



KERNFORSCHUNGSANLAGE JÜLICH
GESELLSCHAFT MIT BESCHRÄNKTER HAFTUNG

Institut für Festkörperforschung

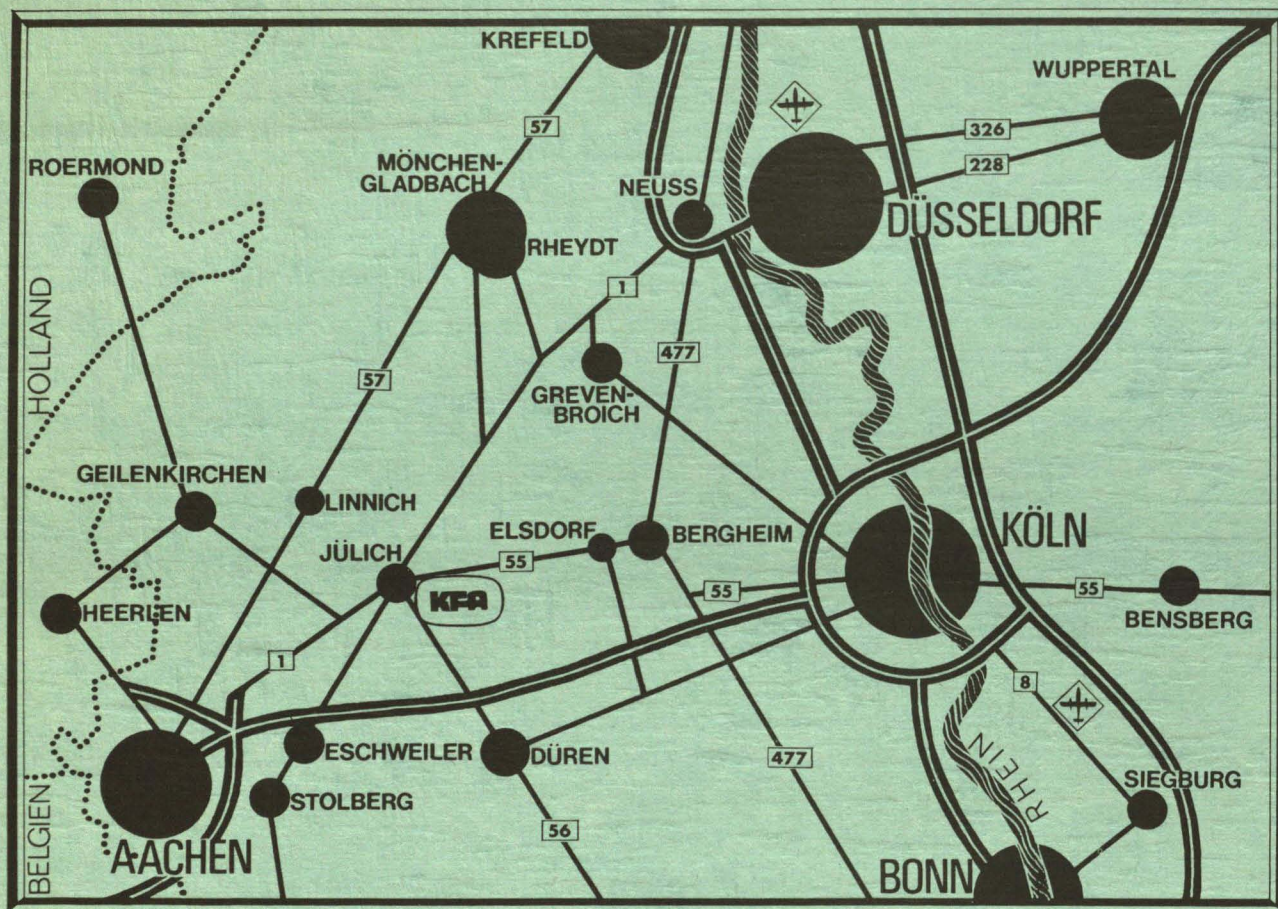
**Sequential Bayes Estimation Algorithm
with Cubic Splines on Uniform Meshes**

by

F. Hoßfeld, K. Mika and E. Pleßer-Walk

Jül - 1249
November 1975

Als Manuskript gedruckt



Berichte der Kernforschungsanlage Jülich – Nr. 1249

Institut für Festkörperforschung Jülich - 1249

Im Tausch zu beziehen durch: ZENTRALBIBLIOTHEK der Kernforschungsanlage Jülich GmbH,
Jülich, Bundesrepublik Deutschland

Sequential Bayes Estimation Algorithm with Cubic Splines on Uniform Meshes

by

F. Hoßfeld, K. Mika and E. Pleßer-Walk

Summary

After outlining the principles of some recent developments in parameter estimation, a sequential numerical algorithm for generalized curve-fitting applications is presented combining results from statistical estimation concepts and spline analysis. Using Bayes recursive estimation theory cubic cardinal splines are applied as a flexible linear model, which is designed as a very suitable tool for scientific data-processing problems. Due to its recursive nature, the algorithm can be used most efficiently in online experimentation: on the one hand the resulting experimental information can be easily updated after each new observation, on the other hand the recursive procedure allows to control the experiment by means of sequential experimental design techniques. In addition, the algorithm is numerically stable since matrix inversions are avoided. Using computer simulated and experimental data, the efficiency and the flexibility of this sequential estimation procedure is extensively demonstrated.

Contents

1. Introduction	1
2. Principles of Estimation	3
3. Bayes Estimation	9
4. Sequential Bayes Estimation	17
5. Nilson's Representation of Cubic Cardinal Splines on Uniform Meshes	26
5.1 Nonperiodic cardinal splines	28
5.2 Periodic cardinal splines	36
6. Applications	44
6.1 Curve fitting examples based on nonperiodic splines	44
6.2 An example based on periodic splines	48
6.3 Deconvolution of experimental data	51
6.3.1 The effect of increased smearing	54
6.3.2 Dependence on the number of knots	54
6.4 Integral equations of non-convolution type	56
7. Conclusions	58
Acknowledgment	59
Appendix	60
References	65
Figure Captions	69
Figures	

1. Introduction

We assume that the m -dimensional sample vector z consists of statistically independent observations of a stochastic process at the experimental variables x_i , $i=1, \dots, m$, where the statistical behaviour of the stochastic variables z_i is described by some specific probability distribution $p(z_i)$, and the variance of z_i may be given by σ_i^2 . The problem which we are interested in and which frequently occurs in experimental data processing is stated quite generally by the fact that this experimental sample z is to be approximated by a function

$$(1.1) \quad \phi(x; \theta) = \int g(x, x') f(x'; \theta) dx',$$

according to some optimality criterion yielding a statistical estimate /1,2/ of the n components θ_v , $v=1, \dots, n$ of the parameter vector θ .

The kernel $g(x, x')$ is a given non-random function which may be determined either by the experimental device or by some basic theoretical aspects /3/. In the special case where the kernel is equal to the δ -function, the approximation procedure is equivalent to a simple curve fitting problem. The function $f(x; \theta)$ may represent the fundamental theoretical relationship, which may be the main interesting information to be confirmed by experiment. The functional form of the model function $f(x; \theta)$ is assumed to be known explicitly, while the unknown parameter vector θ has to be explored from the experimental data by an appropriate estimation method /4/.

The problem therefore is to select an estimator for these parameters optimal according to a certain criterion. In this report we make an attempt to solve this problem by combining new results from statistical estimation theory and numerical analysis. Especially, we try to realize an effective union of the

following concepts:

- (i) Bayes estimation,
- (ii) sequential algorithm,
- (iii) cubic splines as linear model.

In chapters 2 and 3 the fundamentals of estimation theory, in particular of maximum likelihood and least squares, are outlined and the principles of Bayes methods /4/, the concept of cost functions and prior information, are introduced. Under quite general assumptions the maximum a-posteriori estimator is deduced as the optimal solution taking account of the available a-priori information in a most appropriate manner. In chapter 4 the characteristics of sequential estimation algorithms /5,6/ are presented yielding the recursive equations which are equivalent to a special case of the well-known Kalman-Bucy filter equations /4/. In chapter 5, Nilson's representation of cubic cardinal splines on uniform meshes /7/ is extensively discussed and adapted as an efficient tool to the requirements of sequential Bayes estimation algorithms for linear models. The spline formalism comprises nonperiodic cardinal splines with certain boundary conditions and periodic cardinal splines. In chapter 6 computer simulations are used to demonstrate the efficiency of the developed algorithm in different applications.

2. Principles of Estimation

Unfortunately, the problem of fitting a curve to a set of experimental data is often looked at from the point of view of mathematical approximation theory defining the "most general least-squares problem" as the job to find the minimum of the expression

$$(2.1) \quad Q = \sum_{i=1}^m w_i (z_i - \phi(x_i; \theta_1, \dots, \theta_n))^2$$

with respect to the parameter vector θ , where z_i , $i=1, \dots, m$ represents the set of experimental data at "positions" x_i and $\phi(x; \theta)$ is equivalent to some model function assumed to describe the course of the underlying process in some theoretical sense. The weights $w_i > 0$, $i=1, \dots, m$ may be chosen arbitrarily within approximation theory and are used to represent some kind of relevance of the respective experimental data.

Thus formulated this problem, of course, is just a special case of (nonlinear) optimization with no relationship to any statistical interpretation. Defining $z_i^+ = \sqrt{w_i} z_i$ and $\phi_i^+ = \sqrt{w_i} \phi(x_i; \theta)$, we obtain the minimization problem connected with (2.1) as a specialization of the L_p -approximation

$$(2.2) \quad \theta_L: \min_{\{\theta\}} \|z^+ - \phi^+\|_p,$$

which for $p=2$ coincides with minimizing the Euclidean norm and for $p=\infty$ results in the Tchebycheff approximation

$$(2.3) \quad \theta_T: \min_{\{\theta\}} \max_{\{i\}} |z_i^+ - \phi_i^+|.$$

The approximation procedure yields a parameter vector θ_L which is optimal in some postulated sense but which does not necessari-

ly take into account any statistical properties of the experimental data.

From the point of view of mathematical statistics the procedure to draw conclusions from experimental data on the quantitative behaviour of the underlying process is a problem of statistical inference /8/. With this respect, curve fitting is equivalent to the method of point estimation /1/. In terms of statistics, the set of experimental data $\{z_i, i=1, \dots, m\}$ is called a statistical sample obtained by experiment out of the sample space or population of the experimental process. The estimation procedure defines an estimator $\hat{\theta}$ which yields an estimate as a definite parameter vector, if the estimator is applied to a specific sample $\{z\}$. Thus, an estimator is a random variable, and we can assign to it a probability distribution $p(\theta|z)$, an expectation value, and so on.

The fact, that the experimental data as a statistical sample represent in general only a subset of the observation space, makes it impossible to take any estimate as the exact value describing completely the process under study. As a matter of fact, it is very unlikely that both values will coincide because of all kinds of chance and random effects. Nevertheless, in the long run, experimental data sampled at random should reflect the characteristics of the experimental system. However, the practice of experimentation usually cannot afford for "the long run"; decisions and inference have to be made in spite of limited experimental evidence.

It is important to know how to use the available experimental data in the best possible ways to infer the characteristics of the underlying process. Therefore general principles for statistical estimation have been developed /2,4,9,10,11/.

There are several properties that have turned out to be desirable for a procedure to show in order to represent a "good" estimator of some parameter or parameter vector involved

in the probabilistic description of an experiment. Four of these "optimal" properties are comprised by unbiasedness, consistency, efficiency, and sufficiency /9,10/.

If $\hat{\theta}$ is a definition of an estimator for the parameter θ , unbiasedness is described by the equivalence

$$(2.4) \quad E(\hat{\theta}) = \theta,$$

where E represents the operation of generating the expectation value of the random variable $\hat{\theta}$ /4/.

For an estimator $\hat{\theta}$ to be a "good" one, the sample estimator $\hat{\theta}$ should have a higher probability of being close to the true value θ the larger the sample size m . Statistics /11/ which have this property are called consistent estimators. More formally, the estimator $\hat{\theta}$ is a consistent estimator of θ , if for any arbitrary positive ϵ

$$(2.5) \quad \lim_{m \rightarrow \infty} P(|\hat{\theta} - \theta| < \epsilon) = 1,$$

where P means probability.

It should be noticed that there is an important distinction between unbiasedness and consistency: whereas unbiasedness is a fixed-sample property, consistency is an asymptotic property.

If we restrict our considerations to unbiased estimators, efficiency of an estimator $\hat{\theta}_e$ means that the (generalized) variance of this estimator of θ is minimal with respect to all other unbiased estimators $\hat{\theta}$ of θ based on the same sample.

Sufficiency is a concept of great importance in estimation theory. An estimator $\hat{\theta}$ is called a sufficient estimator of the parameter θ if $\hat{\theta}$ contains all of the information available in

the data about the value of θ . A sufficient estimator $\hat{\theta}$ is a "best" estimator of θ in the sense that $\hat{\theta}$ cannot be improved by taking into account any other aspects of the sample data not yet included in $\hat{\theta}$ itself. A formal definition of sufficiency is given by the following. Let $\hat{\theta}'$ be any other estimator based on the sample. Consider the conditional probability distribution of $\hat{\theta}'$ given the value of $\hat{\theta}$. If then the conditional distribution of $\hat{\theta}'$ given $\hat{\theta}$ does not depend in any way on the value of θ , $\hat{\theta}$ is a sufficient estimator.

It should be stated that it is generally not possible to find a "best" estimator for any particular situation. The methods of estimation theory will not determine "best" estimators, but they represent alternatives possessing some of the desirable properties; it is up to the statistician or experimentalist to decide if these properties are satisfactory for the specific problem at hand. Many aspects will enter into this decision like complexity or time consumption with respect to the computing of the individual estimate.

One of the most important estimation methods is based on the principle of maximum likelihood which provides a general strategy for such decisions. Let us suppose that a random variable z is controlled by a probability distribution that depends upon some parameter θ , only. The analytic form of the density function is assumed to be known, whereas the value of θ is unknown. By experiment a statistical sample $\{z_i, i=1, \dots, m\}$ of m independent observations may be obtained. Let $L(z_1, \dots, z_m | \theta)$ be the so-called likelihood function of this specific sample which is equivalent to the probability density function for z given θ . In general, the likelihood function related with a definite sample set $\{z_i\}$ will be different for each possible value of θ . Then: the principle of maximum likelihood requires to accept as the estimate that particular value of θ , if it exists, which makes $L(z_1, \dots, z_m | \theta)$ a maximum value. This maximum-likelihood estimate is that parameter value which "would have made the sample actually obtained have the highest probability" /9/.

As an example for a maximum-likelihood estimator, let us suppose that a statistical sample $\{z_i, i=1, \dots, m\}$ is drawn from a normally distributed population with known covariance matrix V_z where the underlying theoretical model is assumed to be linear in the n -dimensional parameter vector θ

$$(2.6) \quad z = H\theta,$$

where z is the m -dimensional random vector and H the $m \times n$ measurement or modulation matrix, $m \geq n$. Under these assumptions, the corresponding likelihood function takes the form /4/:

$$(2.7) \quad L(z|\theta) = \{(2\pi)^m |V_z| \}^{-1/2} \exp \left\{ -\frac{1}{2} (z-H\theta)^T V_z^{-1} (z-H\theta) \right\},$$

where $|\cdot|$ means the determinant. The maximum-likelihood estimator is defined by

$$(2.8) \quad \hat{\theta}_{ML} : \underset{\{\theta\}}{\text{Max}} L(z|\theta),$$

which yields

$$(2.9) \quad \hat{\theta}_{ML} = (H^T V_z^{-1} H)^{-1} H^T V_z^{-1} z.$$

The principle of maximum likelihood is quite useful in theoretical statistics and applications, since general methods exist for finding the value of θ that makes the likelihood function of a specific sample a maximum. This maximum-likelihood estimator turns out to have important sufficiency and large sample properties. Further, maximum-likelihood estimators show certain invariances /9/. On the other hand, maximum-likelihood estimators are not always unbiased.

Nevertheless, the maximum-likelihood principle has proved a most important concept in mathematical statistics theory and appli-

cation; it is in the spirit of heuristics running throughout the methods of statistical inference. Its point of view is that good methods should make empirical results likely.

Whereas the method of maximum-likelihood requires the knowledge of the explicit functional form of the probability density which is hard to be realized in a large number of experimental situations, the well-known method of least squares requires only the existence of the covariance matrix V_z corresponding to z /12/. It is important to note that the related estimator $\hat{\theta}_{LS}$ obtained from the principle of least squares, i.e. from minimizing

$$(2.10) \quad Q = (z - H\theta)^T V_z^{-1} (z - H\theta)$$

with respect to θ , is equivalent to (2.9), where the covariance matrix of the estimator $\hat{\theta}_{LS}$ is represented by

$$(2.11) \quad V_{\hat{\theta}_{LS}} = (H^T V_z^{-1} H)^{-1}.$$

Least-squares estimators turn out to be unbiased; finally, they are optimal in the sense that they provide estimates with minimum variance /4/.

The estimation methods outlined so far are based on the probability distribution for z with parameter θ and on the experimentally obtained statistical sample as the only information. Let us now turn to a concept which allows the implication of a-priori knowledge and of aspects of costs generated by wrong conclusions drawn from any sample.

3. Bayes Estimation

The excursion into the field of the "classical" principles of estimation theory based on sample evidence only was undertaken in order to get a better understanding of a second approach to statistical inference and decision which is different from the classical methods of maximum likelihood and least squares with respect to the formal utilization of other than sample information. In this approach, sample information is used together with other available information. The basis of this combination of information for inferential and decision-making procedures is provided by Bayes' theorem from which the general characterization of these methods has been adapted as "Bayesian" /8,13/.

Of course, the intention for Bayesian methods comes from the essential need to base inference and decision on any available information whether it might be of experimental or some other nature.

Erroneous decisions may be quite costly; therefore it will not be acceptable to ignore some kind of information on the process in question just because it is not considered objective in the classical sense.

In order to express the total available information in probabilistic terms, Bayesian statistics follows the "subjective" interpretation of probability /8/, where probability is interpreted as the "degree of believe". Since Bayesian statistics utilizes experimental evidence, represented by a statistical sample, and other information, the Bayesian approach must be understood as an extension of the classical procedures more than a replacement.

With respect to estimation theory the new aspect arising with the Bayesian approach consists in a probabilistic description of the parameters which are now considered as statistical variables, too, in the sense of subjective probability theory.

Thus, in addition to the so far used probability density of the process z with parameter vector θ , a new density function $p(\theta)$ for the parameter vector θ is introduced which describes the prior information about θ .

According to Bayes' theorem /8/:

$$(3.1) \quad p(\theta|z) = \frac{p(\theta) p(z|\theta)}{p(z)},$$

which can be derived from the equivalence of the joint probabilities of z and θ :

$$(3.2) \quad p(z, \theta) = p(z) p(\theta|z) = p(\theta) p(z|\theta),$$

the posterior probability function for the parameter vector θ after taking a sample z , $p(\theta|z)$, can be determined. This may be an extremely complex procedure.

The method of maximum likelihood has been remarkably successful in practice mainly because of its simple applicability: once $p(z|\theta)$ is known, straight-forward application of (2.8) yields the parameter estimate which is optimal in the previously discussed sense. In general, this simplicity can no longer be expected with Bayes estimation.

To establish an optimality criterion, a new concept has been introduced by means of functions which are known as loss or cost functions $C(\theta - \hat{\theta})$ /2,4/, where $\hat{\theta}$ is assumed to be the Bayes estimator which now depends on the choice of the cost function.

The meaning of the cost function is very realistic, and from its definition some necessary properties should be satisfied by any $C(\theta)$: If a "wrong" estimation is made this may cause some loss or costs; if the "correct" estimation θ is made, no

costs arise which leads to $C(0)=0$. It is common to postulate /2,4/ that:

$$(3.3) \quad C\left(\frac{1}{2} (\theta_1 + \theta_2)\right) \leq \frac{1}{2} \{C(\theta_1) + C(\theta_2)\}.$$

In addition to this convexity property, $C(\tilde{\theta})$ will be assumed to be symmetric:

$$(3.4) \quad C(\tilde{\theta}) = C(-\tilde{\theta}).$$

For the Bayes optimum estimator to be obtained here, another special condition concerning the probability density $p(\theta|z)$ will be needed: $p(\theta|z)$ should be symmetric with respect to a certain value θ_M (which then necessarily has to be the conditional mean value /4/ for given z):

$$(3.5) \quad p(\tilde{\theta}|z) = p(-\tilde{\theta}|z), \quad \tilde{\theta} = \theta - \theta_M.$$

The most general formulation of the Bayes optimum criterion is now stated as:

$$(3.6) \quad \hat{\theta} : \text{Min}_{\{\theta\}} \iint C(\theta' - \theta) p(\theta', z) d\theta' dz.$$

This double integral is referred to as the Bayes risk. It can also be written as:

$$(3.7) \quad \hat{\theta} : \text{Min}_{\{\theta\}} \int p(z) \{ \int C(\theta' - \theta) p(\theta'|z) d\theta' \} dz,$$

and, since the total integrand is always greater than or equal to zero, and $p(z)$ independent of θ , the alternative form results:

$$(3.8) \quad \hat{\theta} : \min_{\{\theta\}} \int C(\theta' - \theta) p(\theta' | z) d\theta',$$

i.e. the optimal Bayes estimator will minimize the conditional Bayes risk. This formulation is rather general to be useful for special applications. With assumptions (3.3) - (3.5), however, a large class of estimation problems can be solved as will be shown /4/. That means we require that the cost function $C(\tilde{\theta})$ be convex and symmetric, not necessarily continuous, and that the conditional probability density $p(\tilde{\theta}|z)$ be symmetric in $\tilde{\theta} = \theta - \theta_M$.

With:

$$B(\theta|z) = \int C(\theta' - \theta) p(\theta' | z) d\theta',$$

substituting $\theta' = (\theta_M + \theta'')$ and using the postulated equivalence of the conditional probabilities for θ' and θ'' /4/ we find:

$$(3.9) \quad B(\theta|z) = \int C(\theta'' + \theta_M - \theta) p(\theta'' | z) d\theta'' \quad (i)$$

$$= \int C(-\theta'' + \theta_M - \theta) p(\theta'' | z) d\theta'' \quad (ii)$$

$$= \int C(\theta'' - \theta_M + \theta) p(\theta'' | z) d\theta'' \quad (iii)$$

$$= \int \frac{1}{2} \{C(\theta'' + \theta_M - \theta) + C(\theta'' - \theta_M + \theta)\} p(\theta'' | z) d\theta''. \quad (iv)$$

Step (ii) follows from (3.5), step (iii) follows from (3.4), and (iv) is equal to $\frac{1}{2}\{(i)+(iii)\}$. If we write $\Delta = \theta_M - \theta$, then from the convexity condition (3.3) we have:

$$(3.10) \quad \frac{1}{2}\{C(\theta'' + \Delta) + C(\theta'' - \Delta)\} \geq C(\theta''), \text{ for all } \theta''.$$

Inequality (3.10) inserted into (3.9) yields:

$$B(\theta|z) \geq B(\theta_M|z).$$

In (3.3) we have also admitted the equal sign in order to allow the cost function to have piecewise linear (or constant) parts over its range of definition. But since (3.10) is valid for all θ ", we conclude that:

$$(3.11) \quad B(\theta|z) > B(\theta_M|z), \quad \text{for } \theta \neq \theta_M,$$

where we exclude the trivial case of a totally constant cost function.

We obtain the important result that the solution of (3.6), with assumptions (3.3)-(3.5), is given by:

$$(3.12) \quad \hat{\theta} = \theta_M.$$

This result can be put into another form, if a further condition is satisfied by the probability density $p(\theta|z)$:

$$(3.13) \quad p(\theta_M|z) > p(\theta|z), \quad \text{for } \theta \neq \theta_M,$$

which means that at the conditional mean value of θ with respect to $p(\theta|z)$, $p(\theta|z)$ also assumes its maximum value. Then, the result (3.12) as the solution to (3.6) is equivalent to the solution of the problem:

$$(3.14) \quad \hat{\theta} : \text{Max}_{\{\theta\}} p(\theta|z).$$

For the general class of cost functions satisfying (3.3), (3.4), and for the general class of posterior probability functions satisfying (3.5), (3.13), the optimal Bayes estimator is given by the maximum a-posteriori estimator of θ with respect to $p(\theta|z)$. It should be noted that (3.14) has a striking similarity to the maximum-likelihood estimator (2.8).

Formula (3.14) can be derived in a different way /4/, /14/, less restrictive because conditions (3.5) and (3.13) are no longer needed, but more restrictive in that a special type of cost function is chosen:

$$(3.15) \quad C(\tilde{\theta}) = \begin{cases} 0, & \|\tilde{\theta}\| < \epsilon \\ \frac{1}{\epsilon}, & \|\tilde{\theta}\| \geq \epsilon \end{cases} \quad \text{with } \epsilon \rightarrow 0.$$

This special cost function is based on the philosophy that any estimate which will not provide the true value of the parameter causes the same amount of cost no matter how large the deviation in fact might be. Since for the applications to be discussed here conditions (3.5) and (3.13) will be automatically fulfilled, the two derivations lead to the same result and both aspects can be assumed (see also the discussions in /4/, p. 182, and in /15/).

Since the probability density $p(z)$ is not dependent on θ , we may, instead of maximizing the a-posteriori density, maximize the joint probability density:

$$(3.16) \quad \hat{\theta} : \underset{\{\theta\}}{\text{Max}} p(\theta, z).$$

This formula will be used in the following.

From now on we shall always assume a linear model as in (2.6), and Gaussian distribution for the statistically independent data z_i with variances σ_i^2 , where the (diagonal) covariance matrix related with z is defined as $V_z = \text{Diag}(\sigma_i^2)$. We then have: $p(z|\theta) = L(z|\theta)$ with $L(z|\theta)$ from (2.7). We further assume that the a-priori distribution $p(\theta)$ is also Gaussian with mean μ and covariance matrix V_μ :

$$(3.17) \quad p(\theta) = \{(2\pi)^n |V_\mu|\}^{-\frac{1}{2}} \exp\{-\frac{1}{2}(\theta-\mu)^T V_\mu^{-1} (\theta-\mu)\}.$$

To maximize the joint probability density according to (3.16) is then identical with:

$$(3.18) \quad \hat{\theta} : \underset{\{\theta\}}{\text{Min}} \{ (\theta - \mu)^T V_{\mu}^{-1} (\theta - \mu) + (z - H\theta)^T V_z^{-1} (z - H\theta) \}.$$

The expression to be minimized is a quadratic and positive definite function $Q(\theta)$, which may be written as:

$$(3.19) \quad Q(\theta) = (\theta - \hat{\theta})^T V_{\theta}^{-1} (\theta - \hat{\theta}) + Q_{\min},$$

with

$$(3.20) \quad V_{\theta}^{-1} = V_{\mu}^{-1} + H^T V_z^{-1} H.$$

The matrix V_{θ} is the covariance matrix for the parameters, combining the a-priori and the measurement information:

$$(3.21) \quad V_{\theta} = E\{(\theta - E(\theta)) (\theta - E(\theta))^T\},$$

where the expectation values have to be taken with respect to $p(\theta, z)$.

The optimal Bayes estimator for the linear model is obtained from (3.18), yielding with (3.20):

$$(3.22) \quad \hat{\theta} = V_{\theta} \{ H^T V_z^{-1} z + V_{\mu}^{-1} \mu \}.$$

If no a-priori information is available, one has to put $V_{\mu}^{-1} = 0$ in eqs. (3.20) and (3.22), and the Bayes estimator reduces to the maximum-likelihood estimator (2.9). If available, prior information concerning the covariance matrix V_{μ} usually reduces

to the variances of the parameters, i.e. V_{μ} is assumed to be diagonal, the parameters being uncorrelated before any observations are taken.

Numerically, the addition of a positive definite diagonal matrix to the original least squares matrix according to (3.20) has a stabilizing effect with respect to the inversion, as is well-known from the Levenberg-Marquardt-Morrison method for non-linear least squares problems /1,16/. If a-priori information is to be neglected, its formal influence in (3.22) is reduced by choosing comparatively large variances for the parameters μ ; yet, the advantage of inverting more stable matrices is preserved.

Bayes estimation techniques find their main field of applications, where the prior information is of great significance, as in the social sciences /8/. For our purpose, the essential advantages, besides those already mentioned, will become apparent when dealing with sequential estimation as derived in the following chapter.

4. Sequential Bayes Estimation

In principle the bulk estimation procedures discussed so far impute that the experimentalist behaves in some unnatural way since they expect that he will not try to analyze his data before having completed his experiment. Of course, a natural behaviour would involve the permanent judgement of the actual state of his experimental results with respect to the objectives of the measurement by accumulating and updating the information with every new observation. This purpose is achieved by the sequential estimation techniques /4,5,15,17/.

In this chapter we shall give a short derivation of the recursive equations governing the Bayes sequential estimation scheme, which turns out to be a specialization of the general linear Kalman-Bucy filter equations. These filter equations are rather involved and would require a detailed discussion of filter and control theory, which would exceed the range of applications covered in this report (the reader is referred to /4,15/ for an extensive analysis, and for an introduction into this field to /18/ or /19/). The derivation of the sequential algorithm presented in this report will be restricted to the problem of estimating time-independent parameters within the framework of linear Bayes estimation as put forward in chapter 3. The generalization to Kalman filter theory will be discussed only qualitatively.

In the following we assume a model linear in the unknown parameters θ_v , $v=1, \dots, n$:

$$(4.1) \quad f(x; \theta) = h^{+T}(x) \theta ,$$

where the components of the vector $h^{+}(x)$ constitute a set of n functions, which yields from (1.1):

$$(4.2) \quad \phi(x; \theta) = h^T(x) \theta ,$$

with

$$(4.3) \quad h(x) = \int g(x, x') h^+(x') dx'.$$

Under the assumption that the observation vector z with statistically independent components is normally distributed:

$p(z|\theta) = L(z|\theta)$ with $L(z|\theta)$ according to (2.7), and that the a-priori information with respect to the parameter vector θ can be described by a normal distribution: $p(\theta)$ according to (3.17), Bayes estimation maximizing the a-posteriori probability distribution for θ leads to the minimization with respect to θ of the quadratic form (see also (3.18)):

$$(4.4) \quad Q_m(\theta) = (\theta - \mu)^T V_\mu^{-1} (\theta - \mu) + (z - H\theta)^T V_z^{-1} (z - H\theta),$$

with μ , V_μ the a-priori mean and covariance matrix, respectively, $V_z = \text{Diag}(\sigma_i^2)$, and

$$(4.5) \quad H_{iv} = h_v(x_i), \quad v=1, \dots, n, \quad i=1, \dots, m.$$

The subscript m of $Q_m(\theta)$ indicates that m measurements are involved. The minimization of $Q_m(\theta)$ yields the parameter vector $\hat{\theta}$ which is optimal in the Bayes sense.

The sequential estimation proceeds as follows. We assume that k values have been measured, and the optimal parameters $\hat{\theta}(k)$ with corresponding covariance matrix $V_\theta(k)$ have been obtained from $Q_k(\theta)$ according to (4.4). One further measurement is performed with the result z_{k+1} . The aim is to derive expressions for $\hat{\theta}(k+1), V_\theta(k+1)$ directly from the optimal values at stage k , and not to reminimize the total $Q_{k+1}(\theta)$.

We write for brevity σ for σ_{k+1} and h for the vector $h(x_{k+1})$. With the convention that H stands for the $k \times n$ matrix H_{iv} from

(4.5) and z and V_z correspond to the data vector and covariance matrix at step k , we obtain from (4.4) with $m = k+1$:

$$(4.6) \quad Q_{k+1}(\theta) = (\theta - \mu)^T V_\mu^{-1} (\theta - \mu) + (z - H\theta)^T V_z^{-1} (z - H\theta) + \frac{1}{\sigma^2} (z_{k+1} - h^T \theta)^2$$

$$= Q_k(\theta) + \frac{1}{\sigma^2} (z_{k+1} - h^T \theta)^2.$$

With $V_\theta(k+1)$ as the covariance matrix of the parameter vector $\hat{\theta}(k+1)$ after $(k+1)$ observations, expression (4.6) can be written analogous to (3.19):

$$(4.7) \quad Q_{k+1}(\theta) = (\theta - \hat{\theta}(k+1))^T V_\theta^{-1}(k+1) (\theta - \hat{\theta}(k+1)) + Q_{k+1, \min}.$$

If $Q_k(\theta)$ in (4.6) is also expressed as the corresponding expansion around the minimum parameter $\hat{\theta}(k)$:

$$(4.8) \quad Q_{k+1}(\theta) = (\theta - \hat{\theta}(k))^T V_\theta^{-1}(k) (\theta - \hat{\theta}(k)) + Q_{k, \min} + \frac{1}{\sigma^2} (z_{k+1} - h^T \theta)^2,$$

a collection by powers of θ yields, when equating (4.7) and (4.8):

$$(4.9) \quad \theta^T \{ V_\theta^{-1}(k+1) - (V_\theta^{-1}(k) + \frac{1}{\sigma^2} h h^T) \} \theta$$

$$- 2 \{ V_\theta^{-1}(k+1) \hat{\theta}(k+1) - (V_\theta^{-1}(k) \hat{\theta}(k) + \frac{1}{\sigma^2} z_{k+1} \cdot h) \}^T \theta$$

$$+ \{ Q_{k+1, \min} - Q_{k, \min} + \hat{\theta}^T(k+1) V_\theta^{-1}(k+1) \hat{\theta}(k+1) - (\hat{\theta}^T(k) V_\theta^{-1}(k) \hat{\theta}(k) + \frac{1}{\sigma^2} z_{k+1}) \} = 0.$$

Since this identity has to be valid for all θ , we find the following three equations to hold, assuming that the matrix inversion of $V_\theta^{-1}(k+1)$ can be performed:

$$(4.10a) \quad V_{\theta}(k+1) = \{V_{\theta}^{-1}(k) + \frac{1}{\sigma^2} h h^T\}^{-1}$$

$$(4.10b) \quad \hat{\theta}(k+1) = V_{\theta}(k+1) \{V_{\theta}^{-1}(k) \hat{\theta}(k) + \frac{1}{\sigma^2} z_{k+1} h\}$$

$$(4.10c) \quad Q_{k+1, \min} = Q_{k, \min} + \hat{\theta}^T(k) V_{\theta}^{-1}(k) \hat{\theta}(k) + \frac{1}{\sigma^2} z_{k+1}^2$$

$$\hat{\theta}^T(k+1) V_{\theta}^{-1}(k+1) \hat{\theta}(k+1).$$

This system of equations for $k=0,1,\dots,m-1$ together with the initial values

$$(4.11) \quad V_{\theta}(0) = V_{\mu}, \quad \hat{\theta}(0) = \mu, \quad Q_{0, \min} = 0$$

establishes the sequential algorithm. The order in which the solution of (4.10) has to be obtained, is strictly determined: First solve the matrix equation (4.10a) for the covariance matrix, then the vector equation (4.10b) for the parameter vector, and finally the scalar equation (4.10c) for the minimum value of $Q_{k+1}(\theta)$. This has certain consequences: If we are not interested in the minimum value of $Q_{k+1}(\theta)$, we can solve system (4.10) ignoring equation (4.10c); if we were only interested in the covariance matrix the solution of system (4.10) reduces to the solution of equation (4.10a) which is not coupled to (4.10b) and (4.10c), thus the covariance matrix is independent of the actual value of the measurement z_{k+1} and only depends on the variance σ_{k+1}^2 and the experimental variable x_{k+1} . For this reason the linear sequential algorithm proves to be a powerful tool for sequential experimental design [20,21]. In general we want to obtain the estimate of the parameter vector θ ; thus, the corresponding recursive scheme to be solved will consist of equations (4.10a) and (4.10b).

Usually, a more elegant version of equations (4.10) is found in the literature which leads to a better insight into this updating process. To this end we first give a short outline of the

so-called matrix inversion lemma following the derivation given in /4/ or /17/ (a more fundamental proof can be obtained when applying the eigenvalue decomposition to the matrix $(I+aa^T)^{-1}$ for any n-dimensional vector a). We start with:

$$(4.12) \quad V_{\theta}^{-1}(k+1) = V_{\theta}^{-1}(k) + \frac{1}{\sigma^2} hh^T,$$

and multiply this by $V_{\theta}(k+1)$ from left and $V_{\theta}(k)$ from right:

$$(4.13) \quad V_{\theta}(k) = V_{\theta}(k+1) + \frac{1}{\sigma^2} V_{\theta}(k+1) hh^T V_{\theta}(k).$$

Next, another expression for the matrix $V_{\theta}(k+1) hh^T V_{\theta}(k)$ is derived. To do this we multiply the vector h by equation (4.13):

$$(4.14) \quad \begin{aligned} V_{\theta}(k)h &= V_{\theta}(k+1) h + \frac{1}{\sigma^2} V_{\theta}(k+1) hh^T V_{\theta}(k) h \\ &= V_{\theta}(k+1) h \left\{ 1 + \frac{1}{\sigma^2} h^T V_{\theta}(k) h \right\}. \end{aligned}$$

Since $V_{\theta}(k)$ is a covariance matrix, which is always positive definite, the scalar quantity $(\sigma^2 + h^T V_{\theta}(k) h)$ will also be positive. To obtain the alternative expression for $V_{\theta}(k+1) hh^T V_{\theta}(k)$, the dyadic product of vector equation (4.14) with the row vector $h^T V_{\theta}(k)$ yields:

$$(4.15) \quad \frac{1}{\sigma^2} V_{\theta}(k+1) hh^T V_{\theta}(k) = (\sigma^2 + h^T V_{\theta}(k) h)^{-1} V_{\theta}(k) hh^T V_{\theta}(k).$$

This inserted into (4.13) leads to the matrix inversion lemma:

$$(4.16a) \quad V_{\theta}(k+1) = V_{\theta}(k) - (\sigma^2 + h^T V_{\theta}(k) h)^{-1} V_{\theta}(k) hh^T V_{\theta}(k),$$

and this inserted into (4.10b):

$$\begin{aligned}
 \hat{\theta}(k+1) &= \hat{\theta}(k) - (\sigma^2 + h^T V_{\theta}(k) h)^{-1} V_{\theta}(k) h h^T \hat{\theta}(k) + \frac{1}{\sigma^2} z_{k+1} V_{\theta}(k) h \\
 &\quad - (\sigma^2 + h^T V_{\theta}(k) h)^{-1} \frac{1}{\sigma^2} z_{k+1} V_{\theta}(k) h h^T V_{\theta}(k) h \\
 &= \hat{\theta}(k) + (\sigma^2 + h^T V_{\theta}(k) h)^{-1} V_{\theta}(k) h \left\{ \frac{1}{\sigma^2} z_{k+1} (\sigma^2 + h^T V_{\theta}(k) h) \right. \\
 &\quad \left. - \frac{1}{\sigma^2} z_{k+1} h^T V_{\theta}(k) h - h^T \hat{\theta}(k) \right\},
 \end{aligned}$$

or:

$$(4.16b) \quad \hat{\theta}(k+1) = \hat{\theta}(k) + (\sigma^2 + h^T V_{\theta}(k) h)^{-1} V_{\theta}(k) h (z_{k+1} - h^T \hat{\theta}(k)).$$

An alternative expression for equation (4.10c) can best be found when considering equation (4.8) for $\theta = \hat{\theta}(k+1)$:

$$\begin{aligned}
 Q_{k+1, \min} &= (\hat{\theta}(k+1) - \hat{\theta}(k))^T V_{\theta}^{-1}(k) (\hat{\theta}(k+1) - \hat{\theta}(k)) + \frac{1}{\sigma^2} (z_{k+1} - h^T \hat{\theta}(k+1))^2 \\
 &\quad + Q_{k, \min}.
 \end{aligned}$$

Taking $(\hat{\theta}(k+1) - \hat{\theta}(k))$ from (4.16b) one finds:

$$\begin{aligned}
 Q_{k+1, \min} &= Q_{k, \min} + (\sigma^2 + h^T V_{\theta}(k) h)^{-2} (z_{k+1} - h^T \hat{\theta}(k))^2 h^T V_{\theta}(k) h \\
 &\quad + \frac{1}{\sigma^2} \{ (z_{k+1} - h^T \hat{\theta}(k)) - (\sigma^2 + h^T V_{\theta}(k) h)^{-1} (z_{k+1} - h^T \hat{\theta}(k)) h^T V_{\theta}(k) h \}^2 \\
 &= Q_{k, \min} + \frac{1}{\sigma^2} (\sigma^2 + h^T V_{\theta}(k) h)^{-2} (z_{k+1} - h^T \hat{\theta}(k))^2 \{ \sigma^2 h^T V_{\theta}(k) h \\
 &\quad + (\sigma^2 + h^T V_{\theta}(k) h - h^T V_{\theta}(k) h)^2 \}.
 \end{aligned}$$

Finally:

$$(4.16c) \quad Q_{k+1, \min} = Q_{k, \min} + (\sigma^2 + h^T V_{\theta}(k) h)^{-2} (z_{k+1} - h^T \hat{\theta}(k))^2.$$

Formulae (4.16) constitute the famous Wiener-Kalman filter equations for the special case of time-independent parameters /5,6,17/. At this point, a few comments seem to be worthwhile.

Comparing (4.10a) with (4.16a) it turns out that making use of the matrix inversion lemma, equation (4.16a) no longer requires an explicit matrix inversion but only a division by a scalar. This usually will lead to numerically more stable solutions, as discussed in /6/. From (4.12) it can be seen that the matrix $V_{\theta}(k)$ will decrease monotonically - if this term may be used for matrices. More correctly, the quadratic form $y^T V_{\theta}^{-1}(k+1)y$ shows a decreasing behaviour with increasing k , since $y^T h h^T y \geq 0$ for any y , and $V_{\theta}^{-1}(k)$ is positive definite. Therefore, the estimator (4.16b) is statistically consistent /9,10/.

A problem may arise if measurements with high accuracy are processed (if the measurements are "too good"); then the term $V_{\theta}^{-1}(k)$ may be negligible compared to $h h^T / \sigma^2$ in (4.12). The matrix $h h^T$, however, has rank one ($n-1$ linearly independent vectors y can be found with $(h^T y) = 0$ and therefore $(y^T h h^T y) = 0$ in an $(n-1)$ -dimensional subspace), so that the inverse of $V_{\theta}^{-1}(k+1)$ does not exist. If on the other hand equation (4.16a) still is used, the covariance matrix $V_{\theta}(k+1)$ will no longer be positive definite. This can also be seen when forming the vector $V_{\theta}(k+1)h$, which is equal to zero for $\sigma^2 = 0$. This divergence property of the filter equations is well-known /17,19/.

Equation (4.16b) often is written as:

$$\hat{\theta}(k+1) = \hat{\theta}(k) + K(k) (z_{k+1} - h^T \hat{\theta}(k)),$$

where

$$K(k) = (\sigma^2 + h^T V_{\theta}(k) h)^{-1} V_{\theta}(k) h$$

is referred to as the Kalman gain matrix /17/, which in our case, when only one additional measurement is processed, degenerates to an $n \times 1$ matrix. In general, the sequential procedure will work with r data points per step k , then $K(k)$ is an $n \times r$ matrix. The new estimate $\hat{\theta}(k+1)$ is equal to the old estimate $\hat{\theta}(k)$ plus a correction term controlled by the gain matrix. According to our model (2.6), the expression $h^T \hat{\theta}(k)$ is, on the basis of the information of k observation steps, the prediction of the experimental data to be expected at the $(k+1)$ -st measurement step. The correction to this prediction will be significant for large deviations between predicted and actually measured values, thus representing the relative gain of information due to the new measurement. Hence, the name "gain matrix" for $K(k)$ is justified.

Equation (4.16c) may be interpreted along the same lines as (4.16b). Moreover, the minimum value $Q_{k,\min}$ is seen to be a strictly monotonously increasing function with the number of measurements k .

With the introduction of prior information consisting of μ and V_μ , compared to the pure least-squares or maximum-likelihood case, additional influences affect the estimation of parameters. In the limit $V_\mu^{-1} = 0$ the Bayes estimator is equivalent to the least-squares and maximum-likelihood estimator as derived in (2.9). Therefore, if only incomplete prior information is available the sequential algorithm should start with large diagonal elements of V_μ , the off-diagonal elements of V_μ being set equal to zero anyway. It is nevertheless possible to treat the pure least-squares case sequentially, if equations (4.10) or (4.16) are initiated for $k=n$, since this is the smallest number of measurement points to obtain a unique solution for $\hat{\theta}$ with a covariance matrix $V_\theta(n)$, which in this case - contrary to the case with prior information available from the beginning of an experiment - requires the inversion of an $n \times n$ matrix according to (4.10a). Once this matrix is obtained, the matrix inversion

lemma may be applied as before and for observations $k > n$ no more inversions are needed.

From the algorithmic point of view, the origin of the prior information is irrelevant for the sequential estimation procedure; in principle there is no difference between prior information gained from experimental sampling and prior information established by some kind of subjective probability. Of course, the sequential estimator will be biased for any finite number of observations, if the a-priori expectation of θ does not coincide with the true value of θ . However, with increasing k , the disturbing influence of wrong prior information is asymptotically eliminated.

Let us conclude this chapter with a few comments concerning the original Kalman filter equations. These equations were derived for time-varying systems where k means the sampling point at time t_k , and $t_{k+1} > t_k$. In addition to the model equation (2.6) or (4.2), a transition equation for the parameters $\theta(k)$ is set up, where parameters are allowed to vary with time according to state transitions, by control, or due to disturbances. Therefore, one has to distinguish between parameter estimates $\hat{\theta}(k+1|k)$ which are optimal estimates obtained immediately before the measurement at t_{k+1} , and parameter estimates $\hat{\theta}(k+1|k+1)$ which are optimal estimates obtained immediately after the measurement at t_{k+1} has been processed.

5. Nilson's Representation of Cubic Cardinal Splines on Uniform Meshes

The Bayes estimation theory as presented here requires a linear model function, that is a function linear with respect to the parameters to be estimated. As model functions we use cubic spline functions following the representation given by Nilson /7/. These splines are constructed from cardinal splines /22/ with $N+1$ equally spaced knots.

A rederivation of Nilson's results seems to be justified because, on the one hand, we feel that a better understanding of these splines may be obtained, on the other hand a different representation appropriate for the Bayes algorithm will be necessary.

A basic cubic cardinal spline $A(\theta)$ is given by pieces of polynomials with degree three or smaller, which are connected at the knots with a continuous second derivative, and which are zero at the knots except at one knot, where the cardinal spline is equal to unity. In the following we assume equidistant knots which for simplicity are placed at $0, \pm 1, \pm 2, \dots$, and at the zero knot the cardinal spline has the value 1; in addition we postulate symmetry with respect to inversion of θ :

$$(5.1) \quad A(v) = \delta_{v,0}, \quad v = 0, \pm 1, \pm 2, \dots$$

$$A(-\theta) = A(\theta).$$

These conditions define the construction scheme for the basic cardinal spline. We introduce a polynomial $R(\theta)$ to describe $A(\theta)$ on $[0,1]$:

$$R(\theta) = r_0 + r_1\theta + r_2\theta^2 + r_3\theta^3.$$

We have $r_0=1$ from $A(0)=1$, and $r_1=0$ from $A(\theta)=A(-\theta)$ and the continuity of the first derivative. With $A(1)=0$ we obtain:

$$R(\theta) = 1 - (1+r_3)\theta^2 + r_3\theta^3,$$

or, in Nilson notation with $r_3=a+2$:

$$(5.2) \quad R(\theta) = 1 - (a+3)\theta^2 + (a+2)\theta^3.$$

The constant a will be defined later.

Next we consider a polynomial representing $A(\theta)$ on $[v, v+1]$, $v=1, 2, \dots$:

$$A(\theta) = a_0 + a_1\tau + a_2\tau^2 + a_3\tau^3, \quad \tau = \theta - v, \quad 0 \leq \tau \leq 1.$$

Since $A(v) = A(v+1) = 0$, we must have: $a_0=0$ and $(a_1+a_2+a_3)=0$. With $a_1 = A'(v)$, $a_2 = \frac{1}{2} A''(v)$ we find:

$$(5.3) \quad A(\theta) = A'(v) (\tau - \tau^3) + \frac{1}{2} A''(v) (\tau^2 - \tau^3).$$

The continuity of the first and second derivatives therefore requires:

$$(5.4) \quad \begin{aligned} A'(v+1) &= -2 A'(v) - \frac{1}{2} A''(v) \\ A''(v+1) &= -6 A'(v) - 2 A''(v) \end{aligned} \quad v=1, 2, \dots$$

The matrix:

$$L = \begin{pmatrix} -2 & -\frac{1}{2} \\ -6 & -2 \end{pmatrix}$$

with eigenvalues $\lambda_1 = -2 + \sqrt{3}$, $\lambda_2 = -2 - \sqrt{3}$ governs the recursive

behaviour of relation (5.4):

$$\begin{pmatrix} A'(v+1) \\ A''(v+1) \end{pmatrix} = L^v \begin{pmatrix} A'(1) \\ A''(1) \end{pmatrix} = \frac{\sqrt{3}}{4} \begin{pmatrix} \frac{2}{\sqrt{3}}(\lambda_1^v + \lambda_2^v) & -\frac{1}{3}(\lambda_1^v - \lambda_2^v) \\ -4(\lambda_1^v - \lambda_2^v) & \frac{2}{\sqrt{3}}(\lambda_1^v + \lambda_2^v) \end{pmatrix} \cdot \begin{pmatrix} A'(1) \\ A''(1) \end{pmatrix},$$

or, if we insert from (5.2) $A'(1)=a$, $A''(1)=4a+6$:

$$(5.5) \quad \begin{aligned} A'(v) &= \frac{1}{2} \left[\left(a - \frac{1}{\sqrt{3}}(2a+3)\right) \lambda_1^{v-1} + \left(a + \frac{1}{\sqrt{3}}(2a+3)\right) \lambda_2^{v-1} \right], \\ A''(v) &= \sqrt{3} \left[-\left(a - \frac{1}{\sqrt{3}}(2a+3)\right) \lambda_1^{v-1} + \left(a + \frac{1}{\sqrt{3}}(2a+3)\right) \lambda_2^{v-1} \right]. \end{aligned}$$

According to Nilson, "the choice of the constant a is governed entirely by considerations of convenience". From the explicit representation (5.5) we draw the conclusion that because of $|\lambda_2| > 1$ an arbitrary choice of the constant a would almost always lead to a numerically extremely unstable spline representation for not too small N -values. Therefore, we must postulate that the coefficient proportional to λ_2^{v-1} will either disappear or, at most, be proportional to c^N with $0 < c \leq |\lambda_2|^{-1}$. This latter case will play an important role for periodic cardinal splines to be discussed later.

5.1 Nonperiodic cardinal splines

For nonperiodic cardinal splines the constant a will be defined to eliminate the term proportional to λ_2^{v-1} in (5.5):

$$a(\sqrt{3}+2) + 3 = 0,$$

or:

$$(5.6) \quad a = \frac{3}{\lambda_2} = 3\lambda_1.$$

Then we obtain from (5.5), (5.3), and (5.2), where from now on $\lambda_1 = \lambda$:

$$(5.7a) \quad A'(v) = \frac{\sqrt{3}}{2} (\lambda^2 - 1) \lambda^{v-1} = 3\lambda^v$$

$$(5.7b) \quad A''(v) = -3 (\lambda^2 - 1) \lambda^{v-1} = -6\sqrt{3}\lambda^v, \quad v=1,2,\dots$$

$$(5.8a) \quad A(\theta) = 3\lambda^v T(\theta-v), \quad v \leq \theta \leq v+1, \quad v=1,2,\dots$$

$$(5.8b) \quad A(\theta) = R(\theta), \quad 0 \leq \theta \leq 1,$$

with

$$(5.9a) \quad T(\theta) = \theta - (\lambda+2)\theta^2 + (\lambda+1)\theta^3,$$

$$(5.9b) \quad R(\theta) = 1 - 3(\lambda+1)\theta^2 + (3\lambda+2)\theta^3 = 3T(\theta) + (1-\theta)^3.$$

Equations (5.8) together with $A(-\theta) = A(\theta)$ define the basic cardinal spline; this is essentially one and the same cubic arc $T(\theta)$, which for each successive interval is reduced by the factor $\lambda = -2 + \sqrt{3}$.

The function $f(x)$ to be interpolated by the spline may be given on the uniform mesh:

$$x_0 < x_1 < \dots < x_N, \quad x_j - x_{j-1} = h,$$

with function values $f(x_j) = f_j$ at the knots. Then the interpolating spline can be expressed by the $N+1$ cardinal splines [7]:

$$(5.10) \quad A_j(x) = A\left(\frac{x-x_j}{h}\right), \quad A_j(x_k) = \delta_{j,k}, \quad j, k = 0, 1, \dots, N,$$

$$(5.11) \quad S(x) = \sum_{j=0}^N f_j A_j(x),$$

since the sum of piecewise continuous polynomials with coinciding knots is again a piecewise continuous polynomial with the same knots.

Although expression (5.11) interpolates correctly the function $f(x)$, no boundary conditions can be prescribed, especially the natural cubic spline [23] requiring $S''(x_0)=S''(x_N)=0$ cannot be constructed. In order to satisfy boundary conditions, we notice that any cardinal spline having its central peak at $x_k=x_0+kh$, with $k<0$ or $k>N$, does not violate the interpolating property when added to (5.11). Following Nilson we add two cardinal splines with central peaks at (x_0-h) and (x_N+h) to (5.11), thus gaining two free constants to satisfy two boundary conditions:

$$(5.12) \quad S(x) = \sum_{j=0}^N f_j A_j(x) + b_0 B_0(x) + b_N B_N(x),$$

$$B_0(x) = \frac{h}{3\lambda} A\left(\frac{x-x_0}{h} + 1\right), B_N(x) = -\frac{h}{3\lambda} A\left(\frac{x_N-x}{h} + 1\right).$$

The boundary conditions to be considered here are given by:

$$(5.13) \quad \begin{aligned} S''(x_0) - \mu_0 S''(x_1) &= d_0 \\ -\mu_N S''(x_{N-1}) + S''(x_N) &= d_N, \end{aligned}$$

where μ_0, μ_N, d_0, d_N are arbitrary. For $\mu_0=\mu_N=d_0=d_N=0$ we obtain the important case of natural splines. The two equations (5.13) determining the constants b_0 and b_N introduced in (5.12) can be evaluated using

$$A''(v) = -6\sqrt{3}\lambda|v| + 6\delta_{v,0}, \quad v=0, \pm 1, \pm 2, \dots$$

which follows from (5.7b), (5.9b) and $A(\theta)=A(-\theta)$. We obtain for $k=0,1,\dots,N$:

$$(5.14) \quad S''(x_k) = -6\sqrt{3} \left[\frac{1}{h^2} \sum_{j=0}^N f_j \lambda^{|k-j|} + \frac{1}{3\lambda h} (b_0 \lambda^{k+1} - b_N \lambda^{N-k+1}) \right] + \frac{6}{h^2} f_k.$$

Inserting (5.14) into (5.13) we obtain a system of linear equations for b_0, b_N :

$$(5.15) \quad \begin{aligned} -b_0(1-\mu_0\lambda) + b_N \lambda^N (1-\frac{\mu_0}{\lambda}) &= \frac{hd_0}{2\sqrt{3}} + \frac{3}{h} \sum_{j=0}^N f_j (\lambda^j - \mu_0 \lambda^{|j-1|}) - \\ &\quad \frac{\sqrt{3}}{h} (f_0 - \mu_0 f_1) \\ -b_0 \lambda^N (1-\frac{\mu_N}{\lambda}) + b_N (1-\mu_N \lambda) &= \frac{hd_N}{2\sqrt{3}} + \frac{3}{h} \sum_{j=0}^N f_j (\lambda^{N-j} - \mu_N \lambda^{|N-j-1|}) - \\ &\quad \frac{\sqrt{3}}{h} (f_N - \mu_N f_{N-1}). \end{aligned}$$

The solution of this system is given by:

$$(5.16a) \quad b_0 = \frac{1}{D} \left\{ \frac{h}{2\sqrt{3}} \left[d_0 (1-\mu_N \lambda) - d_N \lambda^N (1-\frac{\mu_0}{\lambda}) \right] + \frac{3}{h} \sum_{j=0}^N f_j \left[(1-\mu_N \lambda) (\lambda^j - \mu_0 \lambda^{|j-1|}) - \lambda^N (1-\frac{\mu_0}{\lambda}) (\lambda^{N-j} - \mu_N \lambda^{|N-j-1|}) \right] + \frac{\sqrt{3}}{h} \left[-(1-\mu_N \lambda) (f_0 - \mu_0 f_1) + \lambda^N (1-\frac{\mu_0}{\lambda}) (f_N - \mu_N f_{N-1}) \right] \right\},$$

$$(5.16b) \quad b_N = \frac{1}{D} \left\{ \frac{h}{2\sqrt{3}} \left[d_0 \lambda^N (1-\frac{\mu_N}{\lambda}) - d_N (1-\mu_0 \lambda) \right] + \frac{3}{h} \sum_{j=0}^N f_j \left[\lambda^N (1-\frac{\mu_N}{\lambda}) (\lambda^j - \mu_0 \lambda^{|j-1|}) - (1-\mu_0 \lambda) (\lambda^{N-j} - \mu_N \lambda^{|N-j-1|}) \right] + \frac{\sqrt{3}}{h} \left[-\lambda^N (1-\frac{\mu_N}{\lambda}) (f_0 - \mu_0 f_1) + (1-\mu_0 \lambda) (f_N - \mu_N f_{N-1}) \right] \right\},$$

$$\text{and } D = -(1-\mu_0 \lambda) (1-\mu_N \lambda) + \lambda^{2N} (1-\frac{\mu_0}{\lambda}) (1-\frac{\mu_N}{\lambda}).$$

With the explicit dependence of b_0 and b_N on f_j we now have a representation of (5.12) appropriate for the Bayes algorithm, in which the function values f_j , $j=0, \dots, N$ are treated as parameters θ_v to be estimated, where $v=j+1$ and $n=N+1$ from chapter 4.

To obtain the Bayes representation of (5.12) for each interval $[x_k, x_{k+1}]$, $k=0, 1, \dots, N-1$, we need to specify the cardinal splines $A_j(x)$ according to (5.8) and (5.9). It may be verified that the $A_j(x)$ have the following representation on

$$\begin{aligned}
 & x \in [x_k, x_{k+1}], \quad k=0, 1, \dots, N-1, \quad \theta = \frac{x-x_k}{h} : \\
 & A_j(x) = 3\lambda^{k-j} \cdot T(\theta) + (1-\theta)^3 \cdot \delta_{j,k}, \quad j \leq k \\
 & A_j(x) = 3\lambda^{j-k-1} \cdot T(1-\theta) + \theta^3 \cdot \delta_{j,k+1}, \quad j \geq k+1.
 \end{aligned}
 \tag{5.17}$$

Insertion into (5.12) gives the expression:

$$\begin{aligned}
 (5.18) \quad S(x) = & 3T(\theta) \sum_{j=0}^k f_j \lambda^{k-j} + 3T(1-\theta) \sum_{j=k+1}^N f_j \lambda^{j-k-1} \\
 & + (1-\theta)^3 f_k + \theta^3 f_{k+1} + b_0 h \lambda^k T(\theta) - b_N h \lambda^{N-k-1} \cdot T(1-\theta).
 \end{aligned}$$

With the auxiliary function e_j^k defined by

$$e_j^k = 1 \text{ for } j \leq k, \quad e_j^k = 0 \text{ for } j > k,$$

a collection of all terms proportional to f_j can be easily performed resulting in the final Bayes representation of $S(x)$ for $x \in [x_k, x_{k+1}]$, $k=0, 1, \dots, N-1$:

$$\begin{aligned}
 S(x) = & \frac{h^2}{2D\sqrt{3}} \left\{ T(\theta)\lambda^k \left[d_0(1-\mu_N\lambda) - d_N(1-\frac{\mu_0}{\lambda}) \right] \right. \\
 & \left. - T(1-\theta)\lambda^{N-k-1} \left[d_0\lambda^N(1-\frac{\mu_N}{\lambda}) - d_N(1-\mu_0\lambda) \right] \right\} \\
 & + \sum_{j=0}^N f_j \left\{ 3T(\theta)\lambda^{k-j}e_j^k + 3T(1-\theta)\lambda^{j-k-1}(1-e_j^k) + (1-\theta)^3\delta_{j,k} + \theta^3\delta_{j,k+1} \right. \\
 & + \frac{3}{D} \left\langle T(\theta)\lambda^k \cdot (1-\frac{\mu_0}{\lambda}) \left[\lambda^j(1-\mu_N\lambda) - \lambda^{2N-j}(1-\frac{\mu_N}{\lambda}) \right] \right. \\
 (5.19) \quad & \left. - T(1-\theta)\lambda^{N-k-1}(1-\frac{\mu_N}{\lambda}) \left[-\lambda^{N-j}(1-\mu_0\lambda) + \lambda^{N+j}(1-\frac{\mu_0}{\lambda}) \right] \right\rangle \\
 & + \frac{\sqrt{3}}{D} \left\langle T(\theta)\lambda^k \left[-(1-\mu_N\lambda)(\delta_{j,0}-\mu_0\delta_{j,1}) \right. \right. \\
 & \left. \left. + \lambda^N(1-\frac{\mu_0}{\lambda})(\delta_{j,N}-\mu_N\delta_{j,N-1}) \right] \right. \\
 & \left. - T(1-\theta)\lambda^{N-k-1} \left[(1-\mu_0\lambda)(\delta_{j,N} \cdot (1+6\mu_N) - \mu_N\delta_{j,N-1}) \right. \right. \\
 & \left. \left. - \lambda^N(1-\frac{\mu_N}{\lambda})(\delta_{j,0}(1+6\mu_0) - \mu_0\delta_{j,1}) \right] \right\rangle \Bigg\}
 \end{aligned}$$

$$\text{with: } \theta = \frac{x-x_k}{h}, \quad D = -(1-\mu_0\lambda)(1-\mu_N\lambda) + \lambda^{2N}(1-\frac{\mu_0}{\lambda})(1-\frac{\mu_N}{\lambda})$$

$$\text{and: } T(\theta) = \theta - (\lambda+2)\theta^2 + (\lambda+1)\theta^3.$$

The spline model is a linear model with respect to the unknown parameters f_j plus, in general, a constant term which vanishes for $d_0=d_N=0$.

In case that the parameters f_j are known, i.e. only the interpolating property of the spline is of interest, a very efficient algorithm due to Nilson /7/ can be obtained from (5.18).

Defining

$$\begin{aligned}
 \sigma_k &= 3 \sum_{j=0}^k f_j \lambda^{k-j}, \\
 \tau_k &= 3 \sum_{j=k}^N f_j \lambda^{j-k},
 \end{aligned}
 \quad (5.20)$$

a recursive evaluation of σ_k , τ_k is possible:

$$\sigma_0 = 3 f_0 ,$$

$$\sigma_{k+1} = \lambda \sigma_k + 3 f_{k+1} , \quad k=0,1,\dots,N-1,$$

$$\tau_N = 3 f_N ,$$

$$\tau_{k-1} = \lambda \tau_k + 3 f_{k-1} , \quad k=N,N-1,\dots,1.$$

Equation (5.18) then reads:

$$\begin{aligned} S(x) = & \sigma_k \cdot T(\theta) + \tau_{k+1} \cdot T(1-\theta) + (1-\theta)^3 \cdot f_k + \theta^3 \cdot f_{k+1} \\ (5.21) \quad & + b_0 h \cdot \lambda^k \cdot T(\theta) - b_N h \cdot \lambda^{N-k-1} \cdot T(1-\theta), \end{aligned}$$

with

$$\begin{aligned} b_0 = \frac{1}{D} \left\{ \frac{h}{2\sqrt{3}} \left[d_0 (1-\mu_N \lambda) - d_N \lambda^N \left(1 - \frac{\mu_0}{\lambda}\right) \right] \right. \\ + \frac{1}{h} \left(1 - \frac{\mu_0}{\lambda}\right) \left[(1-\mu_N \lambda) \tau_0 - \lambda^N \left(1 - \frac{\mu_N}{\lambda}\right) \sigma_N \right] \\ + \frac{\sqrt{3}}{h} \left[-(1-\mu_N \lambda) ((1+6\mu_0) f_0 - \mu_0 f_1) \right. \\ \left. \left. + \lambda^N \left(1 - \frac{\mu_0}{\lambda}\right) ((1+6\mu_N) f_N - \mu_N f_{N-1}) \right] \right\} \end{aligned}$$

$$b_N = \frac{1}{D} \left\{ \frac{h}{2\sqrt{3}} \left[\lambda^N \left(1 - \frac{\mu_N}{\lambda}\right) d_0 - (1 - \mu_0 \lambda) d_N \right] \right. \\ \left. - \frac{1}{h} \left(1 - \frac{\mu_N}{\lambda}\right) \left[(1 - \mu_0 \lambda) \sigma_N - \left(1 - \frac{\mu_0}{\lambda}\right) \tau_0 \right] \right. \\ \left. + \frac{\sqrt{3}}{h} \left[(1 - \mu_0 \lambda) ((1 + 6\mu_N) f_N - \mu_N f_{N-1}) - \lambda^N \left(1 - \frac{\mu_N}{\lambda}\right) ((1 + 6\mu_0) f_0 - \mu_0 f_1) \right] \right\}.$$

The calculation procedure then is as follows. Given the set of function values f_j , $j=0,1,\dots,N$, calculate $\lambda^k, \sigma_k, \tau_k$, $k=0,\dots,N$, and b_0, b_N . This is the initialization step. To compute the actual spline value at x , determine first the interval k with $x_{k-1} \leq x < x_k$. With $\theta = (x - x_{k-1})/h$ the four function values $T(\theta)$, $T(1-\theta)$, θ^3 and $(1-\theta)^3$ have to be evaluated, multiplied by the corresponding factors and added up.

If we compare this procedure with the amount of computational effort necessary for the B-spline representation /3,23/:

$$S(x) = \sum_{j=0}^N c_j B_j(x),$$

we find, that once the set of coefficients c_j is determined initially, and the interval of the actual point x is known, again four B-splines have to be evaluated, multiplied by c_j , and summed up.

The significant difference of the B-spline method has its origin in the computation of the coefficients c_j which are the solution of a system of linear equations of the form:

$$(5.22) \quad c_{j-1} \cdot B_{j-1}(x_j) + c_j \cdot B_j(x_j) + c_{j+1} \cdot B_{j+1}(x_j) = f_j.$$

The parameters c_j of the B-spline model therefore cannot immediately be interpreted as a prior information concerning the function $f(x; \theta)$, unless the tridiagonal matrix of (5.22) has been inverted. (This is especially of disadvantage when estimation problems are considered, where special conditions, e.g. symmetry

properties given by physical arguments, have to be implied which most easily can be expressed by some relations holding between the f_k .) On the other hand, the B-splines have a strictly limited range of four intervals (splines with finite support), whereas the cardinal splines, although decreasing exponentially with increasing distance from the central peak, have an infinite range (Fig. 1). This will be more time-consuming, if problems of type (1.1) are considered, where long-ranging kernels $g(x, x')$ are involved.

5.2 Periodic cardinal splines

A periodic cubic spline must by definition fulfil the boundary conditions /22/:

$$(5.23) \quad S^{(r)}(x_0) = S^{(r)}(x_N), \quad r=0,1,2.$$

As in the nonperiodic case, we shall first derive formulae for a periodic basic cardinal spline with knots at $v=0, \pm 1, \pm 2, \dots$ and centered at $v=0$, i.e. which satisfies conditions analogous to (5.1):

$$(5.24a) \quad A(v) = \delta_{v,0}, \quad v=0, \pm 1, \pm 2, \dots, \pm M.$$

$$(5.24b) \quad A(-\theta) = A(\theta).$$

Here M has the following meaning: If the total number of knots is $N+1$, then:

$$\begin{aligned} N=2M, & \quad N \text{ even,} \\ N=2M+1, & \quad N \text{ odd.} \end{aligned}$$

In addition we require:

$$(5.25a) \quad A'(M) = 0, \quad N \text{ even},$$

$$(5.25b) \quad A'(M + \frac{1}{2}) = 0, \quad N \text{ odd}.$$

Because of (5.24b) the second derivative is an even function of θ and therefore:

$$(5.26) \quad \begin{aligned} A^{(r)}(-M) &= A^{(r)}(M), & r=0,1,2, & N \text{ even}, \\ A^{(r)}(-M-\frac{1}{2}) &= A^{(r)}(M+\frac{1}{2}), & r=0,1,2, & N \text{ odd}. \end{aligned}$$

Adding periodic replications of $A(\theta)$ defined on the intervals $[-M, M]$, $[-M-\frac{1}{2}, M+\frac{1}{2}]$, respectively, a periodic basic cardinal spline in the sense of definition (5.23) is obtained, where x_0 may be chosen arbitrarily.

This periodicity property would also have been obtained by satisfying, instead of (5.25), the condition:

$$A(v) = \delta_{v \bmod(N), 0}, \quad v=0, \pm 1, \pm 2, \dots$$

The nonperiodic spline was constructed by cancelling terms proportional to λ_2^{v-1} in (5.5). This selection was not unique because other a -values making this term small might have been chosen as well. Conditions (5.25) for the periodic spline make a unique choice of a necessary. We rewrite (5.5):

$$(5.5') \quad \begin{aligned} A'(v) &= \frac{1}{2\sqrt{3}} \left[a \cdot (\lambda_1^v - \lambda_2^v) - 3 \cdot (\lambda_1^{v-1} - \lambda_2^{v-1}) \right] \\ A''(v) &= \left[-a \cdot (\lambda_1^v + \lambda_2^v) + 3 \cdot (\lambda_1^{v-1} + \lambda_2^{v-1}) \right]. \end{aligned}$$

From $A'(M)=0$ we find with $\lambda_1 \lambda_2 = 1$:

$$(5.27a) \quad a = 3 \cdot \frac{\lambda_1^{M-1} - \lambda_2^{M-1}}{\lambda_1^M - \lambda_2^M} = 3 \cdot \frac{\lambda_1 - \lambda_1^{2M-1}}{1 - \lambda_1^{2M}}, \quad M > 0, \quad N = 2M.$$

We calculate $A'(M+\frac{1}{2})$ from (5.3) and (5.5'):

$$\begin{aligned} A'(M+\frac{1}{2}) &= \frac{1}{4} A'(M) + \frac{1}{8} A''(M) \\ &= \frac{1}{8\sqrt{3}} \left\{ a \left[(1-\sqrt{3})\lambda_1^M - (1+\sqrt{3})\lambda_2^M \right] - 3 \cdot \left[(1-\sqrt{3})\lambda_1^{M-1} - (1+\sqrt{3})\lambda_2^{M-1} \right] \right\}. \end{aligned}$$

Using $\frac{1-\sqrt{3}}{1+\sqrt{3}} = \lambda_1$:

$$A'(M+\frac{1}{2}) = \frac{(1+\sqrt{3})}{8\sqrt{3}} \left[a(\lambda_1^{M+1} - \lambda_2^M) - 3 \cdot (\lambda_1^M - \lambda_2^{M-1}) \right],$$

and $A'(M+\frac{1}{2})=0$ yields:

$$(5.27b) \quad a = 3 \cdot \frac{\lambda_1^M - \lambda_2^{M-1}}{\lambda_1^{M+1} - \lambda_2^M} = 3 \cdot \frac{\lambda_1 - \lambda_1^{2M}}{1 - \lambda_1^{2M+1}}, \quad M \geq 0, \quad N = 2M+1.$$

Therefore both cases (5.27a) and (5.27b) give the same solution:

$$(5.28) \quad a = 3 \cdot \frac{\lambda - \lambda^{N-1}}{1 - \lambda^N},$$

and N even or odd need no longer be distinguished.

$A(\theta)$ can now be calculated from (5.3) and (5.5) for $\theta \geq 1$, and from (5.2) for $0 \leq \theta \leq 1$. With:

$$(5.29a) \quad a - \frac{1}{\sqrt{3}} (2a+3) = \frac{1}{\sqrt{3}} (\lambda a - 3) = \frac{\sqrt{3}}{1 - \lambda^N} (\lambda^2 - 1),$$

$$(5.29b) \quad a + \frac{1}{\sqrt{3}} (2a+3) = \frac{1}{\sqrt{3}} \left(-\frac{1}{\lambda}a+3\right) = \frac{\sqrt{3}}{1-\lambda^N} \cdot \left(\frac{1}{\lambda^2} - 1\right) \cdot \lambda^N,$$

an adequate collection of terms results, because of $\lambda - \frac{1}{\lambda} = 2\sqrt{3}$, in:

$$(5.30) \quad A(\theta) = \frac{\sqrt{3}}{1-\lambda^N} \left\{ \left(\lambda - \frac{1}{\lambda}\right) \lambda^v \cdot \frac{1}{2} \left[(\tau - \tau^3) - \sqrt{3} (\tau^2 - \tau^3) \right] \right. \\ \left. - \left(\lambda - \frac{1}{\lambda}\right) \cdot \lambda^{N-v} \frac{1}{2} \left[(\tau - \tau^3) + \sqrt{3} (\tau^2 - \tau^3) \right] \right\} \\ = \frac{3}{1-\lambda^N} \left\{ \lambda^v \cdot \left[\tau - \sqrt{3} \tau^2 - (1-\sqrt{3}) \tau^3 \right] - \lambda^{N-v} \cdot \left[\tau + \sqrt{3} \tau^2 - (1+\sqrt{3}) \tau^3 \right] \right\},$$

where $\tau = \theta - v$, $v \leq \theta \leq v+1$, $v \geq 1$.

From (5.9a) we have:

$$T(\tau) = \tau - \sqrt{3} \tau^2 - (1-\sqrt{3}) \tau^3,$$

$$T(1-\tau) = (2-\sqrt{3})\tau + (2\sqrt{3}-3)\tau^2 + (1-\sqrt{3})\tau^3 = -\lambda \cdot \left[\tau + \sqrt{3} \tau^2 - (1+\sqrt{3}) \tau^3 \right].$$

Therefore, (5.30) yields:

$$(5.31a) \quad A(\theta) = \frac{3}{1-\lambda^N} \left[\lambda^v T(\tau) + \lambda^{N-v-1} T(1-\tau) \right], \tau = \theta - v, v \leq \theta \leq v+1, v \geq 1,$$

and by direct calculation from (5.2):

$$(5.31b) \quad A(\theta) = (1-\theta)^3 + \frac{3}{1-\lambda^N} \cdot \left[T(\theta) + \lambda^{N-1} \cdot T(1-\theta) \right], \quad 0 \leq \theta \leq 1.$$

Before the computing procedure for periodic splines will be derived, two remarks seem to be useful at this point in order to characterize some interesting features.

Firstly it is seen from (5.29b), that the coefficient proportional to λ_2^{v-1} of the system (5.5) is of $O(\lambda^N)$ with $|\lambda| = |\lambda_2|^{-1} < 1$.

Hence the cardinal spline remains bounded for any N .

Of special interest is furthermore the spline behaviour near the boundary. For $v=M$ the periodic basic cardinal spline is given according to (5.31a):

$$(5.32a) \quad A(\theta) = -\frac{6\sqrt{3}}{1-\lambda^N} \lambda^M \tau^2(1-\tau), \quad N=2M, \quad M \leq \theta \leq M+1,$$

$$(5.32b) \quad A(\theta) = \frac{3(3-\sqrt{3})}{1-\lambda^N} \lambda^M \tau(1-\tau), \quad N=2M+1, \quad M \leq \theta \leq M+1.$$

As would have been expected the spline with even N satisfying (5.25a) has a stationary point at $\theta=M$, whereas the odd spline degenerates to a polynomial of second degree on $[M, M+1]$.

The representation of a function $f(x)$ with $f_j = f(x_j)$, $f_N = f_0$, at the equidistant knots x_j , $j=0, 1, \dots, N$, will be given by periodic cardinal splines as defined in (5.31) and (5.10):

$$(5.33) \quad S(x) = \sum_{j=1}^N A_j(x) f_j$$

with

$$A_j(x_k) = \delta_{j \bmod(N), k}, \quad j, k=0, 1, \dots, N.$$

The term for $j=0$ is omitted in (5.33) because the periodicity condition yields two peaks for the spline $A_0(x)$ at x_0 and x_N .

Expression (5.33) has already the form needed for the Bayes algorithm, since all free parameters f_j do occur explicitly. For computational purposes the specification of $A_j(x)$ in each interval $[x_k, x_{k+1}]$ analogous to (5.17) is advantageous. For $x \in [x_k, x_{k+1}]$, $k=0, 1, \dots, N-1$ and $j=1, 2, \dots, N$ with the exception of the case $k=0$ and $j=N$, we have with $\theta=(x-x_k)/h$:

$$A_j(x) = \frac{3}{1-\lambda^N} \left[\lambda^{k-j} T(\theta) + \lambda^{N-k+j-1} \cdot T(1-\theta) \right] + (1-\theta)^3 \delta_{j,k}, \quad j \leq k, \quad (5.34a)$$

$$A_j(x) = \frac{3}{1-\lambda^N} \left[\lambda^{j-k-1} \cdot T(1-\theta) + \lambda^{N-j+k} \cdot T(\theta) \right] + \theta^3 \delta_{j,k+1}, \quad j > k.$$

On the interval $[x_0, x_N]$, the only cardinal spline of (5.33), which has two peaks equal to unity, is $A_N(x)$ with peaks at x_0 and x_N . Therefore additional to (5.34a) a further requirement is necessary. For $x \in [x_0, x_1]$:

$$(5.34b) \quad A_N(x) = \frac{3}{1-\lambda^N} \left[T(\theta) + \lambda^{N-1} T(1-\theta) \right] + (1-\theta)^3 = A_0(x).$$

Insertion of (5.34) into (5.33) yields for $x \in [x_k, x_{k+1}]$, $k=1, \dots, N-1$:

$$(5.35a) \quad S(x) = \frac{3}{1-\lambda^N} T(\theta) \left[\sum_{j=1}^k f_j \lambda^{k-j} + \sum_{j=k+1}^N f_j \lambda^{N-j+k} \right] + (1-\theta)^3 f_k \\ + \frac{3}{1-\lambda^N} T(1-\theta) \left[\sum_{j=1}^k f_j \lambda^{N-k+j-1} + \sum_{j=k+1}^N f_j \lambda^{j-k-1} \right] + \theta^3 f_{k+1},$$

and on $[x_0, x_1]$:

$$(5.35b) \quad S(x) = \frac{3}{1-\lambda^N} \left[T(\theta) \cdot \sum_{j=1}^N f_j \lambda^{N-j} + T(1-\theta) \sum_{j=1}^N f_j \lambda^{j-1} \right] + (1-\theta)^3 f_N + \theta^3 f_1,$$

which is equivalent to formula (5.35a) for $k=0$ because of $f_N = f_0$.

Next we develop an expression analogous to (5.19). Again the auxiliary function e_j^k , which is equal to one for $j \leq k$ and zero else, will be used. Then both (5.35a) and (5.35b) can be put into the form:

$$\begin{aligned}
 (5.36) \quad S(x) = & \sum_{j=0}^N \left\{ (1-\delta_{j,0}) \left\langle \frac{3}{1-\lambda^N} T(\theta) \left[\lambda^{k-j} e_j^k + \lambda^{N-j+k} \cdot (1-e_j^k) \right] \right. \right. \\
 & + \frac{3}{1-\lambda^N} T(1-\theta) \left[\lambda^{N-k+j-1} e_j^k + \lambda^{j-k-1} \cdot (1-e_j^k) \right] \left. \right\rangle \\
 & + (1-\theta)^3 \delta_{j,k} + \theta^3 \delta_{j,k+1} \left. \right\} f_j,
 \end{aligned}$$

which is valid for $x \in [x_k, x_{k+1}]$, $k=0,1,\dots,N-1$.

Finally a fast algorithm similar to (5.21) is presented /7/ which will be used when the f_j are known.

With

$$\sigma_k = 3 \sum_{j=0}^k f_j \lambda^{k-j}, \quad \tau_k = 3 \sum_{j=k}^N f_j \lambda^{j-k}$$

all sums appearing in (5.35) can be replaced:

$$\begin{aligned}
 (i) \quad & 3 \sum_{j=1}^k f_j \lambda^{k-j} = \sigma_k - 3f_0 \lambda^k, \\
 (ii) \quad & 3 \sum_{j=k+1}^N f_j \lambda^{N-j+k} = \lambda^k \sigma_N - \lambda^N \sigma_k, \\
 (iii) \quad & 3 \sum_{j=1}^k f_j \lambda^{N-k+j-1} = \lambda^{N-k} \tau_1 - \lambda^N \tau_{k+1}, \\
 (iv) \quad & 3 \sum_{j=k+1}^N f_j \lambda^{j-k-1} = \tau_{k+1}.
 \end{aligned}$$

Since $\sigma_N = \lambda \sigma_{N-1} + 3f_N$, and $f_N = f_0$, we have:

$$\lambda^k \sigma_N - 3f_0 \lambda^k = \lambda^{k+1} \sigma_{N-1}.$$

Therefore:

$$(5.37) \quad S(x) = \frac{1}{1-\lambda^N} \cdot T(\theta) \left[\sigma_k(1-\lambda^N) + \lambda^{k+1} \sigma_{N-1} \right] + (1-\theta)^3 f_k \\ + \frac{1}{1-\lambda^N} \cdot T(1-\theta) \left[\tau_{k+1}(1-\lambda^N) + \lambda^{N-k} \tau_1 \right] + \theta^3 f_{k+1}.$$

Whereas (5.36) is the periodic cardinal spline for the Bayes representation with f_j as parameters to be estimated, (5.37) is used for fast computation with f_j known.

6. Applications

In the first chapters we have derived the Bayes estimation procedure, especially for models linear in the parameters, and in the preceding chapter the explicit form of such a linear model, the representation of the physical relationship by cubic splines, has been developed.

Whereas in the notation of the previous chapter the knots of the spline have been assigned to x_k , throughout this chapter, the knots of the spline will be denoted by ξ_k .

The model function (5.19) can be written as:

$$(6.1) \quad S(x) = \sum_{j=0}^N S_j(x) f_j + S_c(x).$$

The $S_j(x)$ represent a superposition of the cardinal splines $A_j(x)$ and coincide with these, if the additional splines $B_0(x)$ and $B_N(x)$, which were introduced to satisfy boundary conditions, are neglected, since then (6.1) reduces to the representation (5.11). Usually the $S_j(x)$ do not differ very much from the $A_j(x)$ except for a few splines near the boundaries (Fig. 1b). The term $S_c(x)$ in (6.1) is a constant with respect to the free parameters f_j and will be zero for natural splines (more generally for $d_0=d_N=0$ in (5.19)) and in the periodic case. In the following discussion this term will not be taken into account.

6.1 Curve fitting examples based on nonperiodic splines

The first and probably most important application of the Bayes concept developed here will be the problem of curve fitting with cubic splines. Given a set of experimental data z_i at measure-

ment positions x_i , $i=1, \dots, m$, a smooth curve through the data, or even a smooth first derivative of the data, is wanted without the explicit knowledge of the function type of the physical relationship, as e.g. a sum of Gaussians, or Gaussians folded by some resolution function. The limitation to a particular function type, i.e. the usual least-squares problem, can be relaxed by representing the physical relationship by a sum of piecewise cubic polynomials joint at a uniform mesh of knots at $\xi_k = \xi_0 + k \cdot h$, $k=0, 1, \dots, N$, or, after relabelling, at ξ_v , $v=1, \dots, n$. Then the underlying model contains the total flexibility of cubic splines with their corresponding optimality properties. The variability of the model with respect to the parameters in the estimation sense is represented by the number of knots.

Usually, data smoothing with cubic splines results from a constrained variational problem [3,24]:

$$(6.2) \quad \hat{\theta}: \quad \text{Min}_{\{\theta, \lambda\}} \left\{ \int_{x_0}^{x_N} [S''(x)]^2 dx + \lambda \left[\frac{1}{m} \sum_{i=1}^m \frac{1}{\sigma_i^2} (z_i - S(x_i))^2 - D \right] \right\},$$

where the knots ξ_i coincide with the measurement points x_i , λ is a Lagrange multiplier and D is the so-called smoothing parameter. The parameter vector θ contains the coefficients of the B-spline representation. Relation (6.2) is a compromise between smoothing achieved by the integral term - its minimum value is zero for $S(x)$ equal to a straight line - and a weighted least-squares term taking into account the statistical errors via the parameter D , which should be of the order of 1 for Gaussian probability distributions. For $D \gg 1$ the spline will approach a straight line whereas for $D=0$ the spline interpolates the data. A difficulty with smoothing splines obtained according to (6.2) arises, when sharp peaks have to be fitted, since in this case the minimization of the second derivative tends to round off the peak structure, thus influencing the measured information in an undesired way.

In order to apply the Bayes criterion (3.18) to obtain an optimally fitted curve, we have to set up the measurement matrix for our model, given by equation (5.19):

$$(6.3) \quad H_{iv} = S_{v-1}(x_i), \quad i=1, \dots, m, \quad v=1, \dots, n.$$

As next step, one has to make certain assumptions about the prior mean value μ of the parameters, which is an easy task in the curve-fitting case, since a convenient choice of the component μ_v is given by interpolating between the data points adjacent to ξ_v . Finally, assumptions for the prior variance of the parameters are needed. As usually nothing is known about the correlations between the parameters, the off-diagonal elements of V_μ will be set equal to zero. The diagonal element representing the variance of μ_v may be chosen proportional to $(\sigma_i^2 + \sigma_{i+1}^2)$, where σ_i^2 is the known variance of the measurement z_i . The proportionality factor can be used as a measure of uncertainty of the prior mean values, thus reducing an undesirable influence of an incorrect prior information. In any case this factor should be of the order of 1 or greater in order to avoid a long-lasting pinning to the a-priori solution regardless of the information gained by the measurements.

Once all these quantities are given, the optimal Bayes parameters $\hat{\theta}$ are obtained according to (3.22) with covariance matrix V_θ according to (3.20). With $f_v = \hat{\theta}_{v+1}$ the optimal spline is represented by (6.1). As an example, a spectrum of scattered neutrons /25/ is fitted using two different numbers of knots (Figs. 2a,b).

Next, we discuss an example which is based essentially on the sequential character of the estimation procedure given by (4.16). This example is taken from /21/: three (Gaussian) peaks corrupted by noise with standard deviation $\sigma=0.1$. The a-priori approximation is set intentionally to twice the theoretical solution. The diagonal elements of the covariance matrix V_μ are chosen large

enough to allow for convergence to the true solution if measurements are processed. Figs. 3a-c show the state of the experiment after 1, 26, and 300 simulated measurements, where each new position has been selected at random. The plots include the measurement points, the true solution, the actually estimated solution and its corresponding error band, which is defined as in the usual least-squares case:

$$(6.4) \quad \sigma_f^2(k, x) = h^T(x) V_\theta(k) h(x),$$

where the vector $h(x)$ has components $h_v(x) = S_{v-1}(x)$ and $V_\theta(k)$ is the covariance matrix after k measurements, updated according to (4.16a). The actual error band is given by $\pm \sigma_f$.

The evolution of this experiment clearly shows convergence to the true solution. This effect is coupled with a reduction of the error band at those regions where experimental data have been taken.

As the positions of these simulated measurements are randomly distributed, this particular experiment could by no means be regarded as optimal in the sense of statistical experimental design [26, 27]. Following the idea of Wynn [20] the new point to be measured is to be placed where the actual error band has its maximum value. Figs. 4b and 4c give an impression of such an optimally planned experiment, where the same prior assumptions have been made as in the experiment shown in Figs. 3a-c. The positions of the observations concentrate at points called D-optimal points [20]. A D-optimum experimental design minimizes the determinant of the covariance matrix V_θ . As can be seen from Fig. 4c, the sequential D-optimum experimental design applied to the linear spline model (5.19) reveals, that the corresponding D-optimal points are close to the knots of the spline (in our case 26 knots have been used). Moreover, these 26 optimal positions are attained in cycles of length 26, i.e. during the first cycle all positions are found once (Fig. 4b) (some minor devi-

ations arise because the starting point is selected at random and the maximum of the error band is determined only by a grid search). After the first cycle, again all 26 positions are attained in one cycle, however in a somewhat different order, and so on. This phenomenon was also observed by Wynn /20/. Comparison of Fig. 4b with Fig. 3b shows most significantly the completely different distribution of experimental information over the total range of measurement.

This interactive performance of an experiment represents a typical example for sequential estimation /21/.

6.2 An example based on periodic splines

This section deals with an example which looks rather special, but for which the present algorithm turned out to be very successful on the one hand, and very instructive on the other hand. Suppose neutrons are scattered from a single crystal with corresponding symmetry /28/. Then the data should reflect, as a function of the scattering angle x , this symmetry in the average. Therefore, the problem will be to fit some model to the measured data and, simultaneously, to satisfy certain symmetry conditions. More precisely, in our case a cubic spline $S(x)$ is to be constructed with

$$(6.5a) \quad S(x+\pi) = S(x),$$

$$(6.5b) \quad S(\pi-x) = S(x).$$

The conditions imply:

$$(6.6a) \quad S'(0) = 0,$$

$$(6.6b) \quad S'(\frac{\pi}{2}) = 0.$$

In the following, we assume that the data are accumulated within $[0, \pi]$, where $\xi_0=0$ and $\xi_N=\pi$.

If a natural spline is fitted to the data, the above conditions would not necessarily be realized. Also any nonperiodic spline satisfying the more general boundary conditions according to (5.13) cannot be used, since (5.13) affects the second derivatives only. Therefore, we try to fit a periodic spline as defined by (5.23). However, because of the statistical character of the data, a periodic spline will only fulfil the periodicity condition (6.5a), whereas conditions (6.5b) and (6.6) will be true only approximately or, better, in the average.

The symmetry condition (6.5b), which indicates a point of reflection at $x = \frac{\pi}{2}$, suggests the following procedure: All data points corresponding to x -values within $[\frac{\pi}{2}, \pi]$ will be reflected symmetrically to $x = \frac{\pi}{2}$ into $[0, \frac{\pi}{2}]$, and vice versa with the data originally related to $[0, \frac{\pi}{2}]$. Then each interval, $[0, \frac{\pi}{2}]$ and $[\frac{\pi}{2}, \pi]$, will contain the complete measured information. With this set of data a fit has been performed with a natural spline in the region $[0, \pi]$ (Fig. 5a). This spline function satisfies conditions (6.5b) and (6.6b).

At this stage, an important observation can be made: Although the first derivative is zero at $x = \frac{\pi}{2}$ (Fig. 5b) as postulated from the symmetrization of the data, its error band differs from zero at $x = \frac{\pi}{2}$. At first sight this effect seems contradictory, because the symmetrization forces the derivative to be zero, and hence the error band should be degenerated to zero at this point. But the explanation of the finite variance is again the statistical nature of the data. The algorithm does not know the history of the origin of the data and assumes all measurements to be stochastic, i.e. when processing a measured point, this is always considered to have statistical fluctuations. Therefore, only incidentally condition (6.6b) is satisfied.

A further step towards the solution to the posed problem is achieved when fitting a periodic spline through this assembly of artificially symmetrized data (Figs. 6a,b). In this way, all conditions (6.5) and (6.6) are strictly satisfied, but again,

as all data are considered to be statistical, the error band of the first derivative does not vanish at $x=0$ and $x=\frac{\pi}{2}$. It should be noted that a periodic spline does not imply a horizontal derivative at ξ_0 and ξ_N . This condition is only realized here on account of the specific character of the data: The spline is symmetrical and periodic, therefore a nonvanishing first derivative at $x=0$ would lead to a discontinuity in this derivative.

In reality, one certainly should not proceed as described here. The true statistics inherent in the experimental data is no longer conserved when doubling the experimental results as was done above. Therefore, only the symmetrized data of one half of the total interval, say on $\left[0, \frac{\pi}{2}\right]$, should be processed. A symmetrical spline is obtained by treating all θ_v with knots ξ_v in $\left[0, \frac{\pi}{2}\right]$ as free parameters, and setting the parameters on $\left[\frac{\pi}{2}, \pi\right]$ equal to the corresponding θ_v on $\left[0, \frac{\pi}{2}\right]$. Formally, this means that the following modified linear model will be used:

$$\begin{aligned} \zeta_i &= \sum_{v=1}^{n/2} (H_{i,v} + H_{i,n-v+1}) \theta_v, & n \text{ even} \\ (6.7) \\ \zeta_i &= \sum_{v=1}^{(n-1)/2} (H_{i,v} + H_{i,n-v+1}) \theta_v + H_{i,(n+1)/2} \theta_{(n+1)/2}, & n \text{ odd.} \end{aligned}$$

We note that with $\xi_0=0$ and $\xi_N=\pi$ the knots are symmetrical with respect to $x=\frac{\pi}{2}$. The calculation of the matrix elements of (6.7) from a nonperiodic natural spline yields a zero derivative and a corresponding zero error band for the derivative at $x=\frac{\pi}{2}$ (Figs. 7a,b). Finally, if a periodic spline is used for the calculation of H according to (6.7), in addition, at $x=0$ a zero derivative with zero error band is obtained (Figs. 8a,b), thus satisfying all conditions (6.5) and (6.6) in a statistically consistent way. The error band is zero at the relevant points, because conditions (6.5) and (6.6) are fulfilled by construction, and not incidentally.

6.3 Deconvolution of experimental data

Frequently the problem of deconvolution occurs in experimental science where the theoretically interesting information $f(x;\theta)$ cannot be measured directly, but is hidden in the data due to a convolution with some given "resolution function". Mathematically, a special relation of the type (1.1) holds:

$$(6.8) \quad \phi(x;\theta) = \int_{-\infty}^{\infty} g(x-x') f(x';\theta) dx' ,$$

where $g(x)$ is the resolution function - assumed to be known explicitly, and $f(x;\theta)$ involves the scientific relationship to be investigated. From a mathematical point of view the solution of an equation of type (6.8) is not a real problem as, for instance, Fourier transform can be applied /29/. Since, however, $\phi(x;\theta)$ is in fact represented by the stochastic quantity z as discussed in the previous chapters, the purely mathematical solution of (6.8) turns out to be one of the well-known "incorrectly posed" problems /30/.

The numerical solution $f(x;\theta)$ is extremely sensitive to the mathematical procedure and the data. Therefore, statistical principles have to be applied to such problems in order to obtain scientifically reasonable and consistent results out of any mathematical procedure. In those cases where the function type $f(x;\theta)$ is given explicitly, the problem is reduced to the determination of parameters by some estimation technique discussed above.

In some applications solutions of the convolution problem have been developed using the spline concept (6.2) /3/ and the

Bayes concept (3.18) with a linear model stair-case approximation to $f(x;\theta)$ /14/. Extensive literature in this field can be found in these references.

In the following we shall apply the new formalism developed in the previous chapters which represents an efficient synthesis of advanced statistics and estimation theory - Bayesian methods and sequential algorithms - with the flexible features of splines.

It is assumed that $f(x;\theta)$ can be represented by (6.1) with $S_c(x) = 0$:

$$(6.9) \quad f(x;\theta) = \sum_{v=1}^n S_{v-1}(x) \theta_v ,$$

with $f(\xi_\mu;\theta) = \theta_\mu$ at the knots according to a natural or periodic cardinal spline expansion. Expression (6.9) inserted into (6.8) yields:

$$(6.10) \quad \zeta_i = \sum_{v=1}^n \theta_v \int_{-\infty}^{\infty} g(x_i - x') S_{v-1}(x') dx' .$$

This preserves the linear model with the measurement matrix:

$$(6.11) \quad H_{iv} = \int_{-\infty}^{\infty} g(x_i - x') S_{v-1}(x') dx' .$$

Hence, the same arguments hold as in the pure curve-fitting case of 6.1 and 6.2. In addition to the estimate of the solution $f(x;\theta)$ error bands are obtained from the statistical variances for the course of $f(x;\theta)$ and its higher derivatives. If no prior information is available ($V_\mu^{-1} = 0$), this method reduces to the usual least-squares solution with the linear model representing $f(x;\theta)$ as a superposition of n cubic splines. Consequences of this model will be illustrated later.

For numerical purposes it should be noted that usually a realistic resolution function $g(x)$ can be truncated for large

$|x|$, so that its range of definition is $[-A, B]$ with two positive constants A, B . For a Gaussian type resolution function as kernel of the integral equation (6.8), for instance, these constants may be chosen equal to $3 \sigma_G$, where σ_G is the standard deviation of the Gaussian function. If the observations z_i are taken within the experimental range $x_a \leq x_i \leq x_b$, then the knots should be placed such that ξ_0 and ξ_N are outside the range of measurements by at least the corresponding widths of $g(x)$ in order to avoid disturbing truncation effects [3, 31/].

Whenever use is made of the sequential algorithm, the prior information with regard to θ and V_θ must be supplied. In the examples discussed below we have always chosen μ_ν from interpolation between z_i and z_{i+1} with $x_i \leq \xi_\nu \leq x_{i+1}$. This choice, in general, will produce a bias; in order to reduce the biasing effect as much as possible, reasonably large values for the diagonal elements of V_μ^{-1} should be used.

In the following, the influence of certain quantities of interest will be studied, where the model function $f_M(x)$ is assumed to be known in order to get an understanding of what one could expect from the present algorithm when applied to realistic problems. The procedure always follows the same line: A particular model function $f_M(x)$ is folded with a given resolution function $g(x)$ - $g(x)$ is assumed to be known also when dealing with real life problems - then the integral $\phi_M(x_i)$ resulting from (6.8) for a specific measurement point x_i , is corrupted by some Gaussian noise ϵ_i with zero mean and given variance σ_i^2 . The resulting data $z_i = \phi_M(x_i) + \epsilon_i$ are then processed according to formulae (6.9), (6.10), (6.11). The estimate of the spline solution with its error band can be compared with the theoretical model originally inserted.

6.3.1 The effect of increased smearing

The most typical kind of convolution problem is involved with a resolution function $g(x)$ of Gaussian type, with σ_G as a measure of its width; the model function $f_M(x)$ may also be represented by a (normalized) Gaussian with width σ_M . Then the integral $\phi_M(x)$ can be calculated explicitly yielding with $f_M(x)$ centered at $\eta=2$:

$$(6.12) \quad \phi_M(x) = \frac{1}{\sqrt{2\pi} \sqrt{\sigma_M^2 + \sigma_G^2}} \exp \left[- \frac{(x-\eta)^2}{2(\sigma_M^2 + \sigma_G^2)} \right].$$

The following set of parameters was chosen: $\sigma_M=0.25$, $m=150$ as number of observations z_i sampled with randomly generated abscissae x_i on $[1,3]$, normal probability distribution for z_i with relative standard deviation $\sigma=0.01$, $n=20$ as number of knots with $\xi_0=0.5$, $\xi_N=3.5$. With this set, three simulation experiments were performed, where the width of the resolution function had the values $\sigma_G=0.10$, $\sigma_G=0.15$, and $\sigma_G=0.20$ respectively (Figs. 9a-c). For higher values of σ_G a solution $f(x;\theta)$ was no longer meaningful. To achieve a smoother curve $f(x;\theta)$ for $\sigma_G=0.20$, one may reduce the number of knots (Fig. 10a). With further reduction of the number of parameters, the spline will start to deviate systematically from the physical model $f_M(x)$. A reasonable improvement may be obtained by increasing the statistical accuracy, say $\sigma=0.0001$ (Fig. 10b), or by an equivalent increase of the experimental sample.

6.3.2 Dependence on the number of knots

Probably of most importance is the reasonable choice of the number n of knots- or free parameters in the estimation sense. A general rule might be established from the sampling theorem /32/. In the cases discussed here the right number was found by arguing with the structure of the function in question, and by

trial, since the number of necessary knots depends on the structure inherent in the scientific problem. Because of the equidistance of the knots, any pronounced peak may require an adequately large number of splines over the whole range of interest.

This difficulty becomes apparent, when the extreme cases of small and large n are considered. If the interval between the knots is large with respect to the fine structure of the underlying process, the smoothing property of the splines will be dominant and essential features of the physical structure will not be reproduced. If on the other hand the number of splines becomes comparable to the number of observations, the minimizing least-squares criterion is superior, and the splines may follow too closely the statistical fluctuations. The consequences of insufficient knots are illustrated in Figs.11a-c , which in this special example reveal, that only for about 35 and more splines the physical details will be reproduced. The theoretical curve, which this example is based on, is an experimentally verified curve representing a Compton profile of γ -rays to be scattered by electrons in copper, which is folded with a hypothetical resolution function of Gaussian type /33/.

A remark concerning the error band seems quite useful at this point. It is not a contradiction that the error band for knots less than about 35 (Figs. 11a,b) does not completely contain the true solution; at first glance, one might expect that in cases, where the spline solution deviates appreciably from the true solution, the error band should be correspondingly large. But, on the one hand the error band maps the course of the standard deviation corresponding to the solution, and therefore the theoretical function can be expected to be contained within the error band up to about 60% only because of statistical reasons. On the other hand, if the variances of the data z_i are small, a narrow error band will always result, regardless of the underlying model. A significantly weak agreement of solution and theoretical curve indicates in any way, that the small number

of knots is a "wrong" model in the estimation-theoretical sense; however, a statement in this direction would be the subject of hypothesis testing or decision theory /21, 34/.

Another sequence of examples (Figs. 12a-c) with a hypothetical compton profile consisting of two joint parabolas demonstrates that, depending on the model and the statistical quality of the data, the range of the optimal number of knots can be strongly limited. In any case, again one has to perform an adequate statistical hypothesis test in order to establish the validity or failure of the model.

6.4 Integral equations of non-convolution type

Equation (6.8) can be regarded as a linear stochastic integral equation of the convolution type. With the developed formalism it is no problem to generalize the range of applicability of the present estimation concept to the following type of equations:

$$(6.13) \quad \phi(x;\theta) = \int_{-\infty}^{\infty} g(x') f(r(x,x');\theta) dx',$$

where $r(x,x')$ is some given monotonous function of the distance between x and x' . Physical examples are known from the field of small-angle scattering, where the distortion effect in the scattering law is due to the finite collimation of the beam and, especially in the case of neutron small-angle scattering, also due to the polychromatic wavelength distribution. The three cases of interest /3/ include:

$$(6.14) \quad \begin{array}{ll} r(x,x') = x-x', & \text{slit-width correction} \\ r(x,x') = (x^2+x'^2)^{1/2}, & \text{slit-height correction} \\ r(x,x') = x/x', & \text{wavelength correction.} \end{array}$$

These types of equations have already been investigated in detail /3,31,35/. It should be mentioned here that the Bayes algorithm has been applied to these problems, in general providing good results as far as the relative experimental error does not exceed 1%.

A special example for a slit-height correction is shown in Fig. 13 where a theoretical model $f_M(x)$, the hardsphere scattering function /3/ with radius $R=10$ and $S_0=10^5$, was used:

$$(6.15) \quad f_M(x) = S_0 \left[\frac{3(\sin u - u \cos u)}{u^3} \right]^2, \quad u=Rx.$$

The kernel of the integral equation was a step function for $x \in [0,2]$, leading to the following relation:

$$(6.16) \quad \phi_M(x) = \int_0^2 f_M((x^2+x'^2)^{1/2}) dx'.$$

The final example (Fig. 14) involves the wavelength smearing effect, again using the hardsphere model (6.15):

$$(6.17) \quad \phi_M(x) = C \int_0^\infty \exp \left[- \frac{(x' - \eta)^2}{2 \sigma_G^2} \right] f_M(x/x') dx'$$

with $\sigma_G=0.2$, $\eta=1.0$, and C is chosen to normalize the kernel on $0 \leq x < \infty$. Due to the strong distortion effect caused by "long wavelengths" and due to the rather large random noise with a relative error of 2%, the model can be reproduced by the estimation algorithm up to the second minimum only, which nevertheless represents a remarkable progress with respect to practical applications.

7. Conclusions

We have developed a numerical technique based on Bayes sequential estimation theory, which gives optimal cubic cardinal spline functions smoothing normally distributed experimental data. Especially, this algorithm allows the performance of online data processing, where experimental information is updated after each new observation; additionally, as a result of this possibility, the control of the further development of the experiment according to statistical criteria from the theory of experimental designs can be realized.

Further applications include the calculation of smooth derivatives of stochastic data, the use of periodic cardinal splines, and the (generalized) deconvolution of experimental data.

The efficiency of the algorithm has been demonstrated utilizing computer simulated and experimental data in a great variety of important examples from scientific practice.

This sequential technique may prove to be of great use in different fields of experimental data processing. We would like to emphasize the fact that the computer program based on the algorithm which has been developed in this report is extensively used in solid state and scattering research. The program may be obtained from the authors by request.

Acknowledgment

We would like to thank Drs. C. de Polignac and O. Glatter for discussions on this subject and Drs. G. Bauer, B. Lengeler and M. Shaanan for making available to us their experimental data. We would like to express our gratitude to Mrs. U. Ehrhart for contributions to the programming of the spline algorithms.

Appendix

Functional Characteristics of the Computer Program BAYSIC

In its original version, the computer program BAYSIC contains 15 subprograms. It allows to handle the pure curve fitting case ($ITYP=0$) and different types of integral equations ($ITYP \neq 0$). Both cases can be applied to computer simulated data ($IDATA=0$) and to experimental data ($IDATA \neq 0$). Since curve fitting problems arise very frequently, the version BAYSIC/0 containing only 6 subprograms has been coded separately. A brief description of this latter program will be given first.

A1. Input information for BAYSIC/0

Only the meaning of the most important input parameters will be given here, for details the program description should be consulted.

MZ data points $Z(J)$ with abscissae $Y(J)$ are read in with variances $VZZ(J)=Z(J)$ if $IDATA=1$, and $VZZ(J)$ also read in if $IDATA>1$. $IPER=1$ must be set in order to fit a periodic cardinal spline, for $IPER=0$ a nonperiodic natural spline ($KOND=0$) or general spline satisfying condition (5.13) with given DO , DN , $MUNU$, MUN ($KOND=1$) will be used. $NPAR$ denotes the number of knots of the spline, XO and XN correspond to the knots ξ_0 and ξ_N , respectively. The a-priori information $FEST(I)$ and $VNULL(I)$, $I=1, NPAR$, i.e. the parameter vector μ and the variances given by the diagonal elements of V_μ , will be read in if $MPRIO=0$, or calculated by SUBROUTINE (S.R.) APRIL if $MPRIO \neq 0$. The actual value of $MPRIO$ then is a proportionality factor for $VNULL(I)$. In order to perform the interpolation correctly, S.R. APRIL requires an ordered sequence of the data ($Y(J) \leq Y(J+1)$), therefore, if the abscissae are not in ascending order, $ISORT=1$ will cause the sorting, otherwise $ISORT=0$.

The output is controlled by the following input parameters. NYP gives the number of spline points to be calculated equidistantly between YPN and YPM. The calculation of the spline and/or its higher derivatives is determined by the logical variables L0, L1, L2, L3, where .TRUE. means calculation. IPRINT allows various options for printing the data, the spline, and its higher derivatives. IPLOT=0 means that plots are produced for all cases which are marked .TRUE., IPLOT<0 means no plot.

A2. The program structure

The MAIN-program consists of preparatory calculations including the call of S.R. APRIL, the main call of S.R. BAYSAL, and some concluding calculations for the output.

S.R. BAYSAL refers first to S.R. BASEST (see below), which performs the complete estimation procedure yielding $F(I)$, $I=1, NPAR$, as optimal parameters $\hat{\theta}_v$ of eq. (3.22) and calculates the error bands for the spline and its higher derivatives given by $VS(M)$, $VS1(M)$, $VS2(M)$, $VS3(M)$, $M=1, NYP$, according to eq. (6.4). As next step, since the $F(I)$ are known, the spline $S(M)$ and all its derivatives $S1(M)$, $S2(M)$, $S3(M)$, $M=1, NYP$, can now be calculated using the fast algorithms (5.21) or (5.37), which are coded in S.R. SPLESH. Hereafter, S.R. BAYSAL returns to the MAIN-program, where finally the curves to be plotted are given by $S(M)$, $S(M)+SQRT(VS(M))$ for the spline and its error band, and analogous expressions for the derivatives.

S.R. BASEST essentially consists of two parts: a) the execution of the sequential estimation procedure, b) the calculation of all relevant error bands.

In part a) the vector $h(x)$ of (4.16) determined by the linear model is calculated for a single observation point $Y(J)$ in S.R. HAMAT, which is based on the expressions (5.19) and (5.36) for the spline models. The components of $h(x)$ are stored in

HAMA(I,1,1), I=1,NPAR. With these elements, together with the observation value $Z(J)$ and its variance $VZZ(J)$, the updating scheme (4.16) is applied in S.R. BAYEST to $XK(I)$ and $VK(I,L)$, $I,L=1,NPAR$, defining the quantities $\hat{\theta}(k)$ and $V_{\theta}(k)$ of (4.16a) and (4.16b), which yields $XK1(I)$ and $VK1(I,L)$. For $J=1$ we have, of course, $XK(I)=FEST(I)$, $VK(I,I)=VNULL(I)$. After $J=MZ$ observations are processed, the sequential estimation is completed.

Part b) of S.R. BASEST calculates the corresponding error bands according to eq. (6.4). The array $YPOST(M)$, $M=1,NYP$, which is equal to $YP(M)$ of the MAIN-program, designates all points x at which the spline $S(M)$, the derivatives $S1(M)$, $S2(M)$, $S3(M)$, and therefore also the error bands are to be calculated. The expressions for $VS(M)$, $VS1(M)$, $VS2(M)$, $VS3(M)$ require the knowledge of the vector $h(x)$ and its derivatives at positions $YPOST(M)$. This is achieved by calling S.R. HAMAT yielding $HAMA(I,1,K)$, $I=1, NPAR$ and $K=1,4$. Hereafter, return to the MAIN-program ends the curve fitting procedure as performed by BAYSIC/O.

A3. The general program BAYSIC

The principle structure as described so far remains unaltered when treating the general problem (6.13). The only changes consist in generating the simulated data, and computing the modified matrix elements (6.11) or equivalent expressions from (6.13) and (6.14).

The simulation of the experimental data requires a model function $f_M(x)$, which is computed in FUNCTION FUNC, where LFUNC selects between different models. The three arguments $r(x,x')$ of (6.14) are calculated in FUNCTION GG, here $ITYP \neq 0$ is used to choose between these argument types. FUNCTION SKERN allows to apply different kernels $g(x)$ of (6.13) depending on IKERN, with SIG, SIG1, SIG2, SIG3 as characteristic parameters for special kernel types. FUNCTION FINT, the total integrand of (6.13), is used as

external function for the integration routine S.R. DQUADR of KFA-library FORTK which also needs NQU, KQU as number of steps for a specific Newton-Cotes quadrature and AANF, BEND as the integration limits, i.e. beyond which the kernel is truncated to zero. The resulting FUNCTION FANC, which is $\phi_M(x_i)$ of section 6.4, is corrupted with Gaussian noise from S.R. RANDN (this is identical with GAUSS of the IBM scientific subroutine package SSP) with standard deviation WVARN yielding the simulated data $Z(J)$ at position x_i stored in $Y(J)$.

To apply the Bayes algorithm it is necessary to evaluate the new matrix of the modified linear model (6.13). This implies to replace, in part a) of S.R. BASEST, the call of S.R. HAMAT by a call of S.R. HAMAST; part b) remains unchanged, because there S.R. HAMAT refers to the solution. For a measurement position $Y(J)$, S.R. HAMAST returns the vector $HAMA(I,1,1)$ of NPAR integrals according to (6.11) or from (6.13). The integration itself is performed with S.R. QUADA, a modification of S.R. DQUADR, which allows for the fact that several integrals must be calculated at one call. Like S.R. DQUADR, S.R. QUADA needs the integrands as external functions. These are calculated in FUNCTION FCT and consist of the products

$$g(x) \cdot S_v(r(x_i, x)),$$

where $g(x)$ is equivalent to FUNCTION SKERN(X), and the spline $S_v(r(x_i, x))$ is obtained by calling S.R. HAMAT at "position" $GG(X)$ where the actual measurement position $Y(J)$ is supplied to $GG(X)$ via a COMMON-statement. Once the new matrix elements are known, S.R. BASEST continues as described above for BAYSIC/O.

A4. Storage requirements

The standard versions of BAYSIC and BAYSIC/O use arrays dimensioned such that

$$MZ \leq 999, \quad NPAR \leq 40, \quad NYP \leq 200.$$

There are 4 single precision vector arrays of length 999, 12 vector arrays and 2 matrix arrays of dimension 40, most of these arrays in double precision, and 17 single precision vector arrays of length 200. In most cases it should be possible to reduce the total length drastically, if necessary.

References

- /1/ Bard, Y.: "Nonlinear Parameter Estimation", Academic Press, New York-London, 1974.
- /2/ Deutsch, R.: "Estimation Theory", Prentice Hall, Englewood Cliffs, New Jersey, 1965.
- /3/ Schelten, J., and Hoßfeld, F.: "Application of Spline Functions to the Correction of Resolution Errors in Small-Angle Scattering", J. Appl. Cryst. 4 (1971), 210.
- /4/ Sage, A.P., and Melsa, J.L.: "Estimation Theory with Applications to Communications and Control", McGraw-Hill, New York, 1971.
- /5/ Lee, R.C.K.: "Optimal Estimation, Identification, and Control", Research Monograph No. 28, The M.I.T. Press, Cambridge, Ma, 1964.
- /6/ Pachelski, W.: "On the Filtering Technique of the Least Squares Estimation of Parameters", Comp. Phys. Comm. 4 (1972), 40.
- /7/ Nilson, E.N.: "Cubic Splines on Uniform Meshes", Comm. ACM 13 (1970), 255.
- /8/ Winkler, R.L.: "Introduction to Bayesian Inference and Decision", Holt, Rinehart, and Winston, London-New York, 1972.
- /9/ Hays, W.L., and Winkler, R.L.: "Statistics", Holt, Rinehart, and Winston, London-New York, 1970, Vol. I.
- /10/ Kendall, M.G., and Stuart, A.: "The Advanced Theory of Statistics", Vol. 2, Ch. Griffin & Cie., London, 1961.

- /11/ Hoel, P.G.: "Introduction to Mathematical Statistics", Wiley, New York, 1971.
- /12/ Dwyer, P.S.: "Generalizations of a Gaussian Theorem", Ann. Math. Stat. 29 (1958), 106.
- /13/ Savage, L.J.: "Foundations of Statistics", Wiley, New York, 1954.
- /14/ Mendes, M., and de Polignac, C.: "Recursive Bayes Deconvolution in Physical Experiments", Acta Cryst. A29 (1973), 1.
- /15/ Kalman, R.E.: "A New Approach to Linear Filtering and Prediction Problems", Trans. ASME, Series D: Journal of Basic Engineering 82 (1960), 35.
- /16/ Kowalik, J., and Osborne, M.R.: "Methods for Unconstrained Optimization Problems", Am. Elsevier Publ. Comp., New York, 1968.
- /17/ Mendel, J.M.: "Discrete Techniques of Parameter Estimation", Control Theory Vol. 1, Marcel Dekker, New York, 1973.
- /18/ Neuburger, E.: "Einführung in die Theorie des linearen Optimalfilters (Kalman-Filter)", R. Oldenbourg, München-Wien, 1972.
- /19/ Krogmann, U.: "Das Kalman-Filter. Zusammenfassende Darstellung und Diskussion der Filtergleichungen", Intern. Elektronische Rundschau 28 (1974), Pt. 1: 45, Pt. 2: 71.
- /20/ Wynn, H.P.: "The Sequential Generation of D-Optimum Experimental Designs", Ann. Math. Stat. 41 (1970), 1655.
- /21/ Hoßfeld, F.: "Entscheidungstheoretische Aspekte physikalischen Experimentierens", Kernforschungsanlage Jülich, Jül-Report, Jül-930-FF, März 1973.

- /22/ Ahlberg, J.H., Nilson, E.N., and Walsh, J.L.: "The Theory of Splines and Their Applications", Academic Press, New York - London, 1967.
- /23/ Greville, T.N.E.: "Theory and Applications of Spline Functions", Academic Press, New York-London, 1969.
- /24/ Reinsch, C.H.: "Smoothing by Spline Functions", Numer. Math. 10 (1967), 177.
- /25/ Shaanan, M., private communication (1975).
- /26/ Fedorov, V.V.: "Theory of Optimal Experiments", ed. by W.J. Studden and E.M. Klimko, Academic Press, New York - London, 1972.
- /27/ Ogawa, J.: "Statistical Theory of the Analysis of Experimental Designs", Marcel Dekker, New York, 1974.
- /28/ Bauer, G.S.: "Diffuse Neutronenstreuung an Verzerrungsfeldern von substitutionell und interstitiell gelösten Fremdatomen in Metallen", Dissertation, Kernforschungsanlage Jülich, Jül-Report, Jül-1158, Januar 1975.
- /29/ Porteus, J.O.: "Optimized Method for Correcting Smearing Aberrations: Complex X-Ray Spectra", J. Appl. Phys. 33 (1962), 700.
- /30/ Turchin, V.F., Kozlov, V.P., and Malkevich, M.S.: "The Use of Mathematical-Statistics Methods in the Solution of Incorrectly Posed Problems", Soviet Physics Uspekhi, 13 (1971), 681.
- /31/ Hoßfeld, F.: "The Correction of Resolution Errors in Small Angle Scattering Using Hermite Functions", Acta Cryst. A24 (1968), 643.

- /32/ Fischer, F.A.: "Einführung in die statistische Übertragungstheorie", BI-Hochschultaschenbücher, 130/130a, Bibliographisches Institut, Mannheim, 1969.
- /33/ Lengeler, B., private communication.
- /34/ Lehmann, E.L.: "Testing Statistical Hypotheses", Wiley, New York, 1959.
- /35/ Shaffer, L.B., and Hendricks, R.W.: "Optimization of the Kratky Small-Angle X-Ray Scattering Collimation System Parameters", Oak Ridge National Laboratory, Report, ORNL-TM-4278, September 1973.

Figure Captions

- Fig. 1a: A complete set of B-splines for the class of natural cubic spline functions with the knots $(0,1,\dots,5)$, in the upper part.
- Fig. 1b: A complete set of cardinal splines for the class of natural cubic spline functions with the knots $(0,1,\dots,5)$, in the lower part.
- Fig. 2a: Typical curve fitting example for a spectrum of scattered neutrons by a natural cardinal spline with 20 knots, in the upper part.
- Fig. 2b: The same as Fig. 2a, with 40 knots, in the lower part.
- Fig. 3a: Simulated sequential experiment after one observation with the true model (full curve), the a-priori information and its error band, in the upper part.
- Fig. 3b: State of the experiment of Fig. 3a after 26 randomly selected observations, in the middle.
- Fig. 3c: The same as Fig. 3b, after 300 observations, in the lower part.
- Fig. 4a: Simulated sequential experiment after one observation with the true model (full curve), the a-priori information and its error band, in the upper part.
- Fig. 4b: State of the experiment of Fig. 4a after 26 observations selected along the rules of a D-optimum sequential experimental design, in the middle.

- Fig. 4c: The same as Fig. 4b, after 300 observations, in the lower part.
- Fig. 5a: Natural spline with 12 knots fitting a set of symmetrically reflected data of scattered neutrons, in the upper left.
- Fig. 5b: The first derivative of the spline of Fig. 5a, in the lower left.
- Fig. 6a: The same as Fig. 5a using a periodic spline, in the upper right.
- Fig. 6b: The first derivative of the spline of Fig. 6a, in the lower right.
- Fig. 7a: Symmetrized natural spline fitting the left half of the data of Fig. 5a, in the upper left.
- Fig. 7b: The first derivative of the spline of Fig. 7a, in the lower left.
- Fig. 8a: The same as Fig. 7a using a symmetrized periodic spline, in the upper right.
- Fig. 8b: The first derivative of the spline of Fig. 8a, in the lower right.
- Fig. 9a: Deconvolution of a Gaussian model function with width $\sigma_M = 0.25$ from simulated measurements with a relative error of 1% obtained from a convolution with a Gaussian kernel of width $\sigma_G = 0.10$. The solution (with 20 knots) and its derivative are shown separately in the upper part.
- Fig. 9b: The same as Fig. 9a, with $\sigma_G = 0.15$, in the middle.

- Fig. 9c: The same as Fig. 9a, with $\sigma_G = 0.20$, in the lower part.
- Fig. 10a: The same as Fig. 9a, with $\sigma_G = 0.20$ and 15 knots, in the upper part.
- Fig. 10b: The same as Fig. 9a, with $\sigma_G = 0.20$ and a relative error of 0.01%, in the lower part.
- Fig. 11a: Deconvolution of an experimentally given compton profile from simulated measurements with a relative error of 1% obtained from a convolution with a Gaussian kernel 10 times smaller in width than the physical model (20 knots), in the upper part.
- Fig. 11b: The same as Fig. 11a, with 30 knots, in the middle.
- Fig. 11c: The same as Fig. 11a, with 35 knots, in the lower part.
- Fig. 12a: Deconvolution of a hypothetical compton profile of two parabolas from simulated measurements with a relative error of 0.1% obtained from a convolution with a Gaussian kernel of width 0.42 times the width of the physical model. The solution and its derivative for 17 knots are shown at the upper left and the upper right, respectively.
- Fig. 12b: The same as Fig. 12a, with 20 knots, without derivative, in the lower left.
- Fig. 12c: The same as Fig. 12a, with 22 knots, without derivative, in the lower right.

Fig. 13: Slit-height correction for the hardsphere scattering function with simulated data of 1% relative error (25 knots). Range of experimental data: $0 \leq X \leq 2.0$, range of knots: $0 \leq X \leq 2.2$.

Fig. 14: Wavelength correction for the hardsphere scattering function with simulated data of 2% relative error (22 knots). Range of experimental data: $-0.2 \leq X \leq 2.0$, range of knots: $-0.2 \leq X \leq 2.2$.

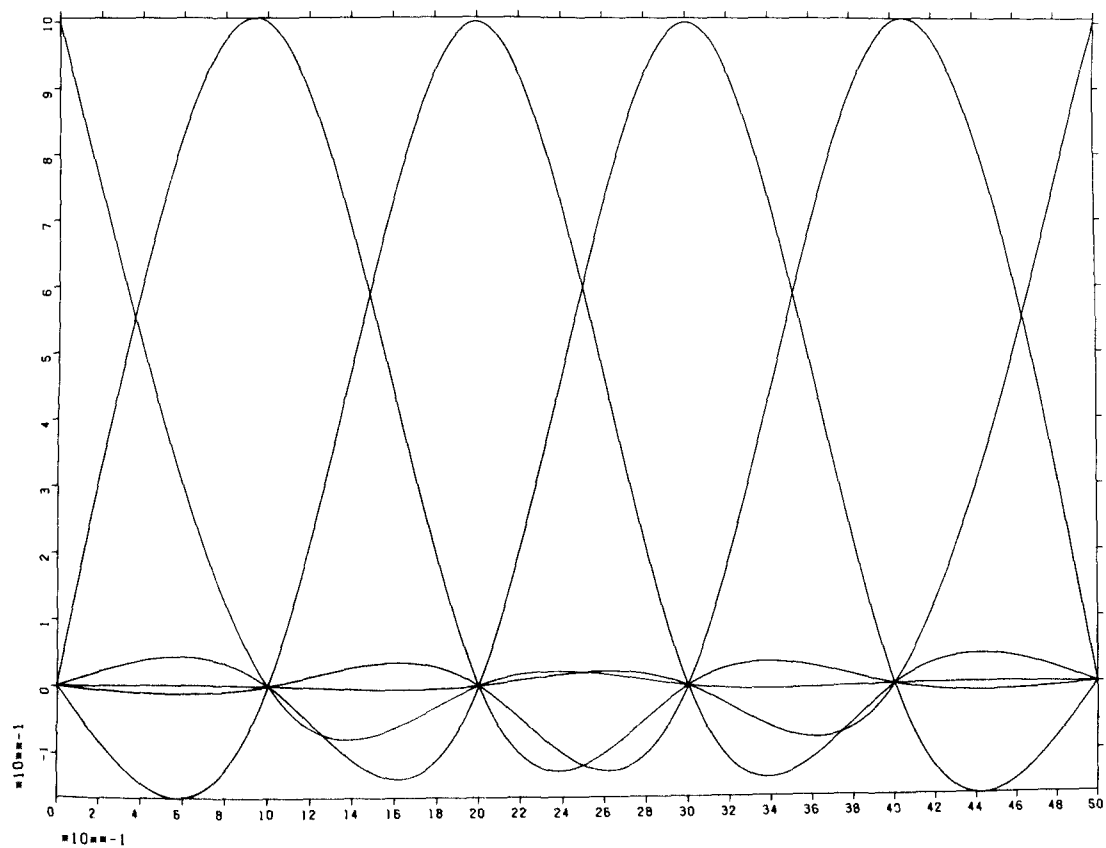
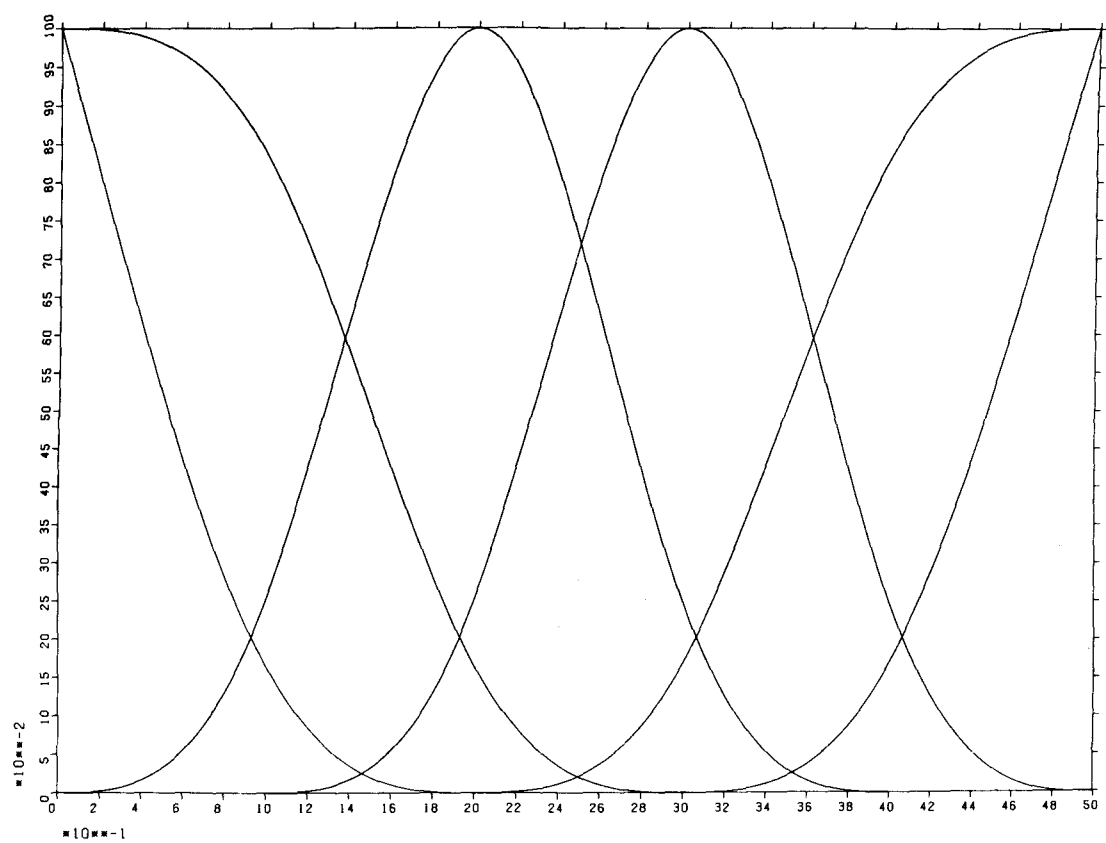


Fig. 1

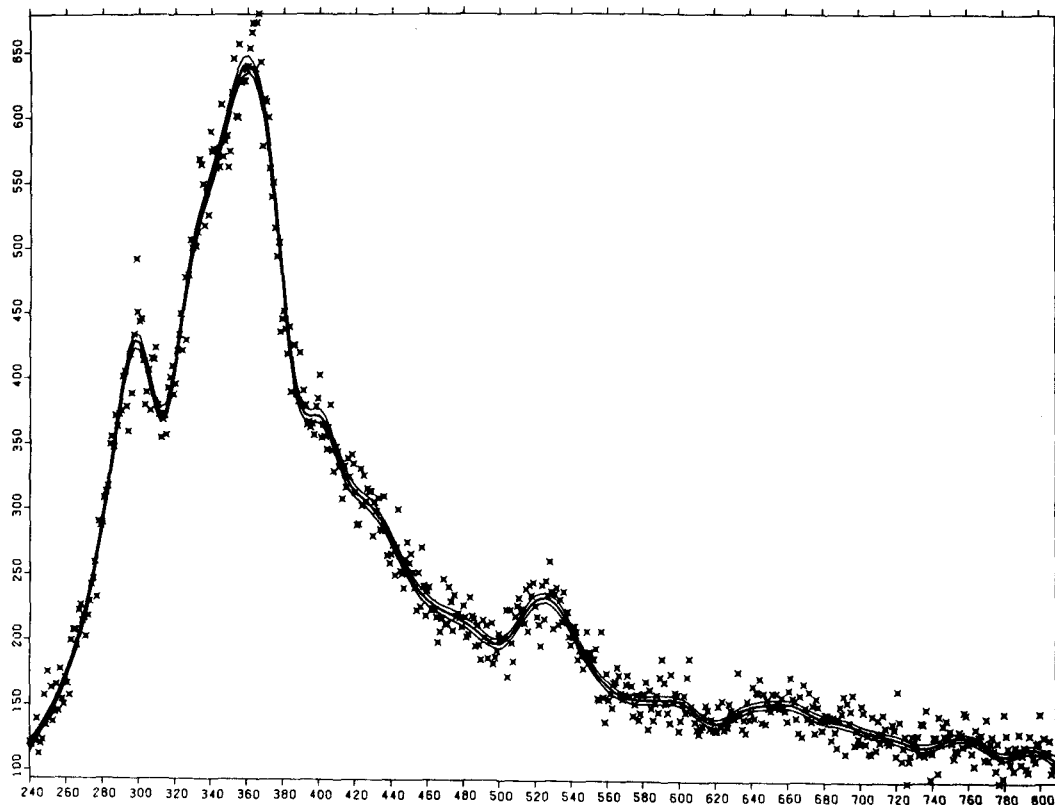
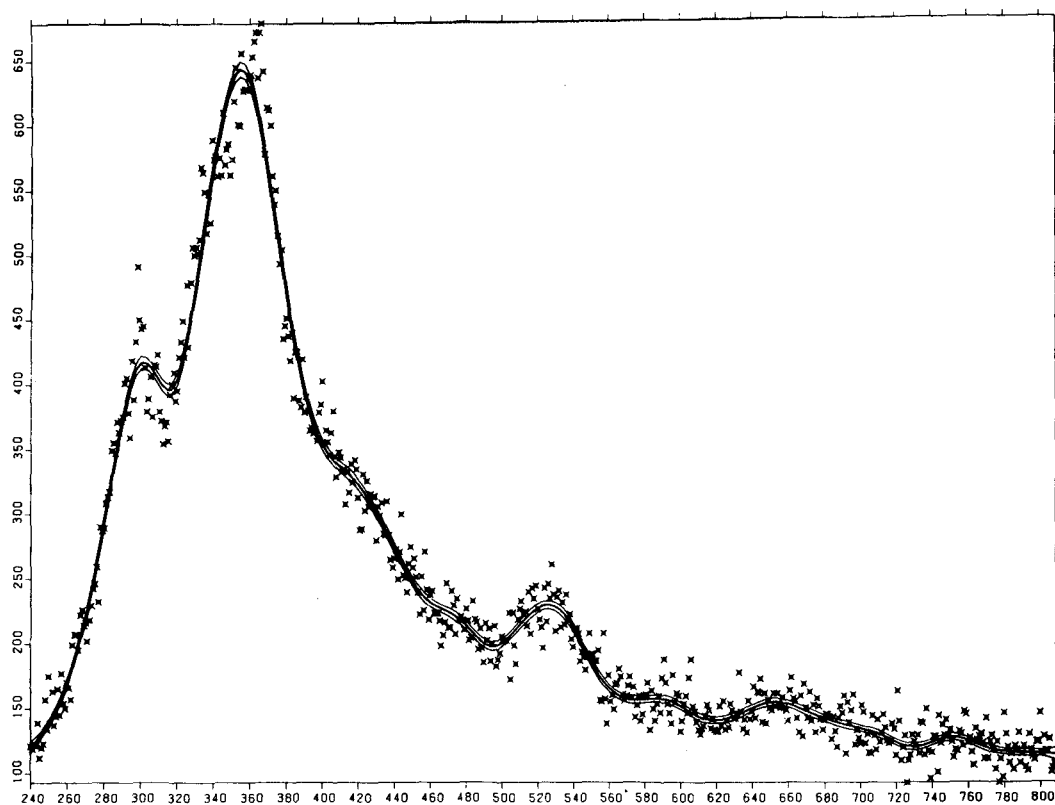


Fig. 2

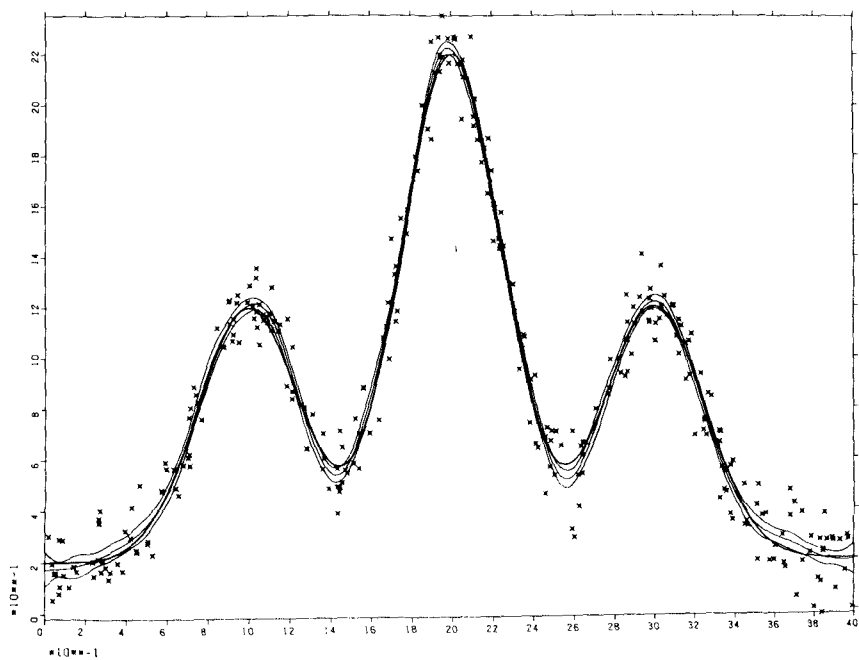
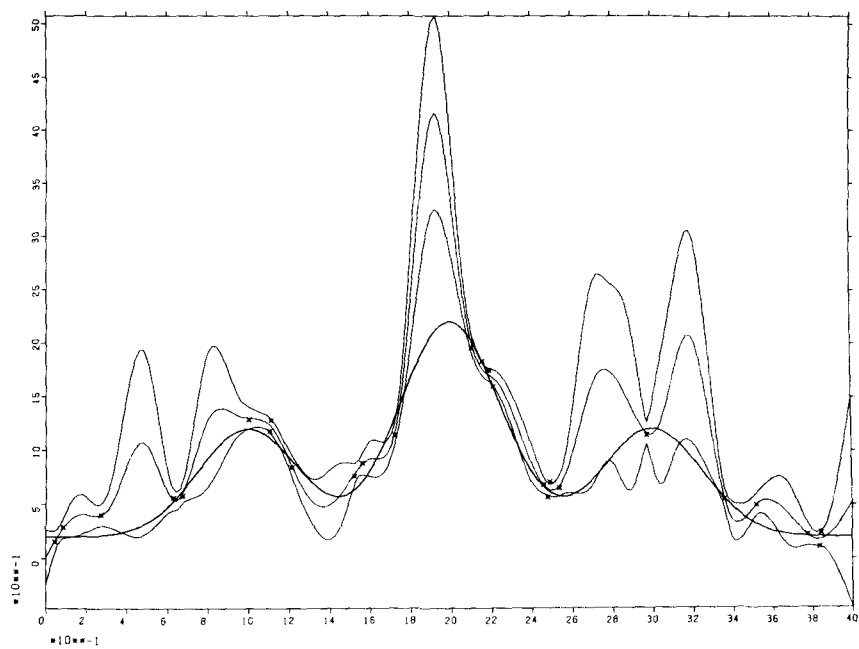
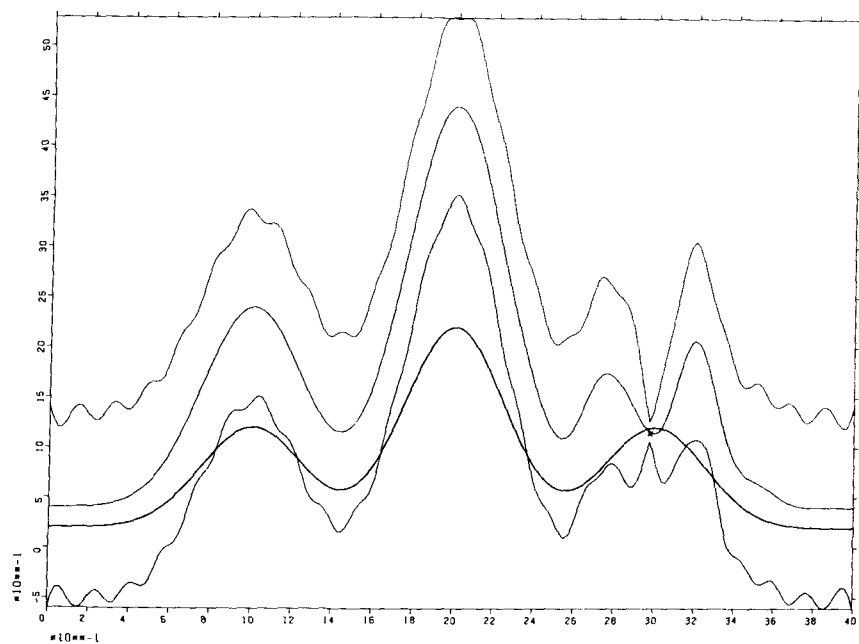


Fig. 3

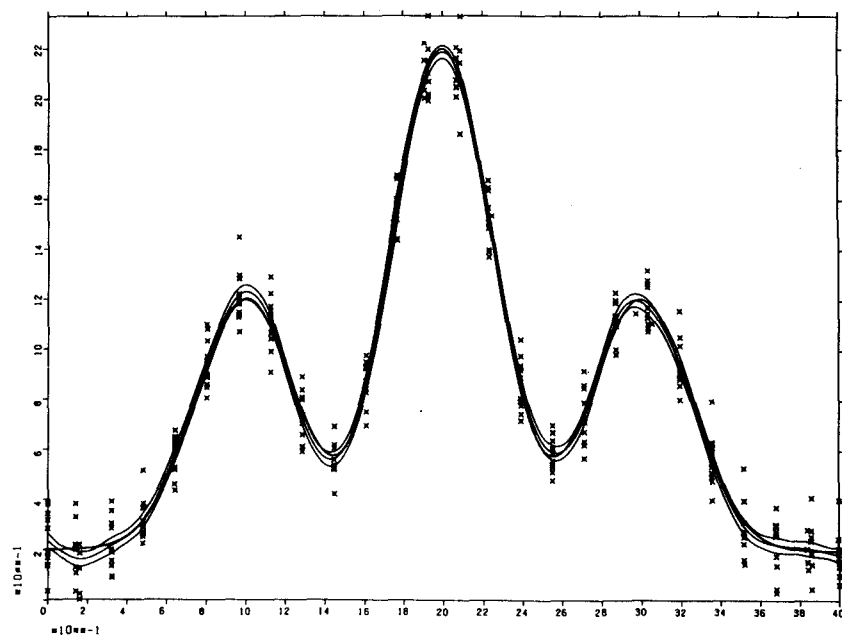
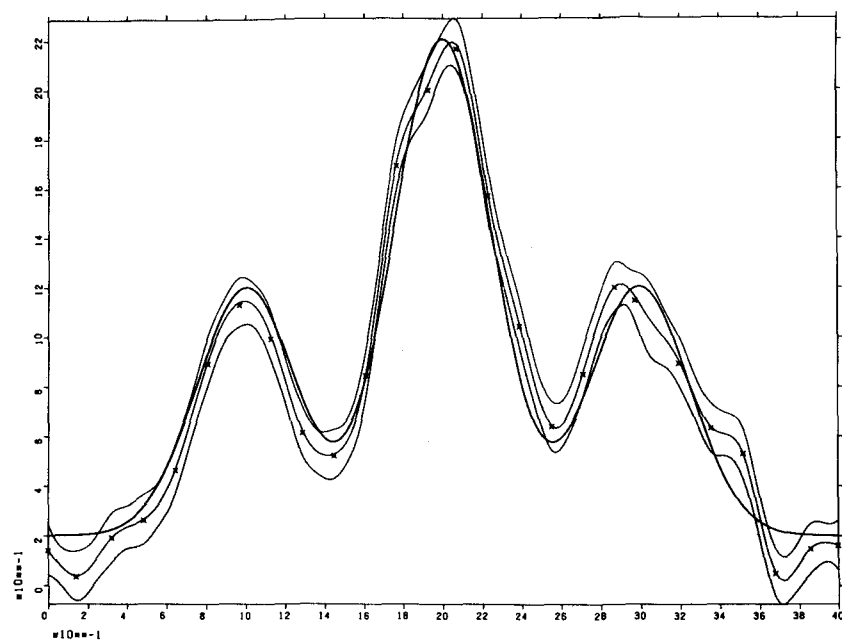
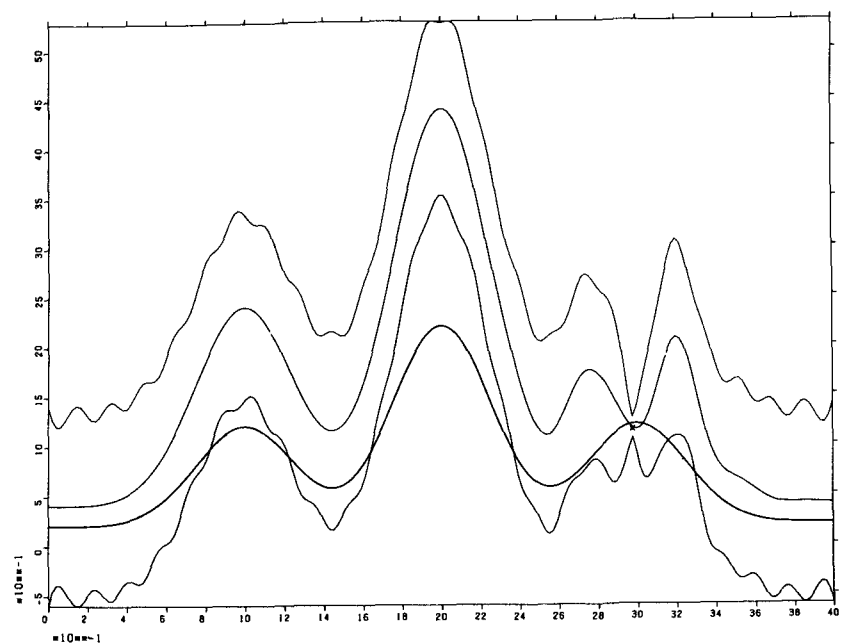


Fig. 4

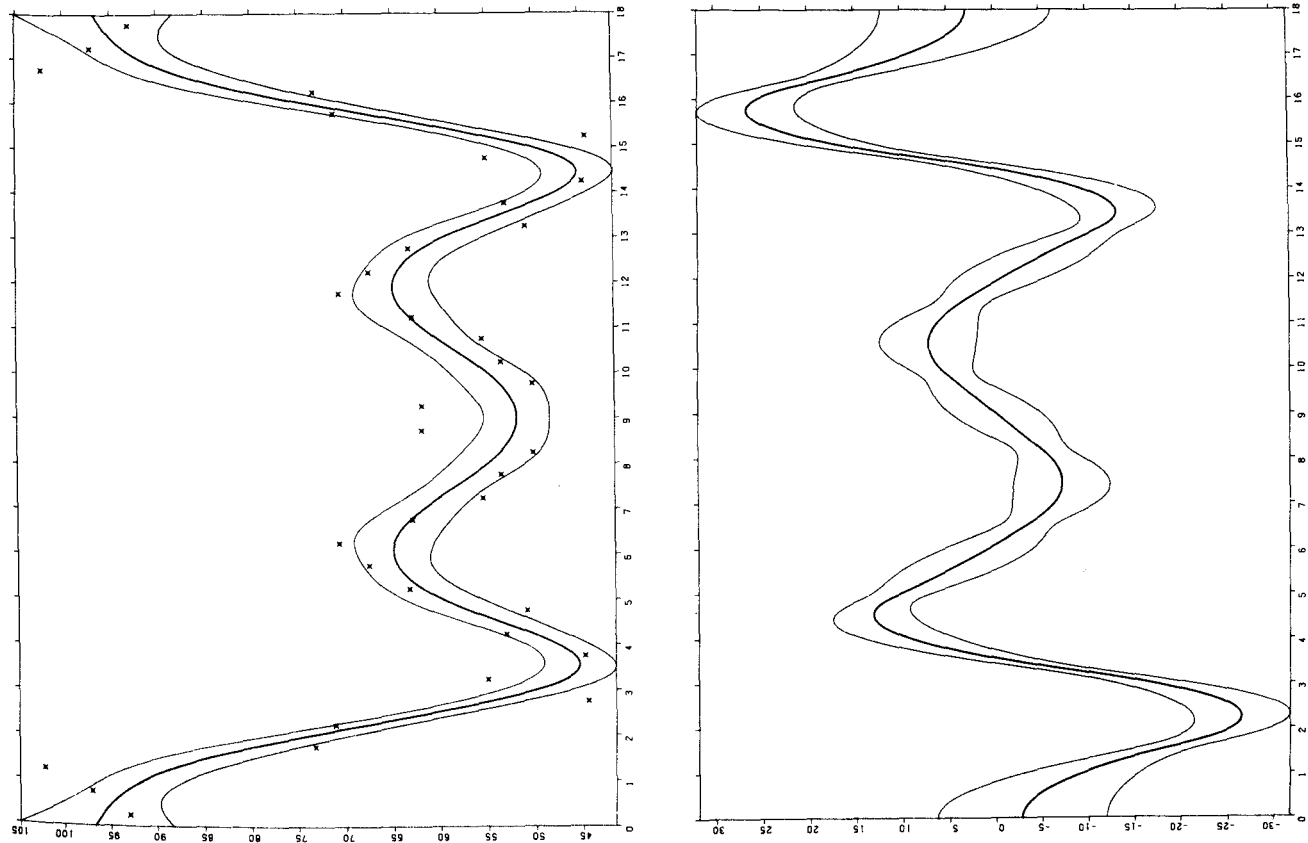


Fig. 5

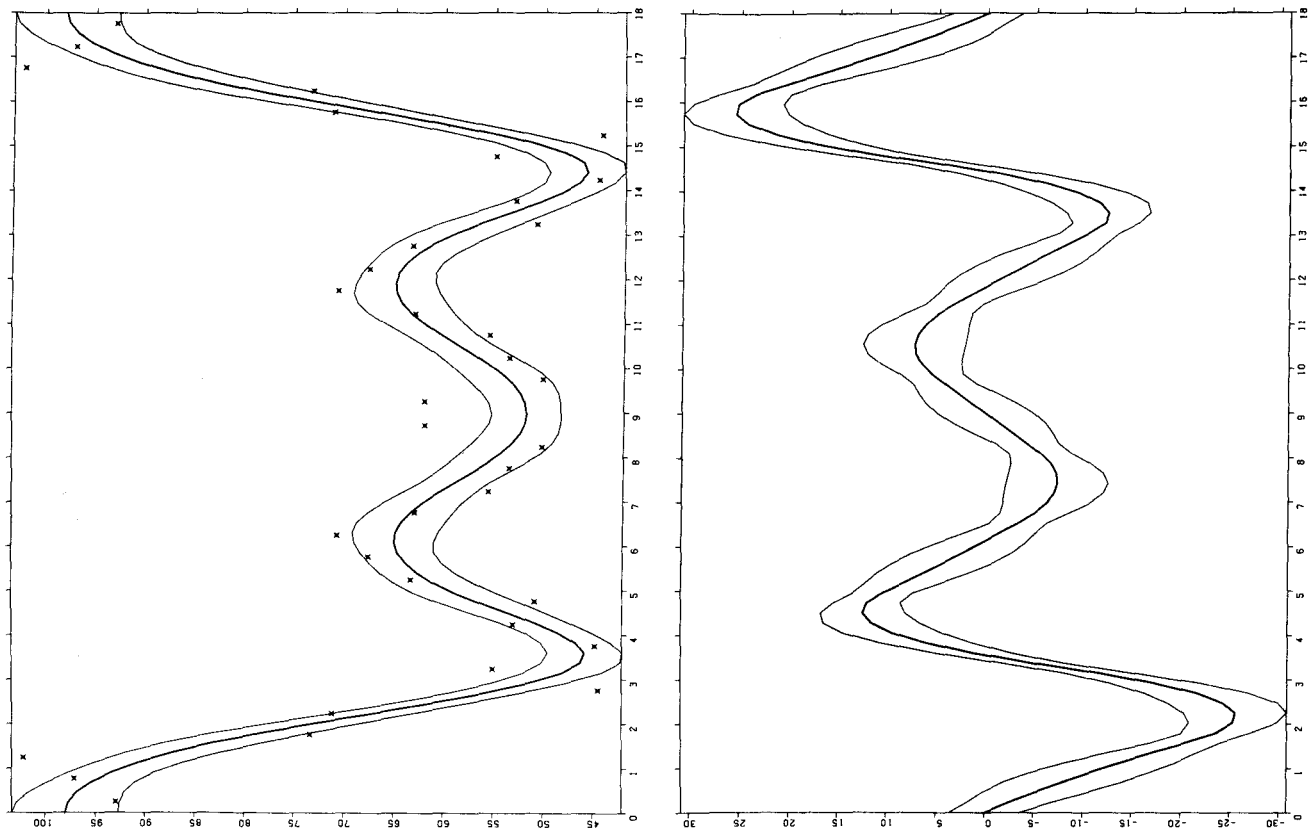


Fig. 6

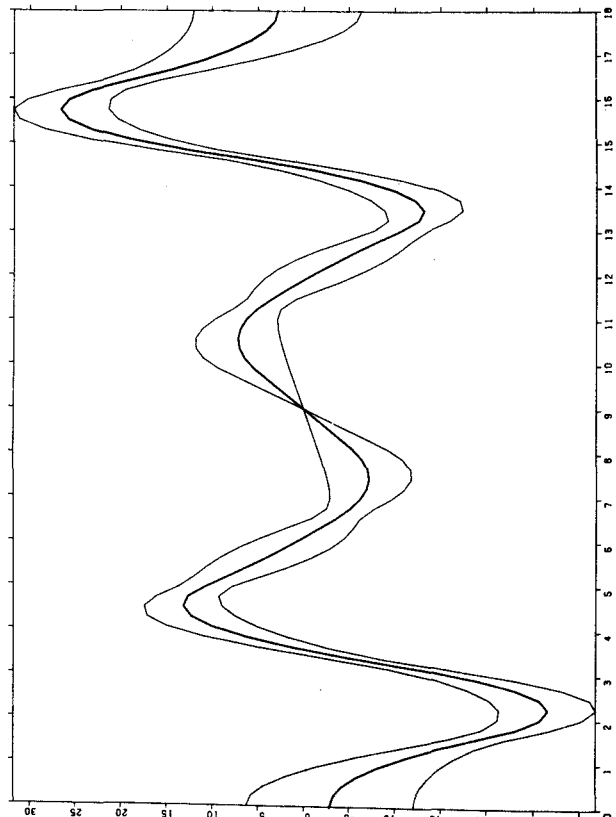
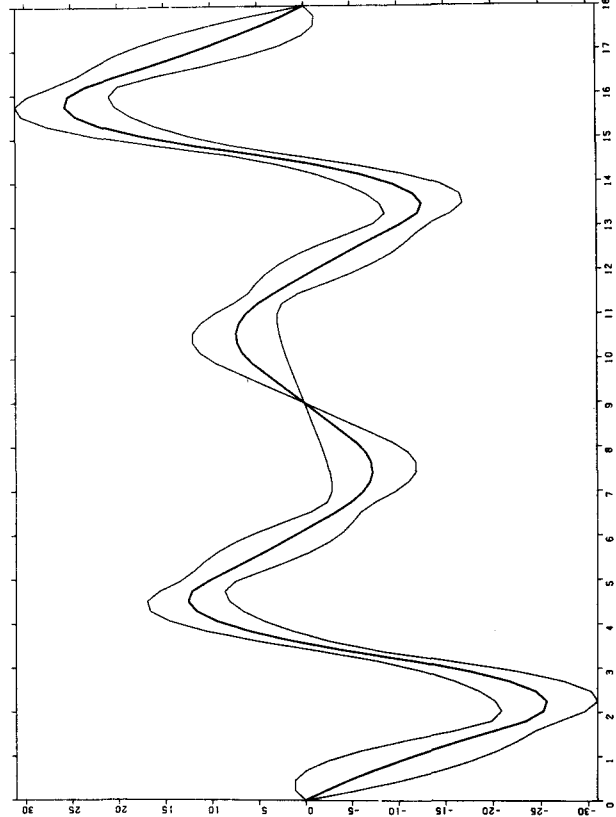
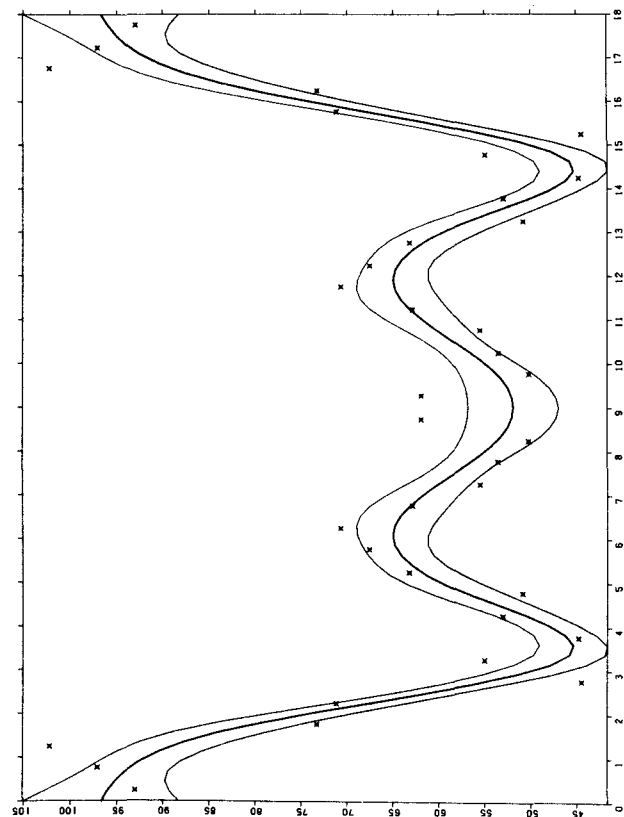
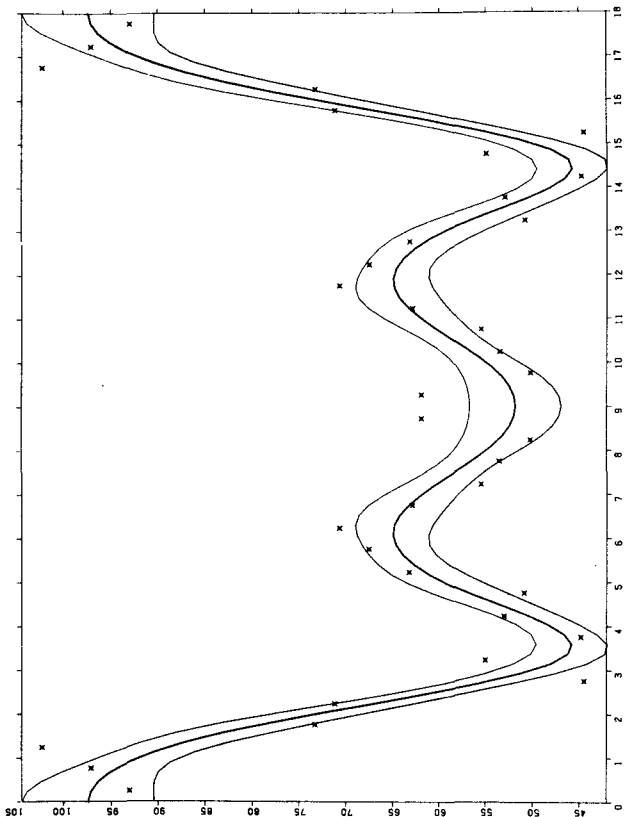


Fig. 8

Fig. 7

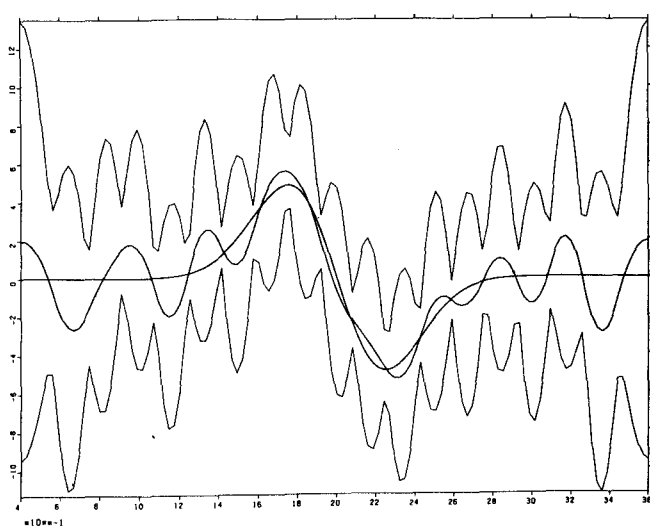
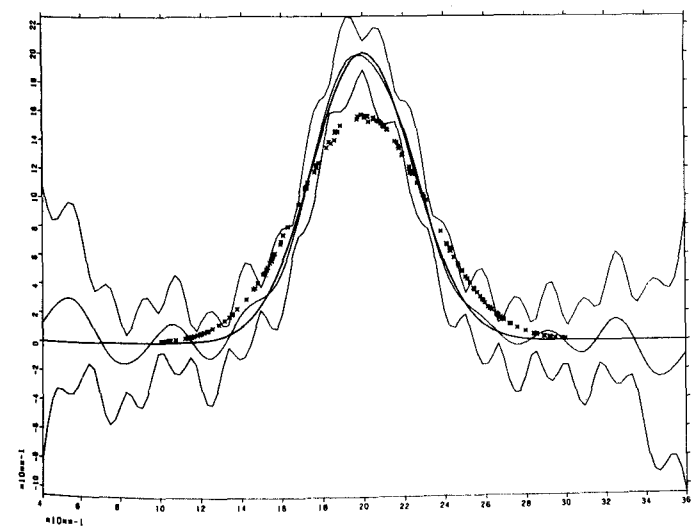
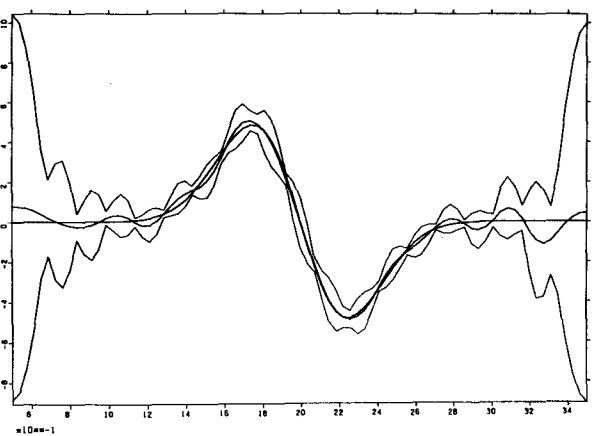
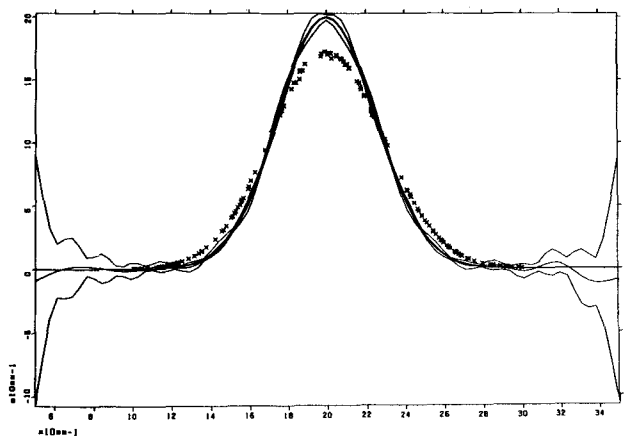
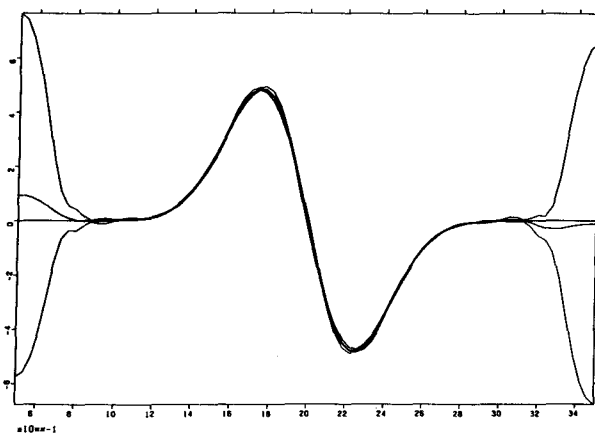
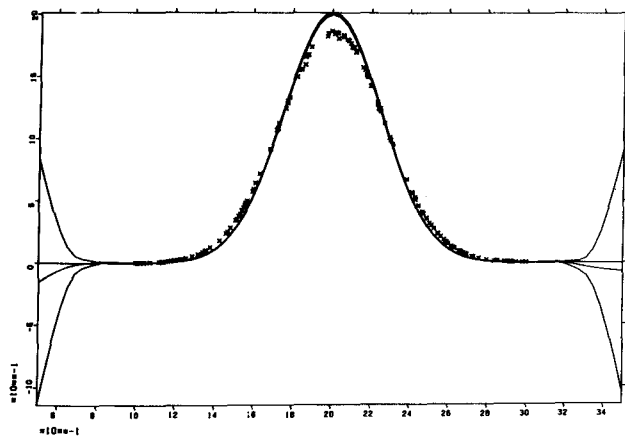


Fig. 9

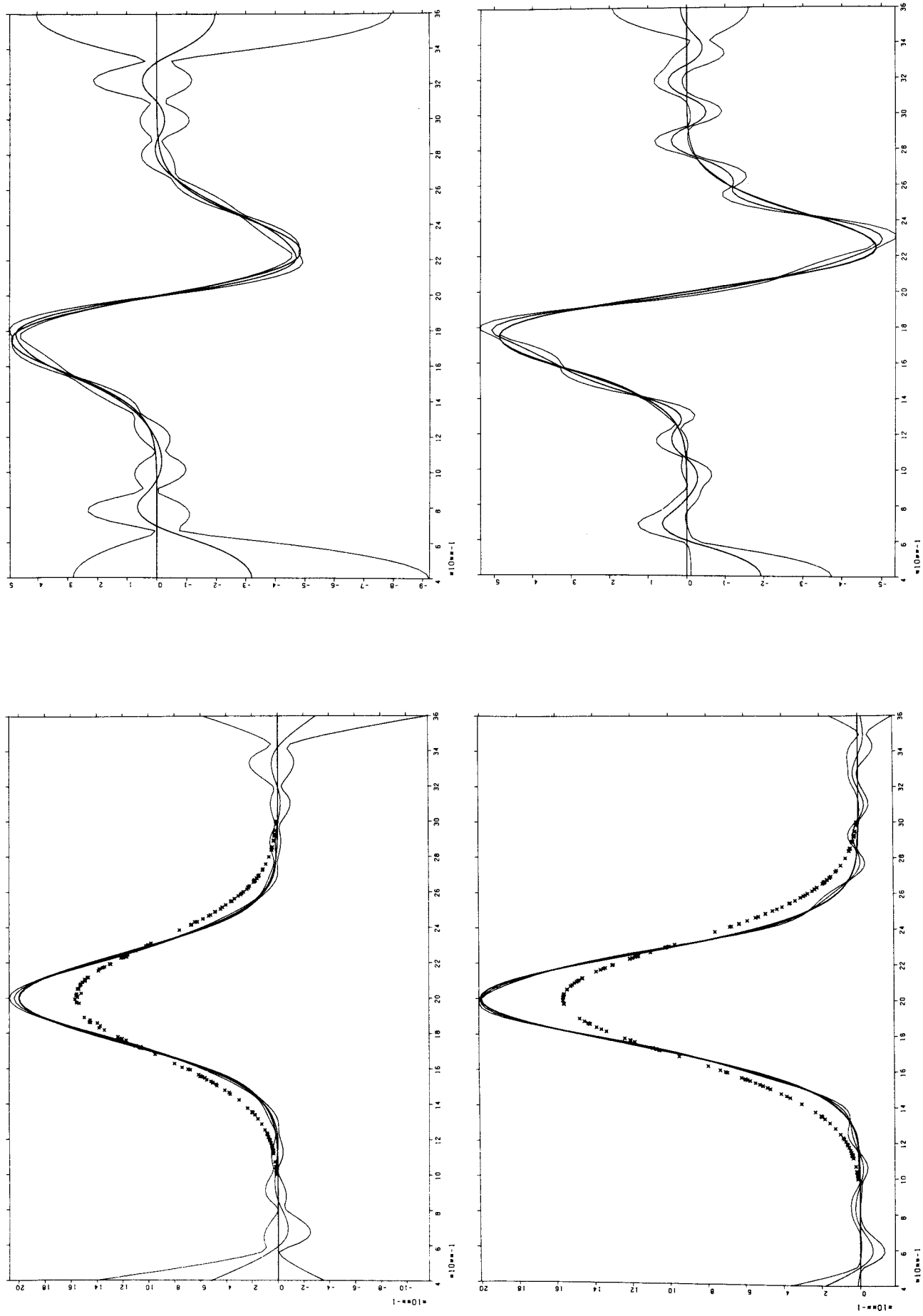


Fig. 10

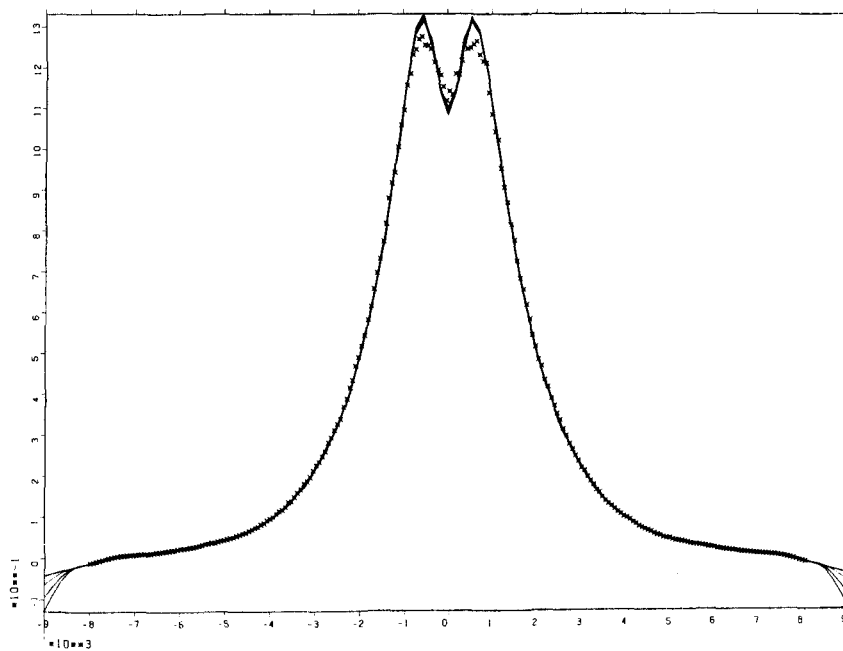
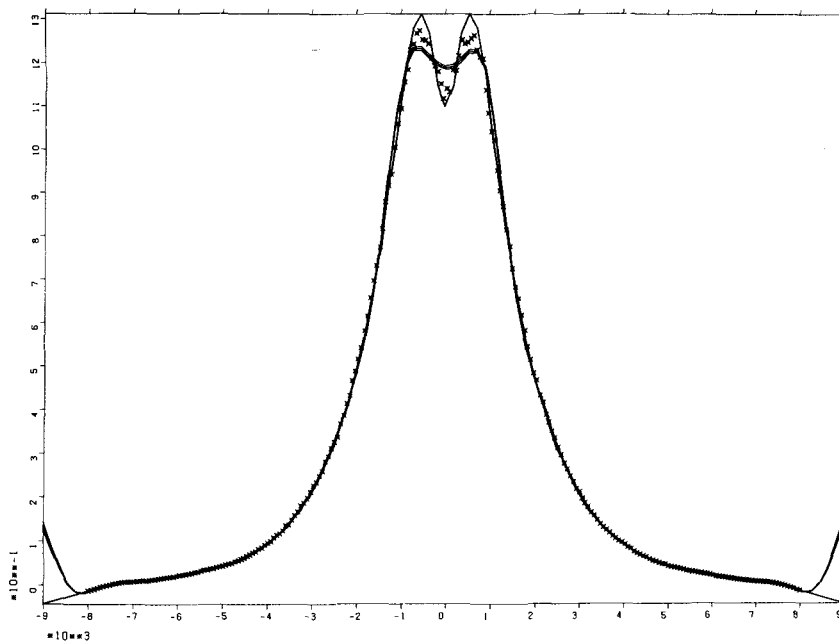
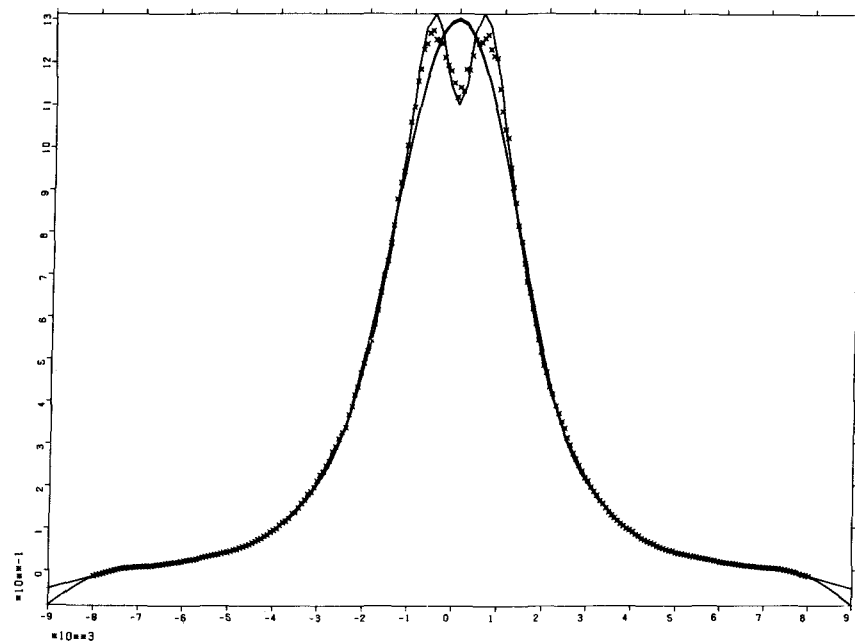


Fig. 11

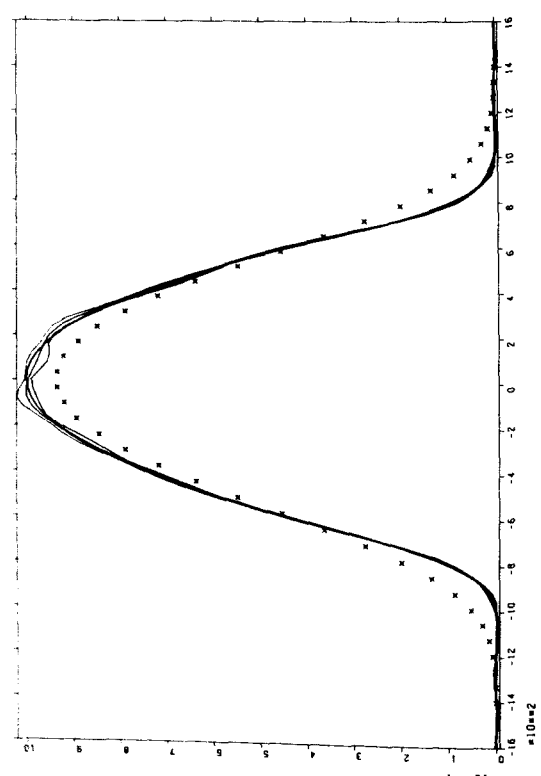
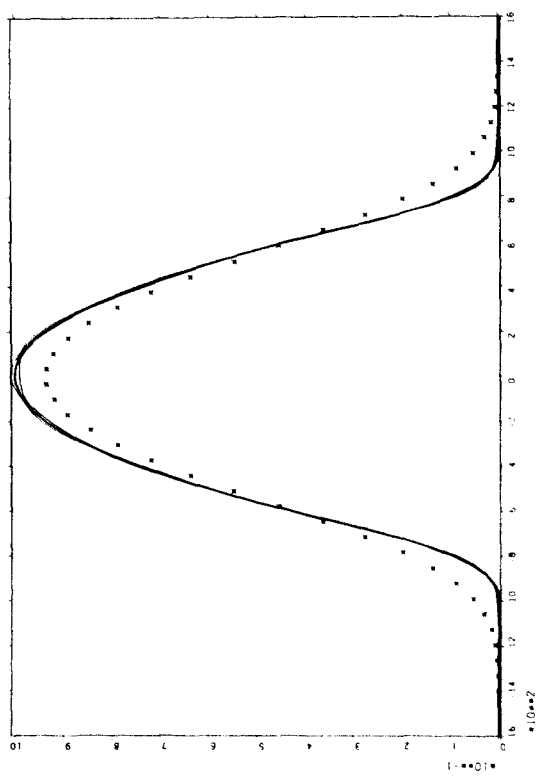
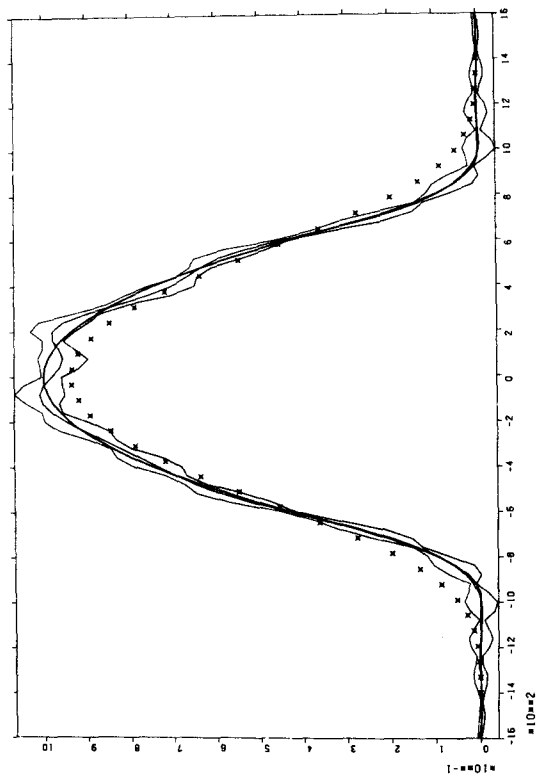
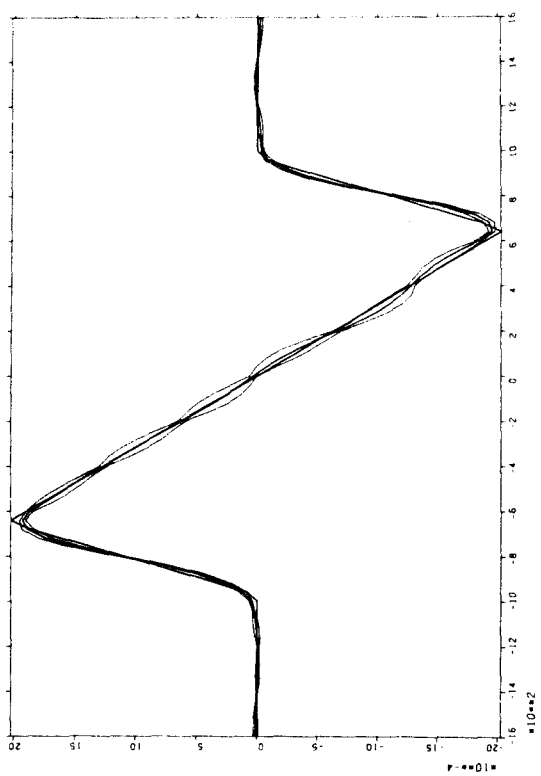


Fig. 12

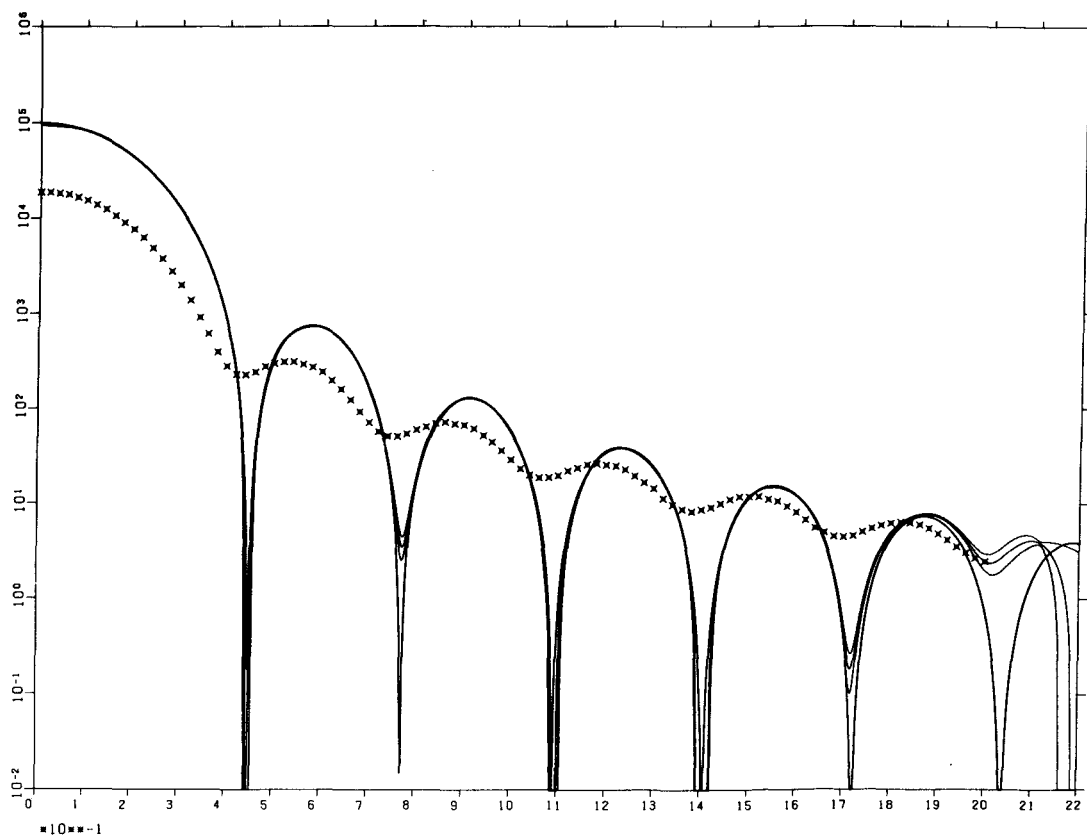


Fig. 13

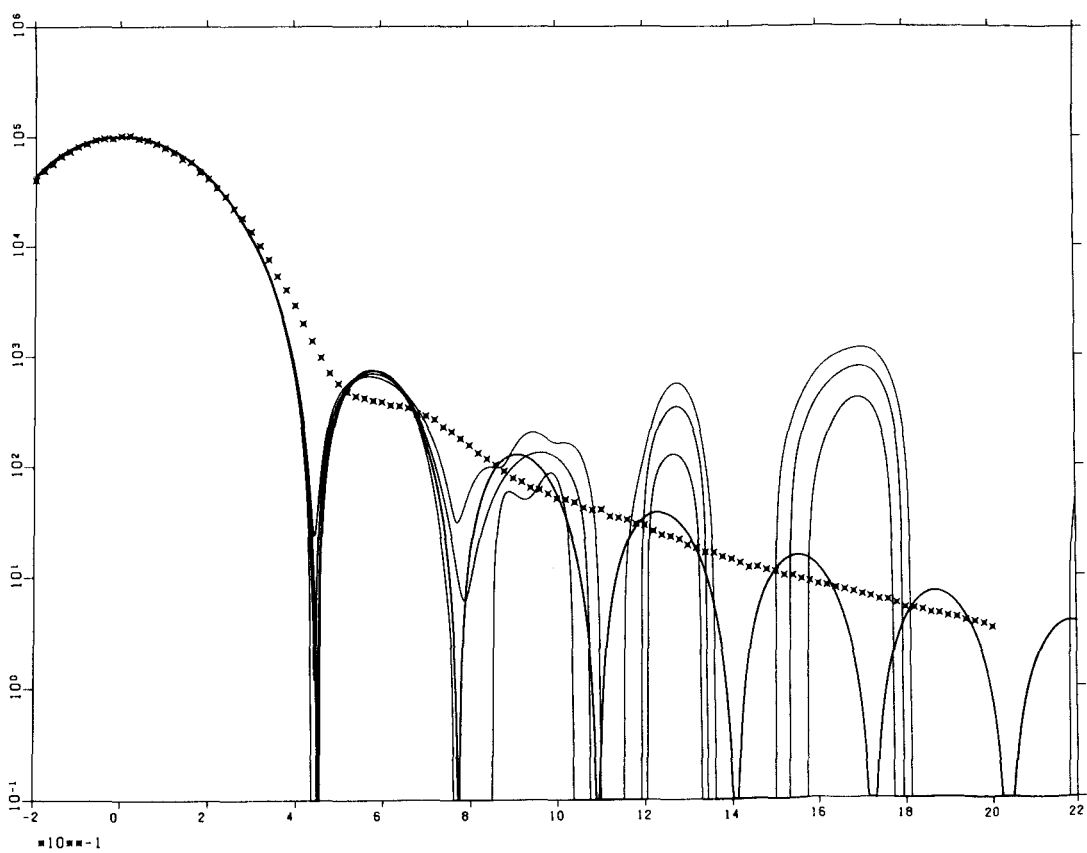


Fig. 14