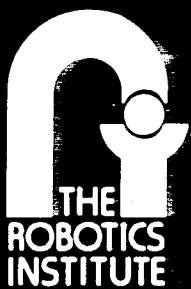
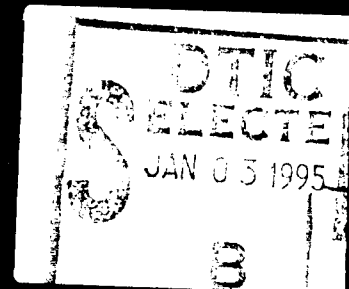


Truncated Gaussians as Tolerance Sets

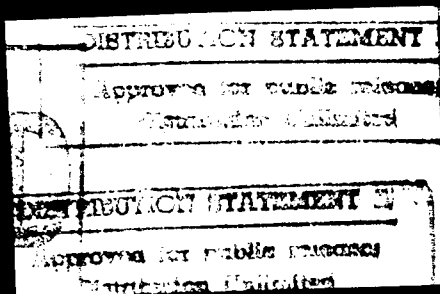
Fabio Cozman Eric Krotkov
CMU-RI-TR 94-35



Carnegie Mellon University

The Robotics Institute

Technical Report



19941228 138

Truncated Gaussians as Tolerance Sets

Fabio Cozman Eric Krotkov
CMU-RI-TR 94-35

The Robotics Institute
Carnegie Mellon University
Pittsburgh, PA 15213

September 28, 1994

DTIC QUALITY INSPECTED 2

©1994 Carnegie Mellon University

This research is supported in part by NASA under Grant NAGW-1175. Fabio Cozman is supported under a scholarship from CNPq, Brazil.



The Robotics Institute
Carnegie Mellon University
Pittsburgh, Pennsylvania 15213-3890

15 December 1994

Defense Technical Information Center
Cameron Station
Alexandria, VA 22314

RE: Report No. CMU-RI-TR-94-35

Permission is granted to the Defense Technical Information Center and the National Technical Information Service to reproduce and sell the following report, which contains information general in nature:

Fabio Cozman and Eric Krotkov
Truncated Gaussians as Tolerance Sets

Yours truly,

A handwritten signature in cursive script that reads "Marcella L. Zaragoza".

Marcella L. Zaragoza
Graduate Program Coordinator

enc.: 12 copies of report

Abstract

This work focuses on the use of truncated Gaussian distributions as models for bounded data - measurements that are constrained to appear between fixed limits. We prove that the truncated Gaussian can be viewed as a maximum entropy distribution for truncated bounded data, when mean and covariance are given. We present the characteristic function for the truncated Gaussian; from this, we derive algorithms for calculation of mean, variance, summation, application of Bayes rule and filtering with truncated Gaussians. As an example of the power of our methods, we describe a derivation of the disparity constraint (used in computer vision) from our models. Our approach complements results in Statistics, but our proposal is not only to use the truncated Gaussian as a model for selected data; we propose to model measurements as fundamentally bounded in terms of truncated Gaussians.

Accession For	
NTIS GRA&I	☒
DTIC TAB	☐
Unannounced	☐
Justification	
By <i>per letter</i>	
Distribution/	
Availability Codes	
Dist	Avail and/or Special
<i>A-1</i>	

1 Introduction

This work presents a new class of statistical models that are well suited for several Robotics applications, such as object recognition or computer vision. Our approach deals with bounded data: measurements that are constrained to appear in a bounded region in the measurement space. Bounded measurements have been studied in Statistics in connection with selection mechanisms; here we propose a different approach, in which the data is considered to be bounded to begin with.

To date, few statistical models for bounded variables are available, none of them satisfactory. The most common approach is to use the Gaussian distribution and model bounds through an *ad hoc* selection mechanism [2, 1]. Another possibility is the uniform distribution [8], but this approach has computational problems: summation of uniform variables does not yield a uniform variable and application of Bayes rule is hard [11]. In short: even though bounds contain a lot of information, they have not received proper attention yet.

Our work uses a class of distributions in the truncated Gaussian family in order to model bounded data. We derive a complete set of *tractable* algorithms for these models: calculation of moments, approximation methods for Bayes rule and summation and noise filtering. Overall, our results make the truncated Gaussian family an operational tool, much more powerful than the uniform or the Gaussian distributions.

In order to illustrate the strength of our approach, we present a statistical derivation of the *disparity constraint* used in Computer Vision. So far no statistical analysis has been given for this constraint.

Our analysis complements results scattered in the literature of Statistics, Information Theory and Control Theory. We contribute to Robotics by indicating a proper way to model bounded measurements and deriving tractable algorithms to handle them. Besides contributing to Robotics, our algorithms demonstrate that Robotics has much to contribute to Statistics itself.

2 The Truncated Gaussian Family

Our basic model is the elliptically truncated Gaussian family. A distribution in this family is proportional to a Gaussian inside an ellipsoid and is zero outside the ellipsoid. The truncated Gaussian model has been proposed in a variety of contexts in Statistics [9] as models of selection mechanisms. In Robotics, the first explicit mention of the possibility of using the truncated Gaussian appears to be by Erdmann [5].

A truncated Gaussian distribution for a n -dimensional random vector x is referred to as

$N_{\nu,M,k}(\mu, P)$; its mathematical expression is:

$$N_{\nu,M,k}(\mu, P) = \frac{(q(\mu, P, \nu, M, k))^{-1}}{|(2\pi)^n \det(P)|^{\frac{1}{2}}} \exp\left(-\frac{1}{2}(x - \mu)^T P^{-1}(x - \mu)\right) I_{\{(x - \nu)^T M^{-1}(x - \nu) \leq k\}}(x), \quad (1)$$

where $I(\cdot)$ is the indicator function and $q(\mu, P, \nu, M, k)$ is a normalizing constant. The set $\{x : (x - \nu)^T M^{-1}(x - \nu) \leq k\}$ defines an ellipsoid in n -dimensional space. Call k the *radius* of the distribution.

As a special case, note that if $\mu = \nu$ and $M = P$, then the normalizing constant depends only on k :

$$q(\mu, P, \mu, P, k) = Pr(\chi_n^2 \leq k) = \int_0^k (2^{n/2} \Gamma(n/2))^{-1} x^{n/2-1} e^{-x/2} dx.$$

In this case, $q(\mu, P, \mu, P, k)$ is the value (at k) of the distribution function of a chi-square variable with n degrees of freedom. Due to the importance of this sub-family for modeling purposes, we call it the *radially truncated Gaussian* family.

2.1 Genesis of a Truncated Gaussian

There are other situations where a truncated model is appropriate because data is purposefully truncated. Consider an n -dimensional source of (unbounded) Gaussian noise and an ellipsoid in this n -dimensional space. If any measurement outside the ellipsoid is discarded, then the resulting data obeys a truncated Gaussian. Examples of these procedures are the algorithms of Cox [2] and Bar-Shalom/Fortmann [1]. These algorithms use the Gaussian distribution associated with *ad hoc* selection mechanisms; none of these algorithms uses appropriate models for bounded data. This approach to bounded data is employed in Statistics in order to model selection mechanisms [9]. Results derived in this work can be understood as new tools for this type of analysis.

There is a different way of looking at bounded data. We can use bounded models in order to model data that is *fundamentally bounded*, not bounded as the result of a selection. We elaborate on that. Suppose we know a random vector x has mean μ , covariance matrix Q , and $Pr(x) = 0$ for all x outside $\{x : (x - \mu)^T Q^{-1}(x - \mu) \leq k\}$. In other words, possible values of x concentrate around the mean in a symmetric fashion, up to the distance k in the metric induced by Q . Under these conditions, we have (proof in Appendix A):

Theorem 1 *Given a expected value is μ , a covariance matrix Q and the fact that a distribution is zero outside the set $\{x : (x - \mu)^T Q^{-1}(x - \mu) \leq k\}$, a maximum entropy distribution that obeys these conditions is a truncated Gaussian $N_{\mu,cQ,k'}(\mu, cQ)$, where $c = Pr(\chi_n^2 \leq k)(Pr(\chi_{n+2}^2 \leq k))^{-1}$ and $k' = kPr(\chi_{n+2}^2 \leq k)(Pr(\chi_n^2 \leq k))^{-1}$.*

This theorem strengthens the parallel between truncated and unbounded Gaussians.

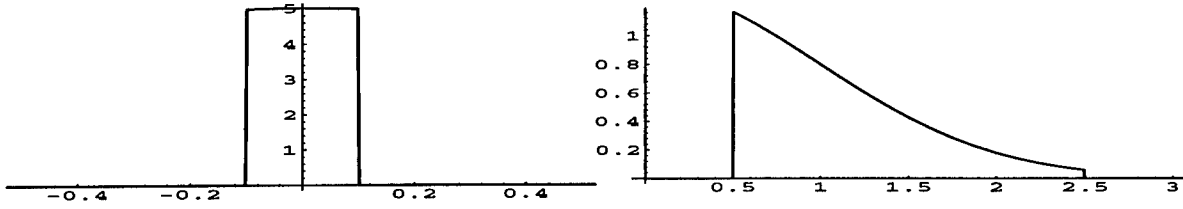


Figure 1: Distributions in the truncated Gaussian family: $N_{0,1,0.01}(0,1)$ (left) and $N_{1.5,1,1}(0,1)$ (right).

The truncated Gaussian family is even more powerful than the unbounded Gaussian. The truncated Gaussian family can represent highly skewed distributions and distributions that are nearly constant. Figure 1 illustrates this claim.

2.2 Numerical Evaluation of $q(\mu, P, \nu, M, k)$

The central problem in the characterization of a truncated Gaussian is the determination of $q(\mu, P, \nu, M, k)$. We introduce an important linear transformation that will be used throughout the paper.

It is always possible to transform the original expression of a truncated Gaussian into the following:

$$N_{\omega, D, k}(0, I) = \frac{(q(\omega, D, k))^{-1}}{(2\pi)^{\frac{1}{2}}} \exp\left(-\frac{1}{2}z^T z\right) I_{\{z: (z-\omega)^T D^{-1}(z-\omega) \leq k\}}(z),$$

where D is a diagonal matrix. The result is obtained through a linear transformation, named *double diagonalization*:

$$z = \Phi^T \sqrt{\Lambda}^{-1} V^T (x - \mu)$$

where: Λ is the eigenvalue matrix of P , V is the corresponding eigenvector matrix of P and Φ is the eigenvector matrix of $(\sqrt{\Lambda}^{-1} V^T) M (V \sqrt{\Lambda}^{-1})$. The transformation is always possible since P and M are positive definite. The vector ω is $\left[\Phi^T \sqrt{\Lambda}^{-1} V^T (\mu - \nu)\right]$.

We work with the transformed variable z , since $q(\omega, D, k) = q(\mu, P, \nu, M, k)$. Then (d_i the inverse of the i th element in the diagonal of D):

$$q(\omega, D, k) = Pr\left(\sum_{i=1}^n d_i (z_i - \omega_i)^2 \leq k\right) = \int_{(\sum_{i=1}^n d_i (z_i - \omega_i)^2 \leq k)} \exp\left(-\frac{1}{2}z^T z\right) dz.$$

Kotz, Johnson and Boyd [9] give a numerical method for the evaluation of this integral based

on its Laguerre expansion. We have:

$$q(\omega, D, k) = Pr(\chi_n^2 \leq k/\beta) + \left(\frac{k}{2\beta}\right)^{n/2} e^{-\frac{k}{2\beta}} \sum_{j=1}^{\infty} c_j L_{j-1}(k/(2\beta)) \frac{(j-1)!}{\Gamma(n/2 + j)}$$

where:

$$\begin{aligned} L_r(y) &= \sum_{i=0}^r \frac{(-y)^i}{i!(r-i)!} \frac{\Gamma(r+n/2+1)}{\Gamma(i+n/2+1)} \\ c_r &= (2r)^{-1} \sum_{j=0}^{r-1} s_{r-j} c_j, \quad c_0 = 1 \\ s_r &= (-r/\beta) \sum_{j=1}^n \omega_j^2 d_j (1 - d_j/\beta)^{r-1} + \sum_{j=1}^n (1 - d_j/\beta)^r. \end{aligned}$$

Convergence is uniform for $\beta > \frac{\max_j(d_j)}{2}$. In general, large values of β yield slow convergence. If we truncate the evaluation of the series at $j = N$, the truncation error is always smaller than:

$$2^{-N} \exp\left(k/(2\beta) + (n-1) \log 2 + 4\omega^T \omega\right).$$

As a one-dimensional example, consider $N_{1.5,1,1}(0,1)$ (shown in figure 1). Numerical integration with 10 significant digits yields $q(0,1,1.5,1,1) = 0.3023278734$. Using $\beta = 1.1$, $N = 9$, we get the same answer.

2.3 Moment Generating Function of a Truncated Gaussian

The moment generating function of any distribution is a fundamental tool in statistical analysis. We give here an expression for the moment generating function of a truncated Gaussian in the most general form (proof in Appendix A):

Theorem 2 *The moment generating function of a truncated Gaussian $N_{\omega,D,k}(0,I)$ is:*

$$\phi(t) = \frac{q(\omega - t, D, k)}{q(\omega, D, k)} \exp\left(\frac{t^T t}{2}\right). \quad (2)$$

For a particular value of t , we can evaluate $\phi(t)$ using Kotz, Johnson and Boyd recursions for $q(\omega - t, D, k)$ and then plugging the result in $(q(\omega, D, k))^{-1} q(\omega - t, D, k) \exp(t^T t/2)$.

2.4 Mean Vector and Covariance Matrix of a Truncated Gaussian

The mean and covariance can be obtained by successive differentiations of $\phi(t)$ [10]. Since Kotz, Johnson and Boyd recursions for $\phi(t)$ are uniformly convergent we can differentiate these expressions term by term. We use $\bar{\mu}$ for the mean and $\bar{P}$ for the covariance of a truncated Gaussian.

Direct differentiation of 2 yields:

Theorem 3 *We have:*

$$\begin{aligned}\bar{\mu} &= k_o \sum_{j=1}^{\infty} k_j g_j \\ \bar{P} &= I - \bar{\mu}\bar{\mu}^T + k_o \sum_{j=1}^{\infty} k_j G_j,\end{aligned}$$

where (k_r are scalars, g_r and h_r are vectors, G_r and H_r are matrices and I is an identity matrix):

$$\begin{aligned}k_j &= \frac{(j-1)!L_{j-1}(k/(2\beta))}{\Gamma(n/2+j)} \quad , \quad k_o = \left(\frac{k}{2\beta}\right)^{n/2} \frac{e^{-\frac{k}{2\beta}}}{q(\omega, D, k)} \\ g_r &= (2r)^{-1} \sum_{j=0}^{r-1} (h_{r-j}c_j + s_{r-j}g_j) \quad , \quad g_0 = 0 \\ h_r &= \frac{2r}{\beta} \text{diag} [d_1(1-d_1/\beta)^{r-1}, \dots, d_n(1-d_n/\beta)^{r-1}]\omega \\ G_r &= (2r)^{-1} \sum_{j=0}^{r-1} (c_j H_{r-j} + g_j h_{r-j}^T + h_{r-j} g_j^T + s_{r-j} G_j) \quad , \quad G_0 = 0 \\ H_r &= -\frac{2r}{\beta} \text{diag} [d_1(1-d_1/\beta)^{r-1}, \dots, d_n(1-d_n/\beta)^{r-1}]\end{aligned}$$

For a distribution $N_{\mu,P,k}(\mu, P)$ in the radially truncated Gaussian family, the mean is μ and the covariance matrix is $Pr(\chi_{n+2}^2 \leq k)(Pr(\chi_n^2 \leq k))^{-1}P$ [9].

2.5 Linear Transformation and Summation of Truncated Gaussians

A non-singular linear transformation applied to random vector with truncated Gaussian distribution produces another truncated Gaussian random vector (proofs of theorems in this section are in Appendix A).

Theorem 4 If $x \sim N_{\nu, M, k}(\mu, P)$ and $y = Ax$ (where A is any non-singular square matrix) then $y \sim N_{A\nu, AMA^T, k}(A\mu, APA^T)$.

It can be seen, through direct manipulation of truncated Gaussians even in one dimension, that the sum of truncated Gaussian random variables is not a truncated Gaussian. We can derive some results for particular cases. Suppose $z = x + y$ where $x \sim N_{\nu_x, M_x, k_x}(\mu_x, P_x)$ and $y \sim N_{\nu_y, M_y, k_y}(\mu_y, P_y)$, x and y independent. Under these conditions:

Theorem 5 The distribution of z has expected value $\bar{\mu}_z = \bar{\mu}_x + \bar{\mu}_y$ and covariance $\bar{P}_z = \bar{P}_x + \bar{P}_y$; the distribution of z is positive only inside the ellipsoid defined by $\{z : (z - \nu_z)^T M_z^{-1} (z - \nu_z) \leq 1\}$ where

$$\nu_z = \nu_x + \nu_y \quad M_z = k_x M_x + k_y M_y$$

A reasonable approximation to the distribution of z is $N_{\nu_z, M_z, 1}(\bar{\mu}_z, \bar{P}_z)$.

A special but important case is represented by $\nu_x = \mu_x$, $M_x = P_x$, $\nu_y = \mu_y$, $M_y = P_y$, $k_x = k_y$. In this special case we can make statements about the distance between the approximation and the correct distribution. Call G_z the correct distribution for z ; then:

Theorem 6 For $\mu_z = \mu_x + \mu_y$, $M_z = M_x + M_y$, $k_z = k_x = k_y$, $\sup_z |G_z - N_{\nu_z, M_z, k_z}(\nu_z, M_z)|$ is $\mathcal{O}(\exp(-k/2))$.

So for this case, the indicated approximation is fairly good.

Consider now the most general case: $z = x + y$, where x and y are arbitrary truncated Gaussians. We shall indicate approximation strategies that work well under specific conditions. We consider the vector u , a linear transformation of z so that defining ellipsoid is $\{u : u^T u \leq 1\}$. Call μ_u the mean of u and P_u the covariance matrix of u . The distribution of u is obtained by convolution, so it is unimodal and closer to a bell-shaped function than the original distributions. We may expect that the distribution of u to be closer to a Gaussian than the distribution of x or y , invoking the Central Limit Theorem.

Consider the approximation $p_1(u) = N_{0, I, 1}(\mu_u, P_u)$. If all eigenvalues of P_u are much smaller than 1, then this approximation is reasonable because the mean will be approximately μ_u and the variance will be approximately P_u . The distribution will always have a maximum inside the defining ellipsoid, so it will always yield a smooth approximation to $p_U(u)$. An example will illustrate this. Consider distribution $N_{1.5, 1, 1}(0, 1)$, depicted in figure 1. Mean is 1.107 and variance is 0.21288. If we add two random variables with this same distribution, we obtain (numerically) the distribution shown in figure 2.a. Mean is 2.213 and variance is 0.4257. The approximation $N_{1.5, 1}(2.213, 0.4257)$ is shown in figure 2.b. Both distributions are shown in figure 2.c. Agreement

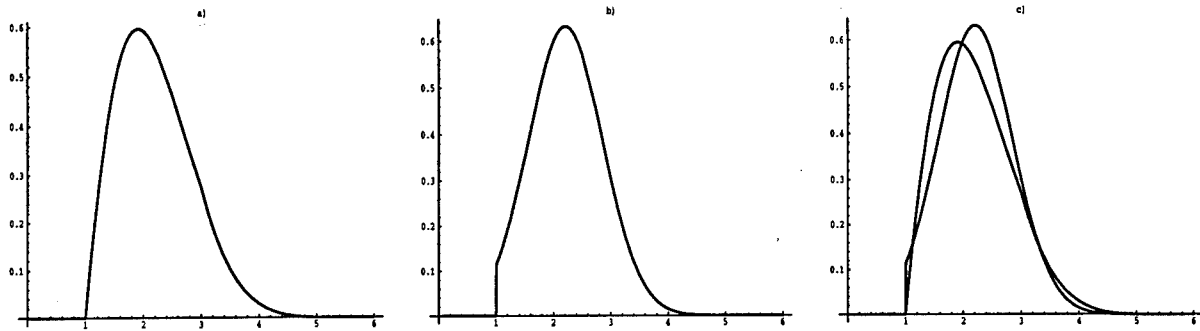


Figure 2: Example: Approximation to Summation

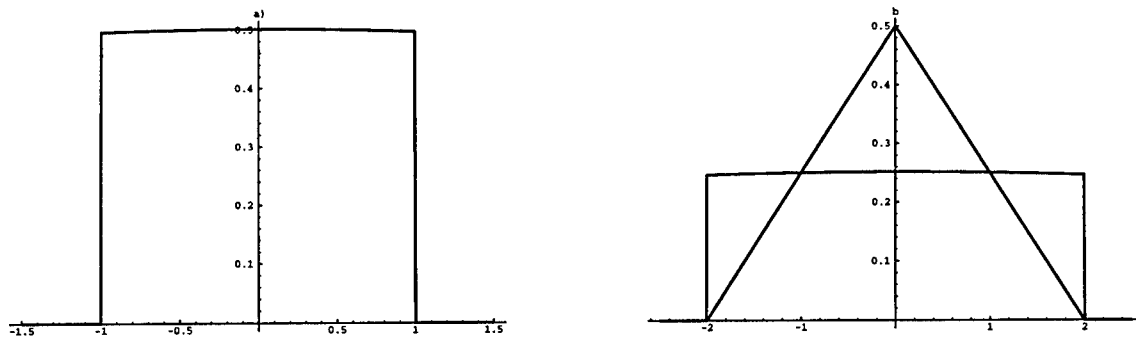


Figure 3: Example: Failure of Approximation to Summation

is remarkable even in this highly skewed distribution; the strategy will work better if the means of the original summands are closer to zero.

We may take the approximation as valid for large variances, but there is a limit to such process. We cannot use this strategy for distributions with arbitrarily large variances. Again, we use an example to illustrate this situation. Consider the distribution $N_{0,1,1}(0, 50)$, depicted at figure 3.a. Consider the summation of two random variables with this same distribution. Figure 3.b shows both the correct distribution and the approximation $N_{0,1,4}(0, 100)$. The disagreement is evident. So we need a specific approximation for the case of large variances.

A different strategy would be to approximate the distribution of u by a truncated Gaussian that matches the mean and variance of u . This approach is explored in Appendix B; approximation derived there are valid under restrictive conditions on the second moment of the distribution of u . Derivation of good approximations for all cases is still an open problem.

3 Inferences with the Truncated Gaussian

Inferences about a random variable are obtained by application of Bayes rule associated with a decision rule. We analyze two decision rules: maximum a posteriori estimate and minimum square

loss estimate.

Since the truncated Gaussian is not closed under multiplication, we should not expect to be able to apply Bayes rule and obtain a truncated Gaussian distribution. The next example reveals the problems that may arise here.

Example 1 Suppose x and ω are related by $z = x + \omega$, $p(x)$ is $N_{0,(1+\epsilon)^2,3}(0, (1+\epsilon)^2)$ and $p(z|x)$ is $N_{x,(1+\epsilon)^2,3}(x, (1+\epsilon)^2)$. Suppose the observation comes out to be $z = 2\sqrt{3}$. Application of Bayes rule for $p(x|z)$ yields a distribution concentrated between $[\sqrt{3} - \epsilon, \sqrt{3} + \epsilon]$, tending to a spike as ϵ goes to zero. $\square$

Fortunately, we can find a good approximation methodology for a situation relevant to practical applications. We consider a linear set-up:

$$z = Hx + \omega_z \quad (3)$$

where $x \in \mathbb{R}^n$ and $z \in \mathbb{R}^m$, A and H are matrices of appropriate dimensions, x is distributed as a truncated Gaussian $N_{\nu_x, M_x, k_x}(\mu_x, P_x)$; ω_z is distributed as a radially truncated Gaussian with zero mean $N_{0, P_z, k_z}(0, P_z)$. As a result, $z \sim N_{Hx, P_z, k_z}(Hx, P_z)$.

The central question is: *given that we observe the value of z , what can we say about x ?*

The posterior is not a truncated Gaussian because the region where the distribution is positive is *not* an ellipsoid (it is the intersection of two ellipsoids). The distribution still is proportional to a Gaussian in the points where it is positive, as we show now. Consider two distributions:

$$q_1 \exp\left(-\frac{1}{2}(x - \mu_x)^T P_x^{-1}(x - \mu_x)\right) I_A(x) \quad q_2 \exp\left(-\frac{1}{2}(z - Hx)^T P_z^{-1}(z - Hx)\right) I_B(x),$$

where A and B are arbitrary convex sets, not necessarily ellipsoids. If we multiply these distributions together, the result is positive only on the set $A \cap B$; on this set the result is proportional to:

$$\mathcal{W} = \exp\left(-\frac{1}{2}(x - \mu_x)^T P_x^{-1}(x - \mu_x) + (z - Hx)^T P_z^{-1}(z - Hx)\right).$$

By using the Matrix Inversion Lemma [1], we get the following standard result of linear systems theory:

$$\mathcal{W} \propto \exp\left(-\frac{1}{2}(x - \mu')^T Q^{-1}(x - \mu')\right)$$

where $Q = P_x - P_x H^T (H P_x H^T + P_z)^{-1} H P_x$ and $\mu' = Q(P_x^{-1} \mu_x + P_z^{-1} H^T z)$. This shows that the shape of the distribution is proportional to a Gaussian. So the final distribution will be:

$$\frac{c}{|2\pi \det(Q)|^{n/2}} \exp\left(-\frac{1}{2}(x - \mu')^T Q^{-1}(x - \mu')\right) I_{A \cap B}(x), \quad (4)$$

where c is a normalizing constant.

So we can easily recover the shape of the distribution; the normalizing constant c still is unknown and depends on $A \cap B$. The main difficulty is to find a good and tractable approximation to this region. This is presented in the next section.

3.1 Intersecting Ellipsoids

In order to approximate the posterior, an ellipsoidal approximation can be built so that the approximate posterior is always a truncated Gaussian. This method was originally proposed by Fogel and Huang [7] in the context of tolerance sets for ARMA processes. Since the algorithm approximates only the geometric intersection of ellipsoids, the values of μ and P do not matter. We indicate truncated ellipsoids by $N(\nu, M, k)$ in this section.

3.1.1 Preliminary Approximation

We make some preliminary approximations:

- Consider expression (3). Suppose $\omega_z \sim N(0, W, 1/a)$ and $x \sim N(\mu, P, b)$.
- Now take a linear transformation A such that $AWA^T = I$ (always possible since W must be positive definite). So we have $Az = AHx + A\omega_z$. Now consider $B = AH$, $y = Az$ and $\omega_y = A\omega_z$; we have:

$$y = Bx + \omega_y$$

and $\omega_y \sim N(0, I, 1/a)$.

- We know that $a\omega_y^T \omega_y \leq 1$; instead of trying to satisfy this inequality, we will satisfy the set of inequalities:

$$a\omega_{yi}^2 \leq 1 \text{ for } i = 1, \dots, m.$$

The meaning of this approximation is the following: the original constraint defines a spheroid in $\mathbb{R}^m$; the new constraint defines the m -dimensional cube that circumscribes the spheroid.

- By rearranging terms we have:

$$a(y_i - B_i x)^2 \leq 1 \text{ for } i = 1, \dots, m. \quad (5)$$

where y_i is the i th element of y and B_i is the i th row of B .

Now the problem is how to approximate the following set:

$$\Theta_m = \left\{ x : \left[(x - \mu)^T P^{-1} (x - \mu) \leq b \right] \cap \left[\bigcap_1^m a(y_i - B_i x)^2 \leq 1 \right] \right\}.$$

3.1.2 Fogel/Huang Algorithm

Consider the following definitions:

$$\begin{aligned} P_o &= P \\ \theta_o &= \mu \\ b_o &= b \\ \Theta_o &= \{x : [(x - \theta_o)^T P_o^{-1} (x - \theta_o) \leq b_o]\} \\ S_i &= \{x : a(y_i - B_i x)^2 \leq 1\} \text{ for } i = 1, \dots, m. \end{aligned}$$

Each set S_i defines a strip (a region in $\mathbb{R}^n$ between two hyperplanes).

The Fogel/Huang strategy is to start with Θ_o and intersect this ellipsoid with one strip at a time, approximating the intermediate intersections by ellipsoids: We could generate Θ_m by the recursion:

$$\Theta_{i+1} = S_{i+1} \cap \Theta_i.$$

- We approximate Θ_{i+1} as:

$$\Theta_{i+1} = \left\{x : [(x - \mu)^T P_i^{-1} (x - \mu) + a q_{i+1} (y_{i+1} - B_{i+1} x)^2 \leq (b_i + q_{i+1})]\right\} \quad (6)$$

where q_{i+1} is a free parameter to be determined. Θ_{i+1} contains the intersection between Θ_i and the strip S_{i+1} .

- We can transform expression (6) into an explicit ellipsoid expression¹

$$\Theta_{i+1} = \{x : (x - \theta_{i+1})^T P_{i+1}^{-1} (x - \theta_{i+1}) \leq \gamma_{i+1}\} \quad (7)$$

where:

$$\begin{aligned} P_{i+1}^{-1} &= P_i^{-1} + a q_{i+1} B_i^T B_i & P_{i+1} &= P_i - a q_{i+1} \frac{P_i B_i^T B_i P_i}{1 + a q_{i+1} B_i P_i B_i^T} \\ \theta_{i+1} &= P_{i+1} (P_i^{-1} \theta_i + a q_{i+1} B_i^T y_{i+1}) & b_{i+1} &= b_i + q_{i+1} - a q_{i+1} \frac{(y_{i+1} - B_i \theta_i)^2}{1 + q_{i+1} B_i P_i B_i^T}. \end{aligned} \quad (8)$$

- We must determine q_{i+1} by some convention. Fogel and Huang prove that it is possible to choose q_{i+1} so that the square of the volume of the ellipsoid Θ_{i+1} is minimized. By doing so, we obtain (proof of this claim is in Appendix A):

$$q_{i+1} = \begin{cases} 0 & \text{if } \alpha_2^2 - 3\alpha_1\alpha_3 < 0 \text{ or } -\alpha_2 + \sqrt{\alpha_2^2 - 4\alpha_1\alpha_3} \leq 0 \\ \frac{-\alpha_2 + \sqrt{\alpha_2^2 - 4\alpha_1\alpha_3}}{2\alpha_1} & \text{otherwise.} \end{cases}$$

¹Using the following standard result: $(x - a)^T A (x - a) + (x - b)^T B (x - b) = (x - c)^T C (x - c)$ where $C = A + B$ and $c = C^{-1}(Aa + Bb)$, and the Matrix Inversion Lemma [1].

where:

$$\begin{aligned}\alpha_1 &= a^2(n-1)(B_{i+1}P_iB_{i+1}^T)^2 \\ \alpha_2 &= (B_{i+1}P_iB_{i+1}(2an-a+a^2b_iB_{i+1}P_iB_{i+1}+a^2(y_i-B_i\theta_i)^2 \\ \alpha_3 &= n(1-a(y_i-B_i\theta_i)^2)-aB_{i+1}P_iB_{i+1}.\end{aligned}\tag{9}$$

Summary of Fogel/Huang Algorithm The algorithm is a sequence of m steps. At each step:

1. Get q_{i+1} from the results of the previous step using expressions (9).
2. If $q_{i+1} = 0$, the measurement did not cause any change in the ellipsoids. Updating is the ellipsoid is not necessary.
3. Update Θ_{i+1} using expressions (7) and (8).
4. If $b_{i+1} < 0$, the likelihood and the prior are in conflict: the sets have empty intersection. This case is discussed below.

We take the defining ellipsoid for the posterior distribution to be $\{x : (x-\theta_m)^T P_m^{-1}(x-\theta_m) \leq b_m\}$. This is equivalent to the set $\{x : (x-\mu)^T P^{-1}(x-\mu) + \sum_{i=1}^m a q_i (y_i - B_i x)^2 \leq b + \sum_{i=1}^m q_i\}$. Our approximation is complete.

3.2 Updated Posterior

The posterior distribution is a combined result of expression (4) and Fogel/Huang algorithm; we use $N_{\theta_m, P_m, b_m}(\mu', Q)$ as the posterior. By looking at the recursions in expressions (7) and (8), we notice that, if $q_{i+1}a = 1$ for every i , then the final approximation is in the radially truncated Gaussian family.

Some examples clarify the use of our algorithm.

Example 2 Consider the situation:

$$z = \begin{bmatrix} 4 & 1 \\ 2 & 3 \end{bmatrix} x + \omega \quad \omega \sim N_{0, I, 1}(0, I) \quad x \sim N_{a, B, 1/4}(a, B)$$

where x is unknown and:

$$a = \begin{bmatrix} 1 \\ 1 \end{bmatrix} \quad B = \begin{bmatrix} 2 & 0 \\ 0 & 2 \end{bmatrix} \quad z = \begin{bmatrix} 2 \\ 3 \end{bmatrix}.$$

Figure 4.a shows the initial ellipsoids of the problem. The region of possible values of x is the intersection of ellipsoids. Figure 4.b shows the strips generated at expression (5); the intersection of these strips with the prior ellipsoid will define the posterior approximation.

The algorithm is applied and results are shown in figure 4.c. The largest hatched ellipse is the result of incorporating $z_1 = 2$; the smallest hatched ellipse is the final result after incorporation of $z_2 = 3$. The approximate posterior is $N_{c,D,e}(f, G)$ where $e = 0.257$ and:

$$c = \begin{bmatrix} 0.363 \\ 0.547 \end{bmatrix} \quad D = \begin{bmatrix} 0.4488 & -0.3272 \\ -0.3272 & 0.7555 \end{bmatrix} \quad f = \begin{bmatrix} 0.321 \\ 0.7418 \end{bmatrix} \quad G = \begin{bmatrix} 0.091 & -0.0867 \\ -0.0867 & 1.8857 \end{bmatrix}. \square$$

The approximations found by the algorithm are correct in that the intersections are always interior to the successive approximations. It should be noted that the original prior distribution for x is almost flat; it is possible, using our approach, to approximate situations where the strict Gaussian model would be brittle.

As a comparison, we analyze what would have happened if we had used an approximation strategy in the spirit of [1] or [2]. We would pretend x and ω to be distributed as unbounded Gaussians (with same mean and variance). Now we find a crucial problem: how to combine the fact that x has an elliptic region with radius $1/4$ and ω has an elliptic region with radius 1 ? By using Bayes rule in the unbounded Gaussians $N(0, I)$ and $N(a, B)$, we obtain an unbounded Gaussian $N(f, G)$. The hatched ellipsis in figure 4.d shows the results of using these values of mean and variance with radius $1/4$, $5/8$ and 1 . All ellipsis fail to cover the region where x is known to lie. Radius $1/4$ is extreme: almost all values of x inferred from this posterior are inconsistent with prior expectations!

Example 3 Consider the situation:

$$z = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} x + \omega \quad \omega \sim N_{0,I,9}(0, I) \quad x \sim N_{0,I,4}(0, I),$$

$z = [3, 7]^T$ and x is unknown.

Figure 5.a shows the ellipsis defined by prior expectations and measurements. It is clear that the measurement is inconsistent, indicating either the presence of outliers or mismodeling. Incorporation of z_1 yields the hatched ellipse in the figure; y_1 is consistent. When z_2 is incorporated, we obtain $b_2 = -11.49$, indicating failure. $\square$

Notice that detection of inconsistency is important and not trivial when using Gaussian distributions. If we were to suppose that $x \sim N(0, I)$ and $\omega \sim N(0, I)$, then the posterior would be $N\left(\begin{bmatrix} 1.5 \\ 3.5 \end{bmatrix}, \begin{bmatrix} 0.5 & 0 \\ 0 & 0.5 \end{bmatrix}\right)$. Again, it is not clear how to choose the radius for the truncated posterior. Figure 5.b shows the resulting ellipsis for radius 4, 6.5 and 9; none of them reflect the true state of affairs.

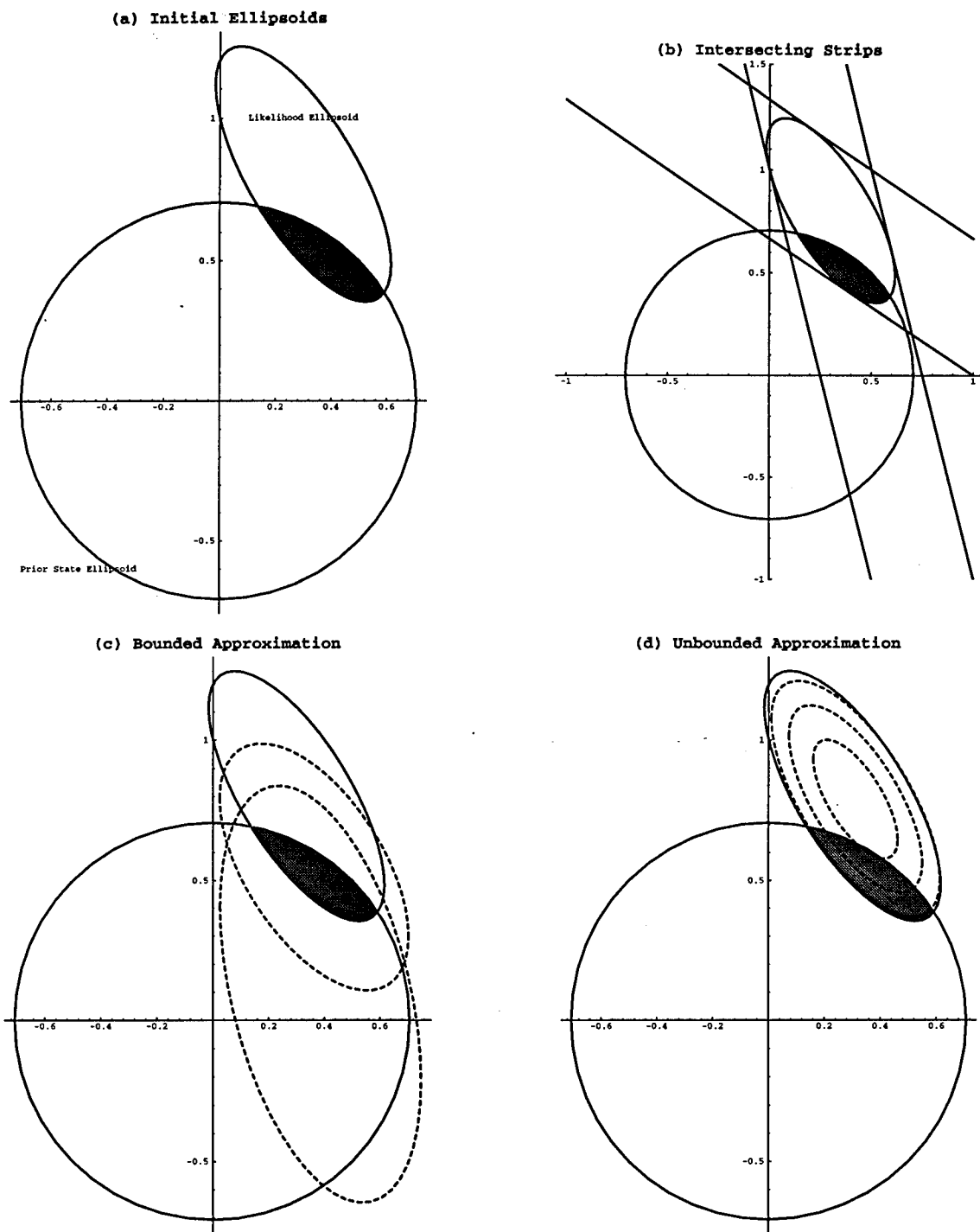


Figure 4: Example of Inference Algorithm for Bounded Distributions

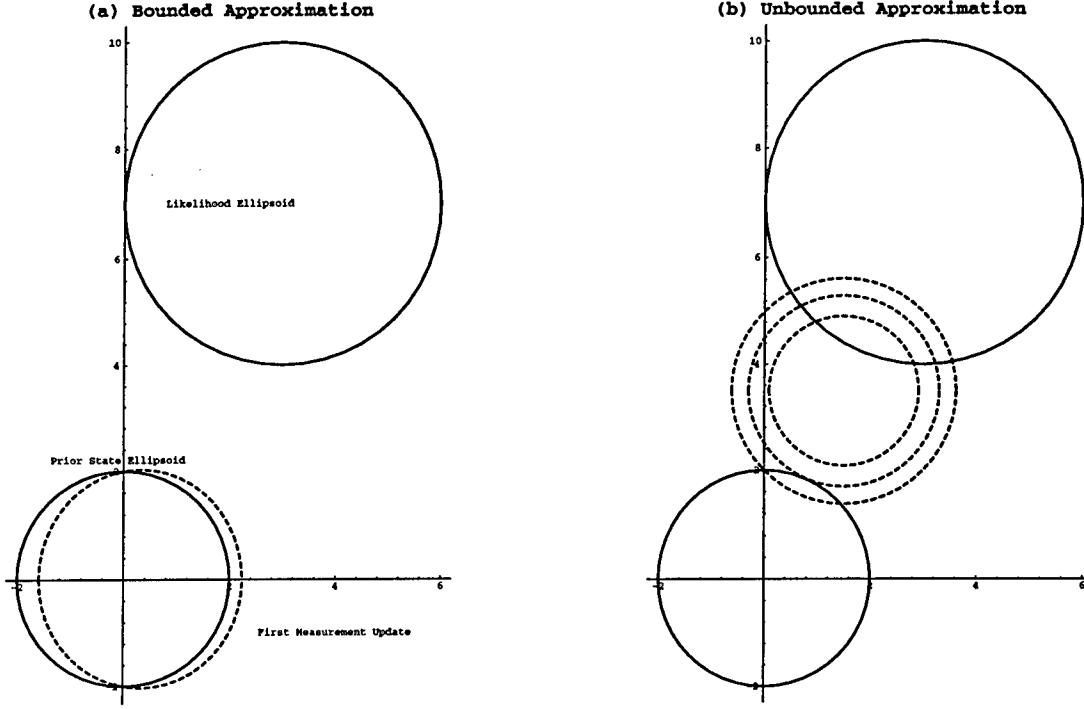


Figure 5: Example of Inference Algorithm for Bounded Distributions in Case of No Intersection

3.3 Inferences based on Maximum a Posteriori

A common estimate is the one obtained by maximizing the posterior density $p(x|z)$. Consider $x \sim N_{\nu, M, k}(\mu, P)$. We take the estimate $\hat{x} = \arg \max_x N_{\nu, M, k}(\mu, P)$. The optimization problem is simplified if we perform a double diagonalization. We look for:

$$\hat{w} = \arg \max_w N_{\xi, D, k}(0, I).$$

There are two possible cases:

- If $\xi^T D^{-1} \xi \leq k$ then 0 is inside the defining ellipsoid. Since the Gaussian is unimodal, the maximum a posteriori is $\hat{w} = 0$ (it must be transformed back to x).
- Otherwise, the maximum occurs on the boundary of the defining ellipsoid. So the optimization problem is:

$$\hat{w} = \arg \min_{\{w: (w-\xi)^T D^{-1} (w-\xi) = k\}} w^T w$$

Defining a Lagrange multiplier λ , we obtain the Lagrangian:

$$\hat{w} = \lambda \left[(w - \xi)^T D^{-1} (w - \xi) - k \right] + w^T w$$

Differentiation of the Lagrangian with respect to w and λ produces a set of equations:

$$w_i + \lambda \frac{w_i - \xi_i}{D_{ii}} = 0 \text{ for } i = 1, \dots, n \quad \text{and} \quad \sum_{i=1}^n \frac{(w_i - \xi_i)^2}{D_{ii}} = k.$$

If we isolate w_i in each one of the linear equations we obtain $w_i = \frac{\xi_i \lambda}{\lambda + D_{ii}}$. Combining all equations:

$$\sum_{i=1}^n D_{ii} \left(\frac{\xi_i}{\lambda + D_{ii}} \right)^2 = k. \quad (10)$$

We must solve this equation for λ . Consider the function $g(\lambda) = \sum_{i=1}^n D_{ii} \left(\frac{\xi_i}{\lambda + D_{ii}} \right)^2$. $g(\lambda)$ is always positive and goes to zero if λ goes either to infinity or minus infinity. $g(\lambda)$ has singularities at $-D_{ii}$ for all i ; at these points it goes to infinity. Suppose we order the values of $-D_{ii}$ so that we have $D_1 \leq D_2 \leq \dots \leq D_n < 0$. We know that the solution for equation (10) that corresponds to a minimum is a positive λ , since this is the only way we can have $w_i < \xi_i$ for all i . So the solution must lie in the interval (D_n, ∞) . So the constrained minimization is solved by performing a double diagonalization, sorting the values of D_{ii} and searching for the solution of equation (10) in the interval (D_n, ∞) . Since the function $g(\lambda)$ is decreasing in this interval, Newton's method will work if started with any value ϵ , larger than D_n but smaller than the solution.

3.4 Inferences based on Square Loss

Another common estimate is obtained by minimizing expected square loss: $L(x, \hat{x}) = (x - \hat{x})^T (x - \hat{x})$. The optimal estimate is the conditional mean $\hat{x} = E[x|z]$ [4]. Since we have the posterior, we can evaluate the mean by the methods previously presented.

4 Filtering with the Truncated Gaussian

Using all results so far presented, we can derive a filtering scheme akin to the Kalman filter but using truncated Gaussians.

Consider first the simple linear system:

$$x_{i+1} = Ax_i \quad (11)$$

$$z_{i+1} = Bx_{i+1} + \omega_i.$$

We assume x_0 distributed as a truncated Gaussian and ω distributed as zero mean radially truncated Gaussian. A and B are matrices of appropriate dimensions.

Given this set-up, we are interested in obtaining $p(x_{i+1}|z_1, \dots, z_{i+1})$:

$$\begin{aligned} p(x_{i+1}|z_1, \dots, z_{i+1}) &\propto p(z_{i+1}|x_{i+1}, z_1, \dots, z_i) \times p(x_{i+1}|z_1, \dots, z_i) \\ &= p(z_{i+1}|x_{i+1}) \times p(x_{i+1}|z_1, \dots, z_i) \end{aligned}$$

The distribution $p(z_{i+1}|x_{i+1})$ is directly obtained from the distribution of ω_i . We need to find the distribution $p(x_{i+1}|z_1, \dots, z_i) = p(Ax_i|z_1, \dots, z_i)$. Using theorem 4 we can obtain $p(Ax_i|z_1, \dots, z_i)$ from the available $p(x_i|z_1, \dots, z_i)$. So the filtering scheme has the following structure:

1. Propagate forward the uncertainty in x_i .
2. Use expression (4) and Fogel/Huang algorithm in order to fuse the incoming information z_{i+1} with x_i .

If necessary, the best estimate of x can be obtained at any time.

Extension of the filter above to a noisy state problem is straightforward. Consider the system model:

$$\begin{aligned} x_{i+1} &= Ax_i + \omega_{x_i} \\ z_{i+1} &= Bx_{i+1} + \omega_{z_i}. \end{aligned} \tag{12}$$

In this case we must obtain $p(x_{i+1}|z_1, \dots, z_i)$ by calculating $p(Ax_i|z_1, \dots, z_i)$ and combining it with $p(\omega_{z_i})$. Having done this, we can propagate forward the uncertainty in x_i (by properly approximating the summation of Ax_i and ω_{z_i}) and fuse it with the incoming information (through the Fogel/Huang algorithm).

5 Using the truncated Gaussian: Disparity Constraint

In this section we analyze a practical problem in Computer Vision using truncated Gaussians. The example will illustrate the power of the method.

Consider two cameras perfectly aligned so that all epipolar lines are parallel. Suppose a feature is detected in one camera at pixel x_2 . The correspondence to the feature must lie in the epipolar line in the other camera in some pixel x_1 . The question is: how much of the epipolar line should we explore in order to find the correspondence x_1 ?

From the basic optics of the problem we can derive the following equation [6]:

$$x_1 = x_2 + \frac{Bf}{z}, \tag{13}$$

where x_2 is the known coordinate of the correspondence (in pixels), B is the baseline distance (in meters), f is the focal length (in pixels) and z is the depth of the feature (in meters). Only z is unknown; as an example, we take $x_2 = 0$, $f = 600$ and $B = 0.5$.

A powerful constraint that can be used is the fact that the point x_1 cannot be arbitrarily far from the given value of x_2 . This constraint has not been justified or analyzed with the help of a statistical model.

The model (13) does not take into account the fact that every camera has finite depth of view. Not all values of z are allowed in a real camera. This can be modeled by a radially truncated Gaussian: $p_Z(z) = N_{\mu, \sigma^2, k}(\mu, \sigma^2)$. As an example, we take $\mu = 5$, $\sigma^2 = 4$ and $k = 3$. Note that the variance is large, reflecting a possible situation of ignorance with respect to z . The distribution of z is shown in figure 6.a. The distribution of x_1 is [10]:

$$p_{X_1}(x_1) = \frac{1}{Bfx_1^2} p_X\left(\frac{1}{Bfx_1}\right).$$

This distribution is graphed in figure 6.b.

A simple calculation will show that $x_1 \in [35.4438, 195.325]$. This could have been obtained from a tolerance set approach: if we simply assume $z \in [1.5359, 8.4641]$, then $x_1 \in [35.4438, 195.325]$. But notice that by using truncated distributions, we obtain much more information beyond the bounds: we know where we should invest time searching if we have only bounded resources; we know the means and variances of the quantities.

Consider now the following extension of the problem: suppose we have additional noise present in the imagery process, so that:

$$x_1 = x_2 + \frac{Bf}{z} + \omega.$$

where $\omega \sim N_{0,1,10}(0, 1)$. How are we to use this valuable information if we work only with bounds? We show now how to easily handle this in the framework of bounded distributions.

We can easily obtain the mean and variance of $x_2 + (Bf)/z$ by numerical integration. If this is considered a burden, we can linearize expression (13) around the mean of z :

$$x_1 \approx x_2 + \frac{Bf}{\mu} - \frac{Bf}{\mu^2}(x_1 - \mu)$$

and from this we can approximate the mean and variance of x_1 . With the values previously outlined, we obtain mean 60 and variance 382.25. Since the standard deviation is much smaller than the domain of definition, we take $N_{115.38, 6490.5, 1}(60, 382.25)$ as an approximation for the distribution of $x_2 + (Bf)/z$. Figure 6.c shows the agreement between approximation and correct distribution. Now we can proceed and combine this with ω ; we can actually use this measurement for filtering purposes through algorithms outlined in previous sections.

Since our approach is mainly approximative, it is interesting to compare its performance with similar methods in the literature. We already mentioned the inability of tolerance sets to capture all the information in the most simple situations. We now analyze another common strategy.

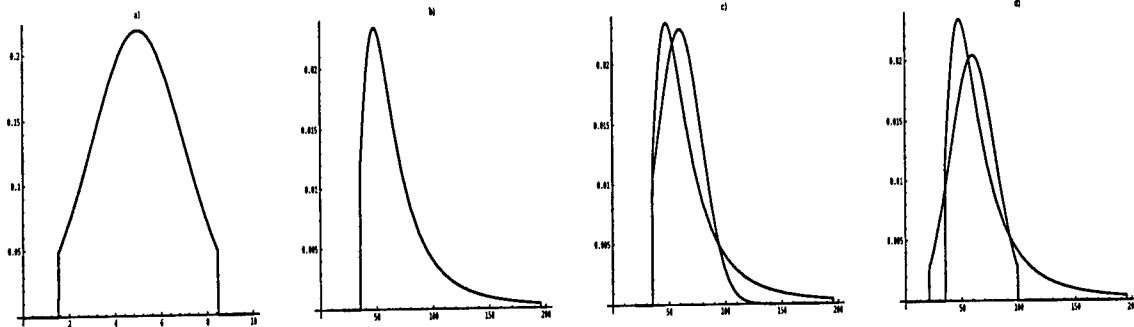


Figure 6: Analysis of the Disparity Constraint

Suppose we consider z as a unbounded Gaussian variable. We approximate the distribution of x_1 through a linearization: $x_1 \sim N(60, 382.25)$. How to justify the disparity constraint? We must arbitrarily set a threshold, based on extraneous heuristic knowledge, if we want to decide whether a value of x_1 is possible or impossible. The model neglects the physical source of the constraint, the camera's depth of field. And finally, the model is also an approximation, a bad one indeed. Figure 6.d shows the roughness of the approximation. The figure depicts the result of truncating the Gaussian $N(60, 382.25)$ after 2 standard deviations: some possible values of x_1 are discarded, many impossible values of x_1 are included. If the correct value of x_1 is located between 100 and 195, a failure will result. We could try to cover all possible values of x_1 , but unless we set the necessary thresholds in a very conservative way, some possible values of x_1 may be discarded as impossible.

6 Conclusion

The central idea of this research is: measurements that are bounded must be properly modeled in order to be consistently used. Our focus is not so much is data produced by selection mechanisms; we focus on data that is fundamentally bounded. We take this type of data to be fundamental in Robotics application such as computer vision or object recognition. Our proposal is to use a family of statistical models (the truncated Gaussian family) that captures measurement boundedness. We offered a comprehensive analysis of estimation aspects for the truncated Gaussian: algorithms for information handling, updating and filtering. An open problem is to find good approximations for summations of truncated Gaussians.

The statistical approach implied by bounded distributions is more interesting than pure tolerance sets in a variety of grounds. First, we can use the powerful language of Probability theory to draw conclusions. Second, usually the mean and variance of disturbances are known or can be estimated. We should use this valuable information if available, as an example illustrates: Consider a variables $a \in [-\pi/2, \pi/2]$ and a derived quantity $b = \tan(a)$, such that $a \in [-\pi/2, \pi/2]$. Without

any statistical knowledge, we only obtain that $b \in (-\infty, \infty)$! If instead we take $a \sim N_{0,1,1}(0, 1)$, then we can obtain $b \sim (1 + b^2)N_{0,1,1}(0, 1)$ [10], from where the analysis can proceed. Situations like this can appear, for example, when using triangulation in order to recover position. The adaptation of tolerance sets to the language of Probability theory is an area in which our models can be applied.

Another possible application of our results is to use a Gaussian model for unbounded data, but define selection mechanisms on the data. We discussed this strategy in section 3.2; this is a common problem in Statistics. Our methods have the advantage that, although approximate, they never assign probability zero to a possible value of the underlying unknowns. Algorithms here developed can find application in any of the usual problems of measurement error analysis, when measurements are bounded and when measurements are selected. In particular, the correct development of a Bayesian analysis of validation gates, as used by Bar-Shalom and Fortmann or Cox, is an immediate example of application.

Acknowledgement

We gratefully acknowledge Ofer Barkai for his help in the proof of theorem 6. This paper benefited from a several suggestions made by Mike Erdmann in the reading of a first draft.

References

- [1] Y. Bar-Shalom and T. E. Fortmann, *Tracking and Data Association*, Academic Press, Boston, 1988.
- [2] I. J. Cox, *A Review of Statistical Data Association Techniques for Motion Correspondence*, International Journal of Computer Vision, 10(1): 53-66, 1993.
- [3] H. Cramér, *Mathematical Methods of Statistics*, Princeton University Press, 1946.
- [4] M. DeGroot, *Optimal Statistical Decisions*, McGraw-Hill Inc., 1970.
- [5] M. Erdmann, *Randomization in Robot Tasks*, International Journal of Robotics Research, 11(5):399-436, October 1992.
- [6] O. Faugeras, *Three-Dimensional Computer Vision: a Geometric Viewpoint*, The MIT Press, 1993.
- [7] E. Fogel and Y. F. Huang, On the Value of Information in System Identification - Bounded Noise Case, *Automatica*, 18(2):229-238, 1982.
- [8] W. E. L. Grimson, *Object Recognition by Computer: the Role of Geometric Constraints*, MIT Press, 1990.

- [9] N. L. Johnson and S. Kotz, *Distributions in Statistics: Continuous Multivariate Distributions*, John Wiley and Sons, Inc., New York, 1972.
- [10] A. Papoulis, *Probability, Random Variables and Stochastic Processes*, 3rd edition, McGraw-Hill, 1991.
- [11] L. D. Servi and Y. C. Ho, Recursive Estimation in the Presence of Uniformly Distributed Measurement Noise, *IEEE Trans. on Automatic Control*, AC-26(2):563-564, 1981.

A Proofs of Theorems

Theorem 1 *Given a expected value is μ , a covariance matrix Q and the fact that a distribution is zero outside the set $\{x : (x - \mu)^T Q^{-1}(x - \mu) \leq k\}$, a maximum entropy distribution that obeys these conditions is a truncated Gaussian $N_{\mu, cQ, k'}(\mu, cQ)$, where $c = \Pr(\chi_n^2 \leq k)(\Pr(\chi_{n+2}^2 \leq k))^{-1}$ and $k' = k\Pr(\chi_{n+2}^2 \leq k)(\Pr(\chi_n^2 \leq k))^{-1}$.*

Proof: It is known that the distribution which is positive in a region $\mathcal{C}$ and which maximizes entropy for given mean and covariance matrix is of the form [10]:

$$p(x) = \Lambda_o \exp \left(-\Lambda_1 x - (x - \mu)^T \Lambda_2^{-1} (x - \mu) \right)$$

for $x \in \mathcal{C}$ and zero otherwise. Λ_o , Λ_1 and Λ_2 are respectively a scalar, a vector and a matrix.

The key to this optimization is to note that $w \log w$ is convex and the constraints are linear, so the solutions for Λ_o , Λ_1 and Λ_2 , if they exist, are unique and optimal. By straightforward substitution, we notice that if $p(x) = N_{\mu, cQ, k'}(\mu, cQ)$ then the mean is μ , the variance is Q [9] and the defining ellipsoid is $\{(x - \mu)Q^{-1}(x - \mu) \leq k\}$. So the truncated Gaussian is the maximum entropy solution. $\square$

Theorem 2 *The moment generating function of a truncated Gaussian $N_{\omega, D, k}(0, I)$ is:*

$$\phi(t) = \frac{q(\omega - t, D, k)}{q(\omega, D, k)} \exp \left(\frac{t^T t}{2} \right). \quad (14)$$

Proof:

$$\begin{aligned} \phi(t) &= \frac{(q(\omega, D, k))^{-1}}{\sqrt{(2\pi)^n}} \int_{(\sum_{i=1}^n d_i(z_i - \omega_i)^2 \leq k)} \exp \left(-\frac{1}{2} z^T z + t^T z \right) dz \\ &= \frac{(q(\omega, D, k))^{-1}}{\sqrt{(2\pi)^n}} \exp \left(\frac{t^T t}{2} \right) \int_{(\sum_{i=1}^n d_i(z_i - \omega_i)^2 \leq k)} \exp \left(-\frac{1}{2} (z - t)^T (z - t) \right) dz \\ &= \frac{(q(\omega, D, k))^{-1}}{\sqrt{(2\pi)^n}} \exp \left(\frac{t^T t}{2} \right) \int_{(\sum_{i=1}^n d_i(w_i + t_i - \omega_i)^2 \leq k)} \exp \left(-\frac{1}{2} w^T w \right) dw \\ &= (q(\omega, D, k))^{-1} \exp \left(\frac{t^T t}{2} \right) q(\omega - t, D, k) \end{aligned}$$

$\square$

Theorem 4 *If $x \sim N_{\nu, M, k}(\mu, P)$ and $y = Ax$ (where A is any non-singular square matrix) then $y \sim N_{A\nu, AMA^T, k}(A\mu, APA^T)$.*

Proof: The distribution of y is given by $p_Y(y) = \frac{1}{\det(A)} f_X(A^{-1}x)$; this gives:

$$\begin{aligned}
p_Y(y) &= \frac{1}{|\det(A)|} \frac{q(\mu, P, \nu, M, k)}{|(2\pi)^n \det(P)|^{\frac{1}{2}}} \times \\
&\quad \exp\left(-\frac{1}{2}(y - A\mu)^T A^{-T} P^{-1} A^{-1} (y - A\mu)\right) I_{\{(y-A\nu)^T A^{-T} M^{-1} A^{-1} (y-A\nu) \leq k\}}(y) \\
&= \frac{q(\mu, P, \nu, M, k)}{|(2\pi)^n \det(A) \det(M) \det(A)|^{\frac{1}{2}}} \times \\
&\quad \exp\left(-\frac{1}{2}(y - A\mu)^T (AP A^T)^{-1} (y - A\mu)\right) I_{\{(y-A\nu)^T (AM A^T)^{-1} (y-A\nu) \leq k\}}(y) \\
&= N_{A\nu, AM A^T, k}(A\mu, AM A^T)
\end{aligned}$$

□

Theorem 5 *The distribution of z has expected value $\bar{\mu}_z = \bar{\mu}_x + \bar{\mu}_y$ and covariance $\bar{P}_z = \bar{P}_x + \bar{P}_y$; the distribution of z is positive only inside the ellipsoid defined by $\{z : (z - \nu_z)^T M_z^{-1} (z - \nu_z) \leq 1\}$ where*

$$\nu_z = \nu_x + \nu_y \quad M_z = k_x M_x + k_y M_y$$

Proof: Given the independence of x and y , the expressions for $\bar{\mu}_z$ and $\bar{P}_z$ are straightforward. It is necessary to prove that $p(z)$ is positive in an ellipsoid $\{z : (z - \mu_z)^T M_z^{-1} (z - \mu_z) \leq k\}$. For this, apply a linear transformation to x and y so that:

$$\begin{aligned}
x' &= \Phi^T(\sqrt{\Lambda}^{-1} V^T)(x - \mu_x) \\
y' &= \Phi^T(\sqrt{\Lambda}^{-1} V^T)(y - \mu_y)
\end{aligned}$$

where: Λ is the eigenvalue matrix of M_x , V is the corresponding eigenvector matrix of M_y and Φ is the eigenvector matrix of $(\sqrt{\Lambda}^{-1} V^T) M_y (\sqrt{\Lambda}^{-1} V^T)$. As a result of this transformation:

$$\begin{aligned}
x' &\sim N_{0, I, k_x}(0, I) \\
y' &\sim N_{0, L, k_y}(0, L)
\end{aligned}$$

where I is an identity matrix of appropriate dimension and L is a diagonal matrix of appropriate dimension. The distribution of $u = x' + y'$ is given by the multi-dimensional convolution between x' and y' . So the distribution of u is positive in a set obtained by sliding a spheroid along an ellipsoid. In other words, u is positive inside an ellipsoid defined by $\{u : u^T (k_x I + k_y L)^{-1} u \leq 1\}$. In order to translate this into a constraint about z , we use two facts:

$$u = \left(\Phi^T(\sqrt{\Lambda}^{-1} V^T)\right)(x + y - \mu_x - \mu_y) = \left(\Phi^T(\sqrt{\Lambda}^{-1} V^T)\right)(z - \mu_z) \quad (15)$$

and

$$\begin{aligned} \left(V\sqrt{\Lambda}^{-1}\Phi(k_x I + k_y L)^{-1}\Phi^T\sqrt{\Lambda}^{-1}V^T \right)^{-1} &= k_x V\sqrt{\Lambda}\Phi\Phi^T\sqrt{\Lambda}V^T + k_y V\sqrt{\Lambda}\Phi L\Phi^T\sqrt{\Lambda}V^T \\ &= (k_x M_x + k_y M_y) \end{aligned} \quad (16)$$

in order to obtain:

$$\begin{aligned} \{u : u^T(k_x I + k_y L)^{-1}u \leq 1\} &= \\ \{z : \left(\Phi^T(\sqrt{\Lambda}^{-1}V^T)(z - \mu_z) \right)^T (k_x I + k_y L)^{-1} \left(\Phi^T(\sqrt{\Lambda}^{-1}V^T)(z - \mu_z) \right) \leq 1\} &= \\ \{z : (z - \mu_z)^T(k_x M_x + k_y M_y)^{-1}(z - \mu_z) \leq 1\} &= \\ \{z : (z - \mu_z)^T M_z^{-1}(z - \mu_z) \leq 1\} \end{aligned}$$

□

Theorem 6 For $\nu_z = \mu_x + \mu_y$, $M_z = M_x + M_y$, $k_z = k_x = k_y$, $\sup_z |G_z - N_{\mu_z, M_z, k}(\mu_z, M_z)|$ is $\mathcal{O}(\exp(-k/2))$.

Proof: It is enough to prove that, for a given k_o , there exists a constant α (which depends only on n) such that

$$\frac{d \sup_z |G_z - N_{\mu_z, M_z, k}(\mu_z, M_z)|}{dk} < -\alpha \frac{\exp(-k/2)}{2} \text{ for } k > k_o.$$

We start with the same double diagonalization used in the previous theorem. Consider the vector $[u, v]^T$ given by:

$$\begin{aligned} u &= x' + y' \\ v &= y' \end{aligned}$$

The distribution of u is given by the multi-dimensional convolution between x' and y' :

$$p_U(u) = \int_{W_u} p_{X'}(u - v) p_{Y'}(v) dv$$

where $W_u = \{v : (u - v)^T(u - v) \leq k\} \cap \{v : v^T L^{-1}v \leq k\}$. Using the distributions of x' and y' we obtain:

$$p_U(u) = \int_{W_u} \frac{q(k)^{-2}}{|(2\pi)^{2n} \det(L)|^{\frac{1}{2}}} \exp\left(-\frac{1}{2} \left((u - v)^T(u - v) + v^T L^{-1}v\right)\right) dv \quad (17)$$

Expression (17) can be simplified by noting two facts:

- Call l_i an eigenvalue of L (L is diagonal):

$$\det(L) = \prod_i l_i = \prod_i (1 + l_i) \frac{l_i}{1 + l_i} = \prod_i (1 + l_i) \prod_i (1 + 1/l_i)^{-1} = \det(I + L) \det(I + L^{-1})^{-1}.$$

- Consider the quadratic form $(u - v)^T(u - v) + v^T L^{-1}v$; it is a summation of terms:

$$\begin{aligned} (u_i - v_i)^2 + \frac{v_i^2}{l_i} &= \frac{u_i^2}{1 + l_i} - \frac{u_i^2}{1 + l_i} + u_i^2 - 2u_i v_i + \frac{v_i^2}{l_i} \\ &= \frac{u_i^2}{1 + l_i} - \frac{1 + l_i}{l_i} \left(\frac{l_i^2 u_i^2}{(1 + l_i)^2} - 2u_i v_i \frac{l_i}{1 + l_i} + v_i^2 \right) \\ &= \frac{u_i^2}{1 + l_i} - \frac{1 + l_i}{l_i} (v_i - (1 + 1/l_i)u_i)^2. \end{aligned}$$

Applying these results, expression (17) can be written as:

$$p_U(u) = I(u) \frac{q(k)^{-1}}{|(2\pi)^n \det(I + L)|^{\frac{1}{2}}} \exp\left(-\frac{1}{2}u^T(I + L)^{-1}u\right)$$

where

$$I(u) = \int_{W_u} \frac{q(k)^{-1}}{|(2\pi)^n \det((I + L^{-1})^{-1})|^{\frac{1}{2}}} \exp\left(-\frac{1}{2}(v - (I + L^{-1})u)^T(I + L^{-1})(v - (I + L^{-1})u)\right) dv$$

so $u \sim I(u)N_k(0, I + L)$. By using expressions (15) and (16), we transform $p_U(u)$ into $p(z)$:

$$p(z) = I(z) \frac{q(k)^{-1}}{|(2\pi)^n \det(M_z)|^{\frac{1}{2}}} \exp\left(-\frac{1}{2}(z - \mu_z)^T M_z^{-1}(z - \mu_z)\right)$$

where W_u is transformed into W_z by straightforward substitution and

$$\begin{aligned} I(z) &= q(k)^{-1} \int_{W_z} \frac{1}{|(2\pi)^n \det((I + L^{-1})^{-1})|^{\frac{1}{2}}} \\ &\quad \exp\left(-\frac{1}{2}\left(v - (I + L^{-1})\left(\Phi^T \sqrt{\Lambda}^{-1} V^T\right)(z - \mu_z)\right)^T (I + L^{-1})\right. \\ &\quad \left.\left(v - (I + L^{-1})\left(\Phi^T \sqrt{\Lambda}^{-1} V^T\right)(z - \mu_z)\right)\right) dv. \end{aligned}$$

Observe that $I(z)$ is equal to $q(k)^{-1}$ times the integral of a Gaussian in a region of $\mathbb{R}^n$ (so this integral is strictly less than 1). We can state $I(z) < q(k)^{-1}$. So we can write:

$$\begin{aligned} \sup_z |G_z - N_{\mu_z, M_z, k_z}(\mu_z, M_z)| &= \sup_z |I(z)N_k(\mu_z, M_z) - N_k(\mu_z, M_z)| < \\ &= (q(k)^{-1} - 1) \sup_u N_{\mu_z, M_z, k}(\mu_z, M_z) = \frac{q(k)^{-1}(q(k)^{-1} - 1)}{|(2\pi)^n \det(M_z)|} \end{aligned}$$

Consider the following function of k :

$$g(k) = q(k)^{-1}(q(k)^{-1} - 1)$$

Then:

$$\frac{dg(k)}{dk} = - \left[\frac{2 - \int_0^k \frac{1}{2^{n/2}\Gamma(n/2)} x^{n/2-1} \exp(-x/2) dx}{\left(\int_0^k \frac{1}{2^{n/2}\Gamma(n/2)} x^{n/2-1} \exp(-x/2) dx \right)^3} \right] \frac{k^{n/2-1} \exp\left(-\frac{k}{2}\right)}{2^{n/2}\Gamma(n/2)}$$

and since the term in brackets is larger than one:

$$\frac{dg(k)}{dk} < - \frac{k^{n/2-1} \exp\left(-\frac{k}{2}\right)}{2^{n/2}\Gamma(n/2)} < - \frac{k_o^{n/2-1} \exp\left(-\frac{k_o}{2}\right)}{2^{n/2}\Gamma(n/2)}$$

for all $k > k_o$. By collecting constants in α , the result is proven. $\square$

Claims in Appendix 3

q_{i+1} is given by equations (9).

Call V_{i+1} the square of the volume of the ellipsoid Θ_{i+1} . Since Θ_{i+1} always contains the desired intersections for any value q_{i+1} , we are interested in the smallest possible Θ_{i+1} . The value of q_{i+1} is then chosen so that V_{i+1} is minimal.

We have $V_{i+1} = c_n b_{i+1}^n \det(P_{i+1})$, where c_n is the volume of a n -dimensional spheroid. Some algebraic manipulations show that in order to minimize V_{i+1} we must minimize

$$\frac{b_i + q_{i+1} - a q_{i+1} (y_i - B_i \theta_i)^2 / (1 + a q_{i+1} B_i P_i B_i^T)}{1 + a q_{i+1} B_i P_i B_i^T}$$

with respect to q_{i+1} .

Claims in Section B

1. Equation (19) has a unique solution if $0 < \sigma^2 < 1/(n+2)$.

Proof: Call $h = c\sigma^2$; there is a one-to-one correspondence between c and h . So we can study the following equation:

$$h \frac{\Pr(\chi_{n+2}^2 \leq 1/h)}{\Pr(\chi_n^2 \leq 1/h)} = \sigma^2.$$

We analyze the function $g(h) = h \frac{\Pr(\chi_{n+2}^2 \leq 1/h)}{\Pr(\chi_n^2 \leq 1/h)}$. Notice that:

- $\lim_{h \rightarrow 0} g(h) = 0$.
- $\lim_{h \rightarrow \infty} g(h) = \frac{1}{n+2}$. This is obtained by using the following expansion of the chi-square probability: $Pr(\chi_n^2 \leq x) = e^{-x/2} (x/2)^{n/2} \sum_{i=0}^{\infty} \frac{(x/2)^i}{\Gamma(n/2+1+i)}$.
- $g(h)$ is equal to:

$$g(h) = \int_{x^T x \leq 1} \frac{1}{\sqrt{(2\pi h)^n}} \exp\left(-\frac{x^T x}{2h}\right) dx,$$

which is increasing with h (variance increases with h).

So $g(h)$ is increasing from 0 to $1/(n+2)$. Any value of σ^2 strictly between these limits will produce a solution for d and hence for c . $\square$

2. Integral $\int_{x^T x \leq 1} x x^T p_4(x) dx$ is equal to expression (23).

Proof: We have four terms to integrate:

- $\int_{x^T x \leq 1} x x^T 1 p_4(x) dx$ is equal to $\sigma^2 I$.
- $\int_{x^T x \leq 1} x x^T \sigma^2 \text{tr}(L) p_4(x) dx$ is equal to $\sigma^4 \text{tr}(L) I$.
- $\int_{x^T x \leq 1} x x^T \frac{\mu^T}{\sigma^2} x p_4(x) dx$ is zero.
- $\int_{x^T x \leq 1} x x^T (x^T L x) p_4(x) dx$ is a diagonal matrix. Consider element i in the diagonal of this matrix. It is a summation of one term $L_{ii} \int_{x^T x \leq 1} x_i^4 p_4(x) dx$, and $n-1$ terms $L_{jj} \int_{x^T x \leq 1} x_i^2 x_j^2 p_4(x) dx$. The fourth moment of any x_i is $6\sigma^4 \gamma$ and the second crossed moments of $x_i^2 x_j^2$ are $2\sigma^4 \gamma$. So the matrix is equal to $4\sigma^4 \gamma L + 2\sigma^2 \gamma \text{tr}(L) I$.

Putting all these terms together, we obtain expression (23). $\square$

B Approximations of Bounded Distributions given First and Second Moments

One may want to approximate a distribution with bounded support by a truncated Gaussian. For example, it may be necessary to obtain an approximation for summations of truncated Gaussians. A natural approximation is to pick a truncated Gaussian that matches the first and second moments of the original distribution.

In this section, we assume that the distribution to be approximated is positive in the spheroid $\{x : x^T x \leq 1\}$ and has a diagonal second moment matrix Q . This may be assumed without loss of generality if the original distribution is positive in an ellipsoid (since we can perform a double diagonalization). Approximations are valid if no eigenvalue of Q exceeds $1/(n+2)$, where n is the dimension of the space.

B.1 Equal Variances Case

In this section we consider a distribution $f_X(x)$ with $Q = \sigma^2 I$ and mean μ . Take an initial approximation of the form:

$$p_1(x) = \frac{1}{q\sqrt{(2\pi c\sigma^2)^n}} \exp\left(-\frac{x^T x}{2c\sigma^2}\right) I_{x^T c \leq 1}(x), \quad (18)$$

where c is the solution of the equation:

$$c \frac{\Pr(\chi_{n+2}^2 \leq 1/(c\sigma^2))}{\Pr(\chi_n^2 \leq 1/(c\sigma^2))} = 1. \quad (19)$$

This equation always has a unique solution if $0 < \sigma^2 < 1/(n+2)$, as proved in Appendix A. The limiting value $1/(n+2)$ corresponds to the second moment of a constant distribution.

The variance of X under the approximation $p_1(x)$ is exactly $\sigma^2 I$. But the mean is zero. We adjust the mean through a linear term:

$$p_2(x) = \frac{1}{q\sqrt{(2\pi c\sigma^2)^n}} \exp\left(-\frac{x^T x}{2c\sigma^2}\right) \left(1 + \frac{\mu^T}{\sigma^2} x\right) I_{x^T c \leq 1}(x). \quad (20)$$

The function $p_2(x)$ is not a distribution, since it may have negative values. So we use another approximation: $(1 + (\mu^T/\sigma^2)x) \approx \exp(\mu^T/\sigma^2 x)$, in order to obtain:

$$p_3(x) = \frac{1}{q'\sqrt{(2\pi c\sigma^2)^n}} \exp\left(-\frac{x^T x - 2c\mu^T x + c^2\mu^T \mu}{2c\sigma^2}\right) I_{x^T c \leq 1}(x). \quad (21)$$

Essentially, we are constructing a Edgeworth-like expansion [3] for $f_X(x)$ and then approximating it by an exponential. If we have higher moments we can increase the number of terms in the approximation by building a sequence of orthogonal polynomials (with norm based on $\exp(-x^2)$), as detailed by Cramér [3].

Example 4 Consider the distribution $f(x) = ((1-x)/2)I_{x^2 \leq 1}(x)$, with mean $-1/3$ and variance $2/9$. This distribution is highly skewed. We obtain $c = 1.51696$ and:

$$p_3(x) = \frac{1}{1.16599\sqrt{2\pi(2/9)c}} \exp\left(-\frac{(x + (c/3))^2}{2(2/9)c}\right) I_{x^2 \leq 1}(x).$$

Figure 7 shows the graphs of $f(x)$ and $p_3(x)$. $\square$

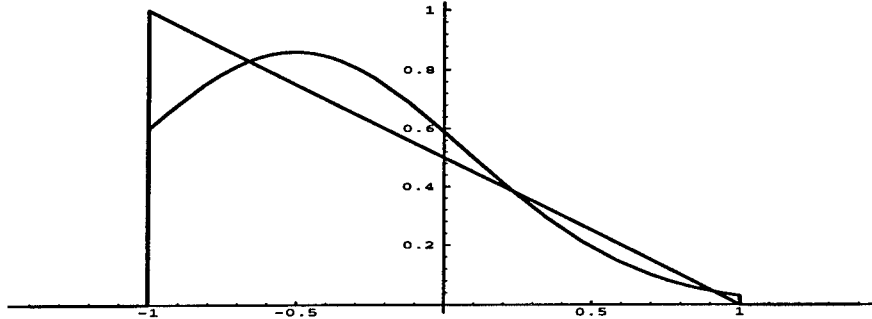


Figure 7: Distribution $f(x)$ and the approximation $p_3(x)$.

B.2 General Case

We relax the assumptions of last section, taking Q to be diagonal. Take function $p_1(\cdot)$ in expression (18) as an initial approximation. There is freedom in the choice of σ^2 . We choose σ^2 to be the largest eigenvalue of Q ; the constant c is the solution of equation (19). Approximation $p_1(\cdot)$ does not produce the correct mean and variance. First we adjust the mean through a linear term obtaining function $p_2(\cdot)$ as in expression (20). Approximation $p_2(\cdot)$ is flatter than the original distribution. We can stop here and obtain $p_3(\cdot)$ as in expression (21).

Alternatively, we can improve the approximation by adding more terms to $p_2(\cdot)$. Pick a quadratic term:

$$p_4(x) = \frac{1}{q\sqrt{(2\pi c\sigma^2)^n}} \exp\left(-\frac{x^T x}{2c\sigma^2}\right) \left(1 + \frac{\mu^T}{\sigma^2}x + (x^T Lx - \sigma^2 \text{tr}(L))\right) I_{x^T x \leq 1}(x).$$

L is a diagonal matrix that produces the correct second moment for $p_4(\cdot)$. Since $p_2(\cdot)$ is flatter than the original distribution, we define L as a diagonal negative definite matrix.

Function $p_4(x)$ may not be a distribution. So we use another approximation:

$$\left(1 - \sigma^2 \text{tr}(L) + \frac{\mu^T}{\sigma^2}x + x^T Lx\right) \approx \exp\left(l' + \frac{l'}{1 - \sigma^2 \text{tr}(L)}\mu/\sigma^2 x + \frac{l'}{1 - \sigma^2 \text{tr}(L)}x^T Lx\right), \quad (22)$$

where l' must be obtained numerically by solving $l' \exp(l') = 1 - \sigma^2 \text{tr}(L)$. Solution is always possible since L is defined to be negative definite. The final approximation is the product of $p_1(x)$ and the approximation in expression (22).

We now indicate how to obtain L . The value of $\int_{x^T x \leq 1} x x^T p_4(x) dx$ is (Appendix A):

$$4\gamma\sigma^4 L + (2\gamma\sigma^4 - \sigma^4)\text{tr}(L)I + \sigma^2 I \quad (23)$$

where $\gamma = c^2 Pr(\chi_{n+4}^2 \leq (1/(c\sigma^2)))(Pr(\chi_n^2 \leq (1/(c\sigma^2))))^{-1}$. Call A the vector of diagonal elements of Q and B the vector of diagonal elements of L ; then ($\mathbf{1}$ is a vector of ones):

$$\left(I + \left(\frac{1}{2} - \frac{1}{4\gamma}\right) \mathbf{1}\mathbf{1}^T\right) B = \frac{A - \sigma^2 \mathbf{1}}{4\gamma\sigma^4}.$$

Closed form solution of this equation is:

$$B = \left(I - \frac{2\gamma - 1}{(4 + 2n)\gamma - n} \mathbf{1}\mathbf{1}^T\right) \frac{A - \sigma^2 \mathbf{1}}{4\gamma\sigma^4}. \quad (24)$$

From B we can obtain L . These formal manipulations do not guarantee that the resulting L matrix is negative definite, as required. Notice that $(A - \sigma^2 \mathbf{1})/(4\gamma\sigma^4)$ is a vector of non-positive numbers. From the geometry of expression (24) it is clear that all eigenvalues of Q must lie in a wedge in the first quadrant. If this fails to happen, we can increase σ^2 . If this fails, we can set the negative elements of L to zero. Albeit crude, the strategy will always improve upon $p_2(\cdot)$.