensional Projection Pursuit

In the preceding sections one- and two-dimensional exploratory projection pursuit algorithms were described. For data exploration the two-dimensional algorithm is likely to prove the most useful owing to the increased richness of structure that can be represented in two dimensions. In principle there is no upper limit to the dimensionality of the solution subspace. One could envision a projection pursuit for finding informative three- and higher-dimensional views, although it is not clear that the richness of the representation would increase as much as in going from one to two dimensions. Three-dimensional representations can be viewed using kinematic graphic techniques such as rigid rotation (see Donoho, Donoho, and Gasko 1986; Fisherkeller, Friedman, and Tukey 1975; McDonald 1984). There are techniques for (approximate) viewing of densities in four dimensions (see Tukey and Tukey 1985). Among the more promising approaches are the "grand tour" methods (Buja and Asimov 1985).

A projection index for higher-dimensional pursuit is easily developed in analogy with the two-dimensional index. The computational expense would be somewhat greater owing to the increased complexity of the product Legendre polynomial expansion and the increased number of optimization parameters. A more serious problem encountered with higher-dimensional projection pursuit is associated with the structure removal process. The difficulty lies in transforming the higher-dimensional (projected) distribution to joint standard normality. A strategy analogous to that for the two-dimensional case would require a great many directions if they are chosen regularly on the unit sphere. A better strategy would be to choose a carefully selected set of directions that depend on the actual (projected) data density. This is accomplished by running a one-dimensional projection pursuit algorithm in the higher-dimensional projected (solution) subspace. As discussed in Section 4, this in fact constructs a transformation of the original data density to standard normality.

7. EXAMPLES

In this section we present the results of running the one- and two-dimensional projection pursuit algorithms on data. The first two examples are simulation studies in which we try to assess the sample size requirements for detecting (known) structure as a function of increasing data dimensionality. The next three examples show the results of applying two-dimensional projection pursuit to real data sets of varying dimension and cardinality. In all examples no robustification was introduced into the sphering. To aid in interpretation all variables were standardized to zero mean and unit variance ("auto-scaled") before projection pursuit was applied. In the applications of two-dimensional projection pursuit the solutions were rotated to maximize the variance of the loadings on the second (vertical) defining axis (see Sec. 6.6). Except for the first real data example (states data, Sec. 7.3), the order of the Legendre expansion J [see (10) and (16)] was taken to be $J = 6$.

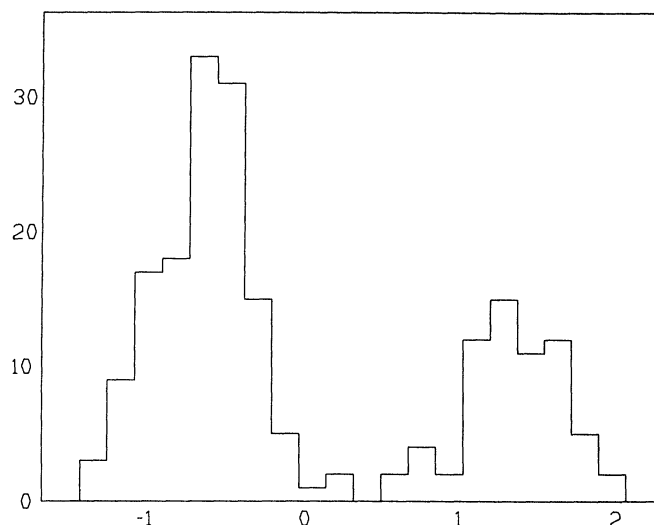


Figure 1. Single Clustered Projection.

7.1 Single Clustered Projection in Several Dimensionalities

The purpose of this study is to get an idea of how the sample size requirements for finding a single structured projection increase with the dimensionality of the data sample. The population for this study is a Gaussian mixture. Two-thirds of the data are generated from a joint standard normal distribution, and the remaining one-third is normal with unit covariance matrix, but with location displaced six units in a random direction. The data are then scaled to have unit variance in this direction so that the structure is not reflected in the linear associations among the variables.

Three experiments were performed at dimensionalities 5, 10, and 15, respectively. Since (by design) the data structuring appears in only one view (the direction defined by the difference of the means), this example tests the projection pursuit algorithm's ability to find structure in

the presence of an increasing number of pure noise variables. From the point of view of projection pursuit this represents a difficult example, since the structure appears in only a single projection. Figure 1 shows a histogram of a random sample of size 200 from this population projected onto the solution direction.

At each dimensionality a series of one-dimensional projection pursuit runs were made to get a rough idea of the threshold sample size at which the algorithm could reliably find the (known) structured projection. Since (by design) the projection pursuit algorithm has some difficulty at these (threshold) sample sizes, a measure of that difficulty is the iteration number (projection pursuit followed by structure removal) at which the algorithm discovers the known structure as opposed to spurious structure (pseudomaxima) induced by the small sample size and/or high dimensionality. If for a given sample size and dimensionality the algorithm repeatedly finds the known structure at the first iteration, then it is having little difficulty. If, on the

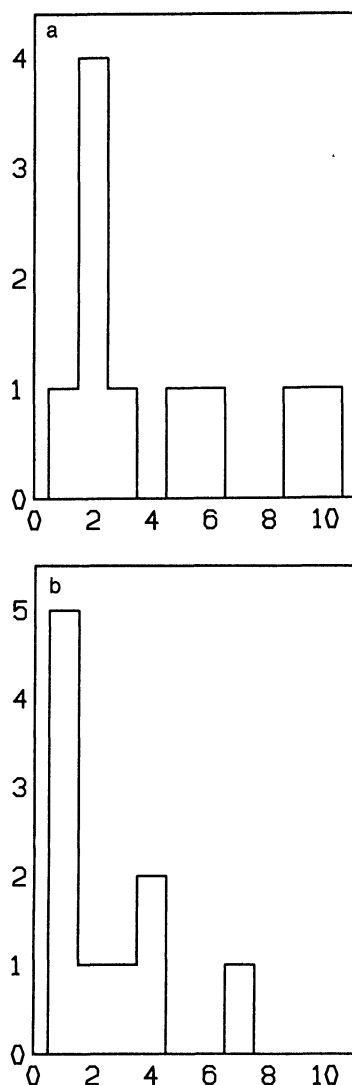


Figure 2. Iteration at Which the True Population Structure Was Found for $p = 5$: (a) $N = 45$; (b) $N = 60$.

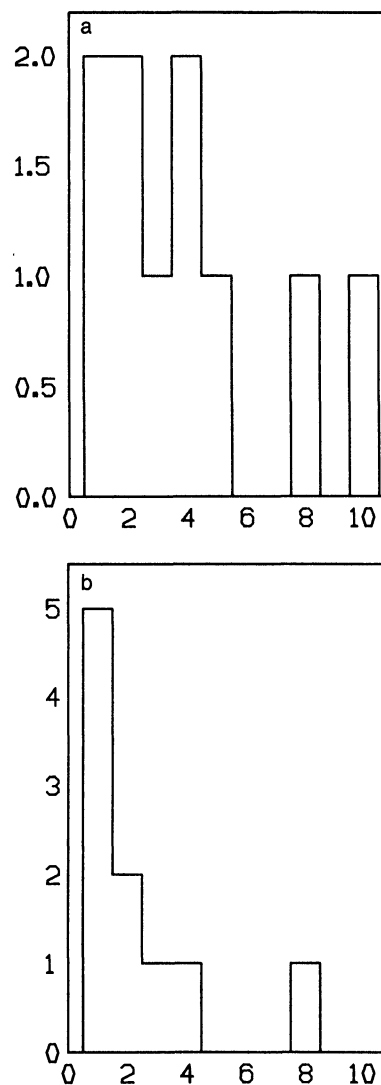


Figure 3. Iteration at Which the True Population Structure Was Found for $p = 10$: (a) $N = 200$; (b) $N = 300$.

other hand, it finds several pseudomaxima (which are subsequently deflated by structure removal) before finding the real structured projection, this is an indication of some difficulty.

Figures 2 through 4 show the distribution of the iteration number at which the true (population) structured projection was found for 10 random samples for each of 6 situations. Each situation consists of a specific dimensionality and sample size. Two sample sizes are shown for each dimensionality. The first (Figs. 2a, 3a, 4a) is a smaller sample size at which the algorithm seems to be having some difficulty and thus represents a minimal cardinality for finding the true structured projection at the corresponding dimensions. The second (larger) sample size (Figs. 2b, 3b, 4b) is seen to be large enough to find the true underlying structure fairly reliably. In all runs the determination as to whether the algorithm found the actual underlying structure was unambiguous. The projection index associated with this solution was typically four to five

times that of the spurious solutions (pseudomaxima), and the solution direction lined up very closely with the direction associated with the underlying (population) optimal projection. It seems that once the optimizer gets close to the true solution (via the course stepping algorithm) the gradient-directed search locks on to it very accurately (and rapidly).

As seen from the figures, the required sample size increases with dimensionality fairly rapidly, but still much slower than the exponential rate associated with the "curse of dimensionality." A qualitative explanation of why the increase is as rapid as it is has to do with the numerical optimization. In most statistical methods the size of the spurious structure associated with sampling fluctuations has to be comparable to that of the real underlying (population) structure to cause trouble. Here it need only be large enough (and numerous enough) to trap the numerical optimizer. Nevertheless, the sample size requirements are seen to be fairly modest for a search dimension of $q \leq 10$. For large samples, search dimensionalities of up to $q = 15$ or larger are possible (see Sec. 6.2).

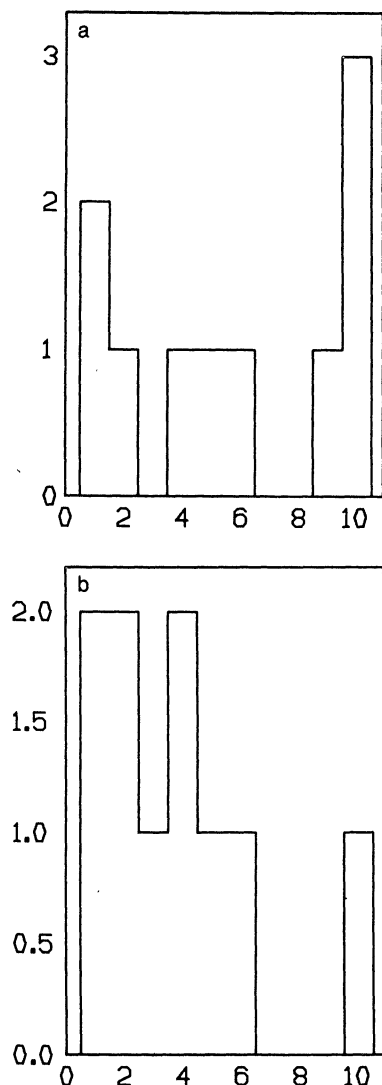


Figure 4. Iteration at Which the True Population Structure Was Found for $p = 15$: (a) $N = 600$; (b) $N = 999$.

7.2 Needle in a Haystack

The population for this example is again a Gaussian mixture. In this case, however, the two components of the mixture have the same location but different covariance structure. A random sample of size 175 is drawn from a 10-dimensional standard normal. Added to this is a sample of 25 observations that are standard normal in a 4-dimensional subspace (through the origin) and spherical normal with covariance matrix .0025 times the identity matrix in the 6-dimensional complement subspace. As in the previous example the data are scaled so that this structure is not reflected in the covariances of the combined data. The problem is to discover the presence of the small (25 observation) 4-dimensional needle in a 10-dimensional haystack, in analogy with finding a 1-dimensional needle in a 3-dimensional haystack.

To this end the 1-dimensional projection pursuit algorithm was applied to these data. This problem is difficult owing to high dimensionality (of the haystack) and the small cardinality of the needle. On the other hand the needle is visible (and thus can be found) in any projection that is orthogonal to its four-dimensional subspace. Figure 5a shows such a projection of these data.

Figure 5b shows the distribution of the iteration number at which the projection pursuit algorithm found the needle in 10 random samples from this Gaussian mixture. As in the previous example the determination of this was unambiguous. The results indicate that the algorithm was fairly well able to find a needle of this size. When the size of the needle is increased to 40 observations (out of 200), the algorithm always found it on the first iteration.

7.3 States Data

The data for this example are seven summary statistics associated with each of the 50 United States (Becker and

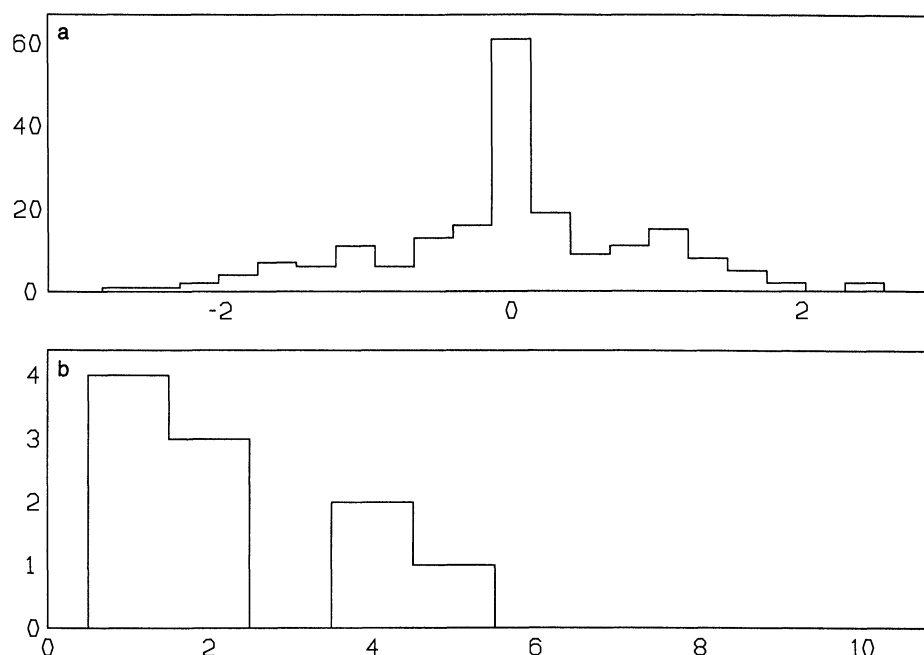


Figure 5. Needle in a Haystack: (a) a true solution projection; (b) iteration number at which a true solution was found.

Chambers 1984). Table 1 describes each of the seven variables. Two-dimensional projection pursuit was applied to this data set. The results of the simulation study (sec. 7.1) indicate that the sample size ($N = 50$) is small for a projection pursuit in seven dimensions. The eigenexpansion of the correlation matrix shows that 92% of the (auto-scaled) data variance is captured by the first four eigenvalues. Therefore, we restrict the projection pursuit to the subspace spanned by the first principal components ($q = 4$). Owing to the small sample size, it is unlikely that we will be able to detect very fine structure, so the order of the Legendre expansion J [see (16)] was set to $J = 2$.

To get an idea of the significance of the resulting solutions the identical procedure was applied to (different) random samples of size 50 drawn from a seven-dimensional standard normal distribution. An ordered list of the solution projection indices for 20 such (null) runs is given in part a of Table 2.

When applied to the states data, four iterations of two-dimensional projection pursuit produced solution projection indexes of .19, .08, .025, and .023, respectively. When referenced to the (estimated) null distribution represented in Table 2 (part a), only the first value is seen to appear significant (at 5%). Table 3 presents the (varimax rotated) variable loadings (α and β) associated with the linear com-

binations defining this solution plane. Figure 6 shows the data projected onto the solution plane.

Viewing Figure 6 shows that the data appear to divide into two clusters mainly along the vertical direction. In addition, two outliers are seen in the upper right corner. A smaller cluster of 12 states seems fairly well separated from the larger group of 38 states. Table 1 and 3 show that the dominant loadings associated with the vertical direction (β) involve income, high-school graduation rate, and (negatively) illiteracy. This index seems to divide the states into two fairly distinct groups. The horizontal axis (α) is dominated by (negative) population, (negative) life expectancy, and (positive) homicide rate. Thus, increasing values along the horizontal direction involve generally lower population and life expectancy, and increasing homicide rate. The states in the upper cluster generally have a lower value of the horizontal index than those in the lower one, with the dramatic exception of the two outliers in the upper right corner.

Table 4 lists the states of the smaller cluster in decreasing

Table 1. Measurement Variables for the States Data

Y_1	population estimate as of July 1, 1975
Y_2	average income (1974)
Y_3	illiteracy rate (1970)
Y_4	life expectancy (1969–1971)
Y_5	homicide rate (1976)
Y_6	high-school graduation rate (1970)
Y_7	average number of days below freezing temperature (1931–1960) in capital or large city

Table 2. Null Projection Index Distributions for the Data Examples

(a) $p = 7, q = 4, N = 50, J = 2$										
.022,	.025,	.028,	.028,	.030,	.030,	.030,	.030,	.030,	.033	
.033,	.038,	.038,	.045,	.058,	.060,	.060,	.065,	.065,	.098	
(b) $p = 10, q = 10, N = 392, J = 6$										
.043,	.045,	.048,	.050,	.050,	.053,	.053,	.053,	.053,	.055	
.055,	.055,	.058,	.058,	.060,	.060,	.060,	.063,	.070,	.070	
(c) $p = 13, q = 13, N = 506, J = 6$										
.043,	.043,	.043,	.045,	.045,	.045,	.045,	.045,	.045,	.048	
.048,	.048,	.048,	.050,	.053,	.053,	.053,	.053,	.055,	.063	

NOTE: Projection index (of order J) values obtained by running two-dimensional projection pursuit in the subspace spanned by the largest q principal components, on 20 random samples of size N , drawn from a p -dimensional standard normal.

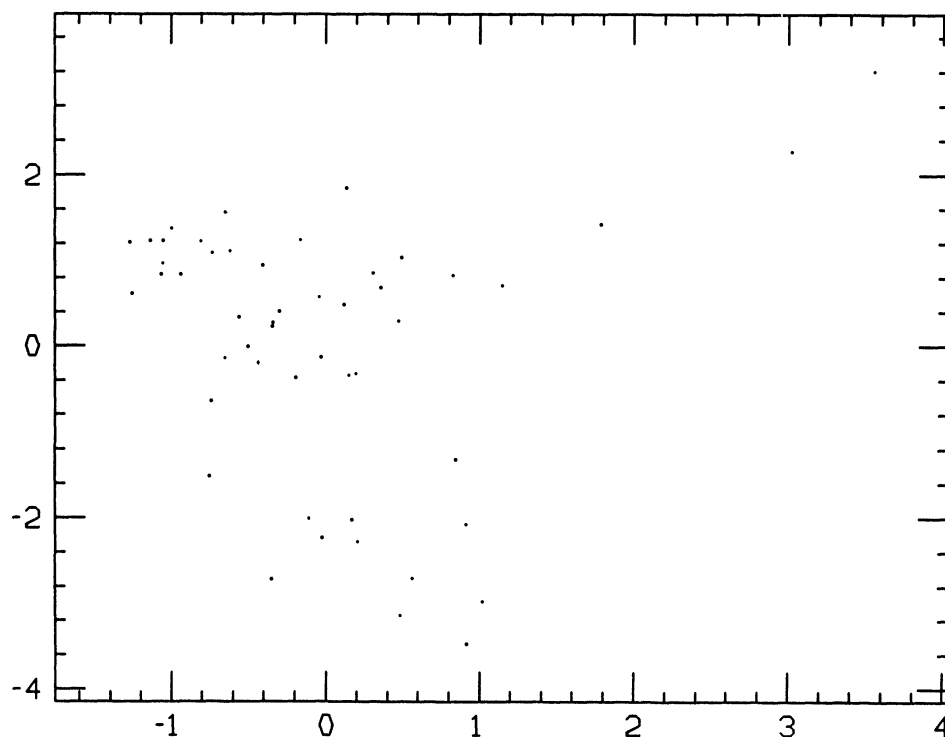


Figure 6. States Data: Solution Data Plot.

value of the vertical index. The extreme outlier in the upper right corner is Alaska whereas the other outlying point (closest to it) corresponds to Nevada.

7.4 Automobile Data

This data set consists of 10 characteristics associated with 392 automobile models sold in the United States and reported in *Consumer Reports* from 1972 to 1982 (Donoho and Ramos 1982). Table 5 lists the 10 variables. The last three are dummy variables for the manufacturing origin of the automobile. These dummy variables have only two distinct values whereas the second variable has only five distinct values (the value three was encoded for rotary engine automobiles). Following the discussion in the last part of Section 6.3, we Gaussianize these variables after randomly ordering all observations corresponding to the same value. Part b of Table 2 lists in order of ascending value the (null) projection index values obtained by 2-dimensional projection pursuit on 20 random samples of size $N = 392$ drawn from a $p = 10$ -dimensional standard normal distribution.

Six iterations of 2-dimensional projection pursuit on the automobile data produced projection index values of .31, .15, .13, .065, .11, and .088, respectively. All but the smallest value seem significant. The solutions corresponding to the largest two projection index values are presented in

Table 6 and Figure 7. Table 6 shows the (varimax rotated) loadings (α and β) defining the two solution planes. Figure 7a shows the data projected onto the first solution plane, and Figures 7b and 7c show the adjusted and original data plots corresponding to the second solution (see Sec. 6.5).

The first solution exhibits a strong clustering along the vertical index (approximately twice engine size minus fuel inefficiency) especially for moderate values of the horizontal index (approximately engine size minus weight). The second solution displays a distinctly trimodal distribution. Note that this structure is not a direct reflection of the clustering shown in the first solution owing to the structure removal.

7.5 Boston Neighborhood Data

This compilation of census data (Harrison and Rubinfeld 1978) is on its way to becoming a standard test bed for multivariate procedures. It was published in its entirety in Belsley, Kuh, and Welsch (1980). Each observation is a neighborhood (standard metropolitan statistical area) in the Boston area. Associated with each of these 506 census

Table 3. First Projection Pursuit Solution for the States Data

Solution 1: Projection index = .19	
α	.52, .26, .08, -.60, .41, .16, .31
β	-.05, .73, -.30, .10, .0, .58, .16

Table 4. States Composing the Smaller Cluster Associated With Lower Values of the Vertical Axis, Listed in Descending Values of the Vertical Coordinate

New Mexico	-1.32	North Carolina	-2.29
Texas	-1.52	Alabama	-2.72
Tennessee	-2.01	Arkansas	-2.72
West Virginia	-2.03	South Carolina	-2.98
Georgia	-2.08	Louisiana	-3.15
Kentucky	-2.24	Mississippi	-3.48

Table 5. Measurement Variables for Automobile Data

Y_1	gallons per mile (fuel inefficiency)
Y_2	number of cylinders in engine
Y_3	size of engine (cubic inches)
Y_4	engine power (horse power)
Y_5	automobile weight
Y_6	acceleration (time from 0 to 60 mph)
Y_7	model year
Y_8	American (0/1)
Y_9	European (0/1)
Y_{10}	Japanese (0/1)

tracts are 13 summary statistics that form the variables associated with each observation. Table 7 lists the quantities that compose the variable set.

These data are well known to contain striking structure, much of which is exhibited in the coordinate marginal distributions. Following the discussion in Section 6.3 we removed the most obvious of this structure (extreme skewness in Y_1 and Y_{11}) by the transformations $Y'_1 = \log(Y_1)$ and $Y'_{11} = \log(.4 - Y_{11})$.

As with the previous examples, we first obtain an estimate for the null distribution of projection index values by running 2-dimensional projection pursuit on 20 random samples of size 506 from a 13-dimensional standard normal distribution. An ordered list of the values so obtained is shown in part c of Table 2. Note that none of the 20 values is greater than .07. Running 10 iterations of 2-dimensional projection pursuit on the Boston neighborhood data produced solution projection index values of .69, .51, .40, .25, .34, .26, .31, .20, .22, and .10, respectively. Clearly, all of these values but the last are highly significant when referenced to the null distribution (Table 2, part c). As the high projection index values indicate, all of these views (but the last) exhibit striking structure. In the interest of brevity only the first five are shown. Table 8 lists the solution linear combinations, $\alpha\beta$, defining the solution planes for each of the five solutions.

Figure 8 shows the data projected onto the first solution plane. Figures 9–12 show both the adjusted and original data projections for each of the subsequent solutions. The first solution shows the data dividing into two groups on the second (“big lots”) variable. Even after this structure is removed, the adjusted plots of the subsequent solutions show that there is considerable (additional) clustering and other nonlinear effects. In this case, projection pursuit has provided a great many views with which to begin exploring and understanding the data.

Table 6. First Two Projection Pursuit Solutions for the Automobile Data

Solution 1: Projection index = .31	
$\alpha =$	-.72, -.08, .56, .30, -.20, .10, -.13, .00, .00, .02
$\beta =$	.00, -.01, .91, .14, -.39, .08, -.01, -.03, -.02, -.01
Solution 2: Projection index = .15	
$\alpha =$	-.10, .18, -.70, -.15, .22, -.06, -.17, -.23, .45, -.32
$\beta =$	.03, .09, .00, -.30, .53, -.01, -.10, .34, .19, -.69

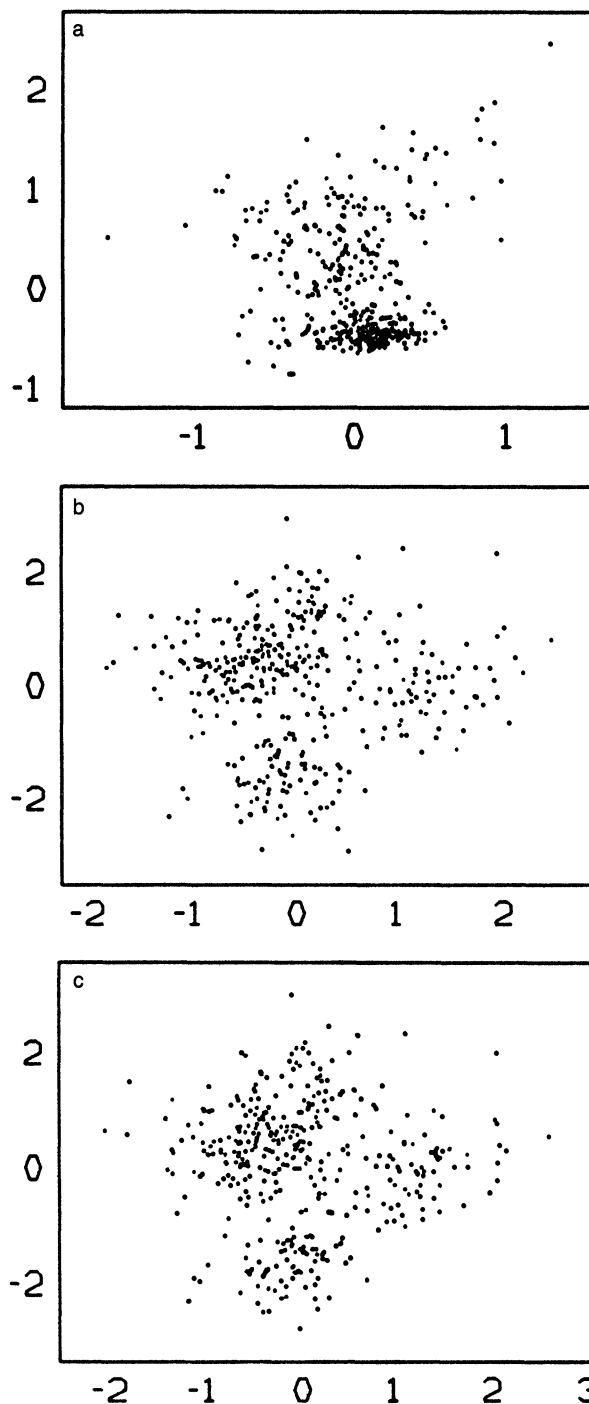


Figure 7. Car Data: (a) First Solution, Data Plot; (b) Second Solution, Adjusted Plot; (c) Second Solution, Data Plot.

8. DISCUSSION

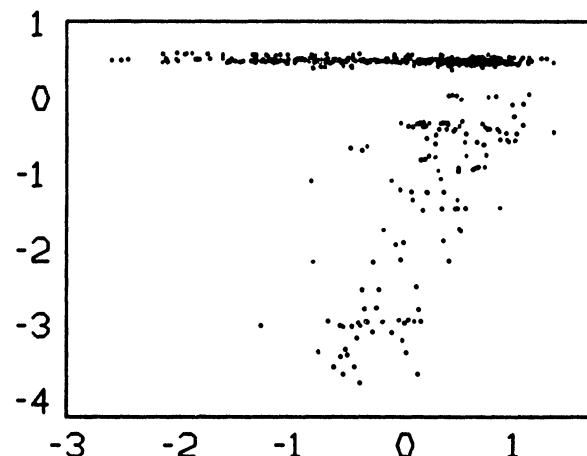
The examples of the previous section are intended to illustrate that the exploratory projection pursuit procedures developed in the earlier sections can effectively discover nonlinear data structuring in fairly high dimensionality with practical sample sizes. The aspects contributing to this are a projection index that measures nonnormality in the main body of the distribution rather than in the tails, an optimization algorithm that combines course stepping

Table 7. *Neighborhood Variables Composing the Boston Housing Data*

Y_1	log (per capita crime rate)
Y_2	fraction of land zoned for big lots
Y_3	fraction of nonretail business land
Y_4	(nitrogen oxide concentration) ² (pphm) ²
Y_5	(average number of rooms) ²
Y_6	fraction of owner-occupied units built before 1940
Y_7	log (weighted distances to five employment centers)
Y_8	log (index of access to radial highways)
Y_9	full-value property-tax rate
Y_{10}	pupil-teacher ratio
Y_{11}	log [.4 - (fraction black population - .63) ²]
Y_{12}	log (fraction of lower status population)
Y_{13}	log (median value of owner-occupied homes)

followed by gradient-directed optimization, and an effective technique for structure removal. This algorithm is also much faster than the Friedman and Tukey (1974) implementation owing to the superior optimization procedure but, much more important, to the rapidly computable form of the projection index (and its derivatives) using the (Legendre polynomial) moment expansion. This should help make the method more practical to those with fairly modest computing resources.

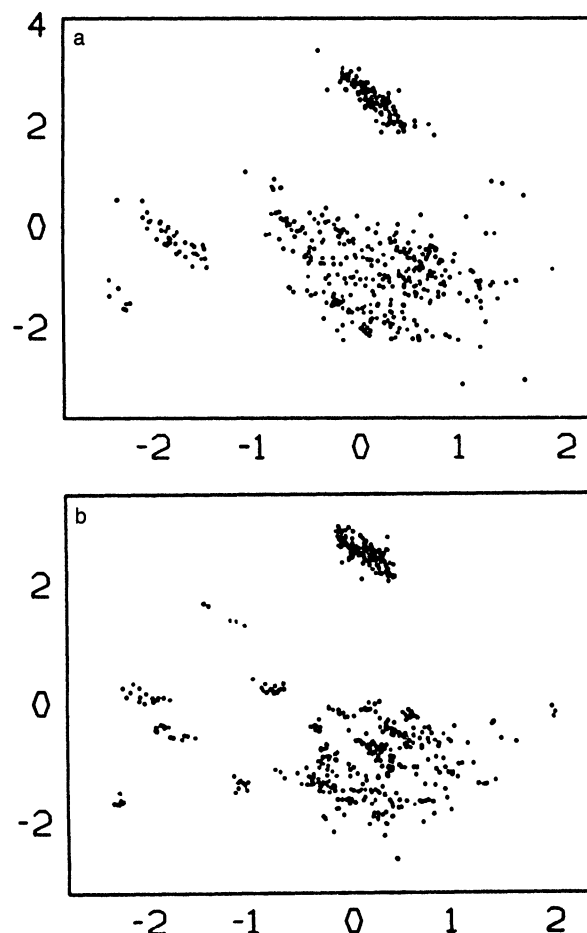
For purposes of data exploration (as opposed to density estimation) the two-dimensional projection pursuit procedure is likely to be more informative than the one-dimensional algorithm. This is due to the increased richness of structure that can be portrayed in a single picture with two dimensions. Both algorithms have the ability to uncover arbitrary nonlinear structure in the high-dimensional data density. The two-dimensional procedure, however, can often present the information in a format that is more easily visualized and interpreted. Features of the density such as bimodality (clustering), extreme skewness, sharp peaking, and some kinds of discontinuities can be seen in both one- and two-dimensional projections. In two di-

Figure 8. *Boston Neighborhood Data, First Solution: Data Plot.*

mensions, however, one can additionally see trimodality (where the modes do not align along a common axis) as well as arbitrary nonlinear relationships between the chosen variables (linear combinations) determining the projection plane. The principal disadvantage of the two-dimensional procedure is computational. In addition, for the special case of small data sets in which the data structuring is almost completely one-dimensional, two-dimensional

Table 8. *First Five Projection Pursuit Solution Planes for Boston Neighborhood Data*

<i>Solution 1: Projection index = .69</i>	
α	.13, -.41, -.50, -.24, -.04, -.02, .20, .26, .24, -.04, -.50, .14, .19
β	.0, .996, .05, -.02, .0, .02, .02, .04, -.04, -.02, .0, .01, .01
<i>Solution 2: Projection index = .51</i>	
α	.12, .23, -.76, -.06, .05, .01, .04, .09, .37, .41, -.09, .10, .12
β	.17, .08, .0, .16, .03, -.13, -.08, .40, .83, .24, .06, -.01, -.11
<i>Solution 3: Projection index = .40</i>	
α	-.21, -.23, -.25, .16, -.19, -.09, .38, -.31, .70, .0, .04, .08, -.15
β	.17, -.44, -.79, .07, -.03, -.02, .0, .10, .29, -.23, .03, -.02, .03
<i>Solution 4: Projection index = .25</i>	
α	-.41, .18, -.23, .22, -.15, .66, .18, -.01, .30, .10, .08, .10, .30
β	.02, -.09, .13, .07, -.45, .0, .22, -.07, .05, .07, .09, .02, .84
<i>Solution 5: Projection index = .34</i>	
α	-.03, -.12, -.60, -.01, -.08, .10, -.42, .05, .06, .57, -.05, -.29, -.12
β	.25, .06, -.55, .71, .0, -.02, .17, -.07, -.29, .0, -.01, .0, .02

Figure 9. *Boston Neighborhood Data, Second Solution: (a) Adjusted Plot; (b) Data Plot.*

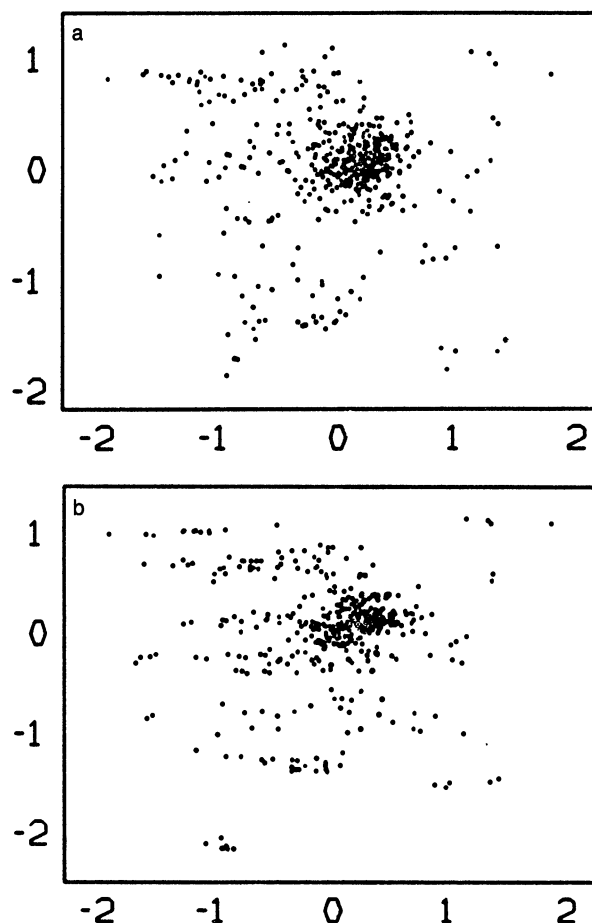


Figure 10. Boston Neighborhood Data, Third Solution: (a) Adjusted Plot; (b) Data Plot.

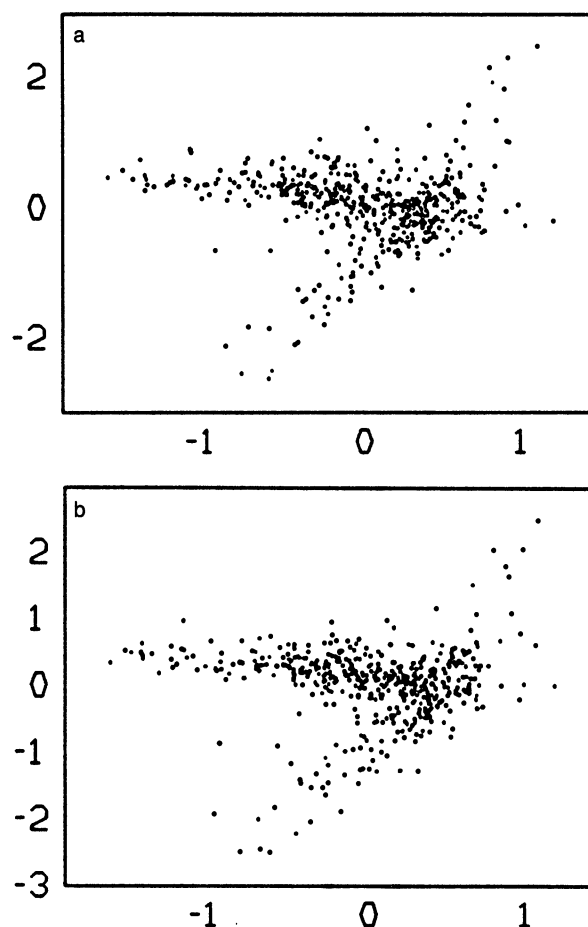


Figure 11. Boston Neighborhood Data, Fourth Solution: (a) Adjusted Plot; (b) Data Plot.

projection pursuit can have less power than the one-dimensional procedure.

A powerful aid in interpreting the output of this projection pursuit procedure would be a means for connecting the various solution plots (views of the data) so that particular observations or groups of observations in one plot could be identified in the other views. One could then easily identify hierarchical clustering as well as many more types of complex structure from the several views provided by the different projection pursuit solutions. With a color terminal supporting dynamic graphics one could use the color m and n plotting technique (McDonald 1984) to great advantage. With a black and white terminal (again supporting dynamic graphics) the scatterplot brushing techniques (Becker and Cleveland 1985) would be very useful. In the absence of either of these alternatives the static m and n plotting technique (Diaconis and Friedman 1983) might be of some use. In the absence of these powerful graphical techniques the varying views can be connected by more laborious methods involving isolating points in one view and plotting their positions in the other views.

For the solution projections presented in the previous section the structure (nonnormality) was fairly striking and easily recognized from simple point (scatter) plot representations of the projected densities. This is not always the case. Human visual perception is not very good at

distinguishing varying densities of points. Only the (local) presence or absence of points is easily recognized. Fairly striking density variation is often difficult to see in a scatterplot. For this reason it is often helpful to view a graphical representation of a smoothed density estimate of the projected solution distributions. Structure easily missed in a scatterplot can be quite evident in such displays. There are several good methods for (smooth) density estimation in two dimensions (Scott 1985). These density estimates can be represented graphically as contour plots, color relief maps, or isometric projections of the three-dimensional surface of the (estimated) density versus the projection coordinates.

As seen in the examples, exploratory projection pursuit solutions can sometimes both discover interesting nonlinear effects and suggest straightforward interpretations for them. More often the interpretation of the discovered structure is elusive and requires a great deal of study and further investigation. In this sense applying projection pursuit to a data set can often raise more questions than it (immediately) answers. This is the primary purpose of an exploratory technique. The discovery of strong (nonlinear) effects will usually cause the analyst to look harder at the data with hopefully a corresponding gain in insight and understanding.

FORTRAN programs implementing the exploratory

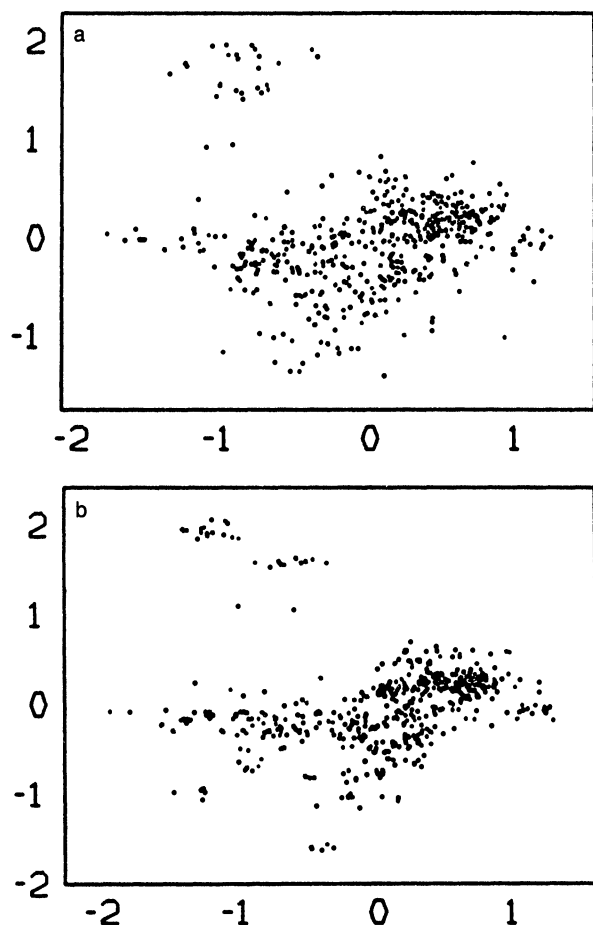


Figure 12. Boston Neighborhood Data, Fifth Solution: (a) Adjusted Plot; (b) Data Plot.

projection pursuit procedures described are available from the author.

[Received December 1985. Revised July 1986.]

REFERENCES

- Anderson, T. W. (1958), *Introduction to Multivariate Analysis*, New York: John Wiley.
- Becker, R. A., and Chambers, J. M. (1984), *S: An Interactive Environment for Data Analysis and Graphics*, Belmont, CA: Wadsworth.
- Becker, R. A., and Cleveland, W. S. (1985), "Brushing a Scatter Plot Matrix: High-Interaction Graphical Methods for Analyzing Multidimensional Data, Statistical Research Report 7, Murray Hill, NJ: AT&T Bell Laboratories.
- Bellman, R. E. (1961), *Adaptive Control Processes*, Princeton, NJ: Princeton University Press.
- Belsley, D. A., Kuh, E., and Welsh, R. E. (1980), *Regression Diagnostics*, New York: John Wiley.
- Buja, A., and Asimov, D. (1985), "Grand Tour Methods: An Outline, in *Proceedings of the 17th Symposium on the Interface of Computer Science and Statistics*, New York: Elsevier.
- Chen, Z., and Li, G. (1981), "Robust Principal Components and Dispersion Matrices Via Projection Pursuit, Report PJH-8, Harvard University, Dept. of Statistics.
- Devlin, S. J., Gnanadesikan, R., and Kettenring, J. R. (1981), "Robust Estimation of Dispersion Matrices and Principal Components," *Journal of the American Statistical Association*, 76, 354-362.
- Diaconis, P., and Freedman, D. (1984), "Asymptotics of Graphical Projection Pursuit," *The Annals of Statistics*, 12, 793-815.
- Diaconis, P., and Friedman, J. H. (1983), "M and N Plots," in *Recent Advances in Statistics*, eds. M. H. Rizvi, J. S. Rustagi, and D. Siegmund, New York: Academic Press.
- Donoho, A. W., Donoho, D. L. and Gasko, M. (1986), "MACSPIN™: Graphical Data Analysis Software," technical report, Austin, TX: D² Software, Inc. (P.O. Box 9546).
- Donoho, D. L., and Ramos, E. (1982), "PRIMDATA: Data Sets for Use with PRIMH," technical report, Harvard University, Dept. of Statistics.
- Fisher, M. A., Friedman, J. H., and Tukey, J. W. (1975), PRIM-9: An Interactive Multidimensional Data Display and Analysis System, Report SLAC-PUB-1408, Stanford, CA: Stanford Linear Accelerator Center.
- Friedman, J. H., Stuetzle, W., and Schroeder, A. (1984), "Projection Pursuit Density Estimation," *Journal of the American Statistical Association*, 79, 599-608.
- Friedman, J. H., and Tukey, J. W. (1974), "A Projection Pursuit Algorithm for Exploratory Data Analysis," *IEEE Transactions on Computers*, Ser. C, 23, 881-889.
- Gill, P. E., Murray, W., and Wright, M. H. (1981), *Practical Optimization*, London: Academic Press.
- Harrison, D., and Rubinfeld, D. L. (1978), "Hedonic Housing Prices and the Demand for Clean Air," *Journal of Environmental Economics and Management*, 5, 81-102.
- Huber, P. J. (1981), "Projection Pursuit," Research Report PJH-6, Harvard University, Dept. of Statistics.
- (1985), "Projection Pursuit," *The Annals of Statistics*, 13, 435-475.
- Jones, M. C. (1983), "The Projection Pursuit Algorithm for Exploratory Data Analysis," unpublished Ph.D. dissertation, University of Bath, School of Mathematics.
- Kennedy, W. J., and Gentle, J. E. (1980), *Statistical Computing*, New York: Marcel Dekker.
- McDonald, J. A. (1984), "Interactive Graphics for Data Analysis," Report ORION 11, Stanford University, Dept. of Statistics.
- Mosteller, F., and Tukey, J. W. (1977), *Data Analysis and Regression*, Reading, MA: Addison-Wesley.
- Scott, D. W. (1985), "Average Shifted Histograms: Effective Nonparametric Density Estimators in Several Dimensions," *The Annals of Statistics*, 13, 1024-1040.
- Switzer, P. (1970), "Numerical Classification," in *Geostatistics*, ed. D. F. Merriam, New York: Plenum Press, pp. 31-43.
- Tukey, J. W., and Tukey, P. A. (1985), "Computer Graphics and Exploratory Data Analysis: An Introduction," in *Proceedings of the Annual Meeting of National Computer Graphics Association*, Alexandria, VA: National Computer Graphics Association.
- Tukey, P. A., and Tukey, J. W. (1981), "Preparation: Prechosen Sequences of Views," in *Interpreting Multivariate Data*, ed. V. Barnett, London: John Wiley, pp. 189-213.