

THEMIS SIGNAL ANALYSIS STATISTICS RESEARCH PROGRAM

PARTIAL PRIOR INFORMATION: SOME EMPIRICAL

BAYES AND G-MINIMAX DECISION FUNCTIONS

by

Stephen L. George

Technical Report No. 48
Department of Statistics THEMIS Contract

Reproduction in whole or in part is permitted
for any purpose of the United States Government.

This document has been approved for public release
and sale; its distribution is unlimited.

THEMIS SIGNAL ANALYSIS STATISTICS RESEARCH PROGRAM

PARTIAL PRIOR INFORMATION: SOME EMPIRICAL

BAYES AND G-MINIMAX DECISION FUNCTIONS

by

Stephen L. George

Technical Report No. 48
Department of Statistics THEMIS Contract

October 30, 1969

Research sponsored by the Office of Naval Research
Contract N00014-68-A-0515
Project NR 042-260

Reproduction in whole or in part is permitted
for any purpose of the United States Government.

This document has been approved for public release
and sale; its distribution is unlimited.

DEPARTMENT OF STATISTICS
Southern Methodist University

PARTIAL PRIOR INFORMATION: SOME EMPIRICAL
BAYES AND G-MINIMAX DECISION FUNCTIONS

Approved by:

PARTIAL PRIOR INFORMATION: SOME EMPIRICAL
BAYES AND G-MINIMAX DECISION FUNCTIONS

A Thesis Presented to the Faculty of the Graduate School

of

Southern Methodist University

in

Partial Fulfillment of the Requirements

for the degree of

Doctor of Philosophy

with a

Major in Statistics

by

Stephen L. George

(B.A., Texas Technological College, 1965)

(M.E.S., North Carolina State University, 1967)

October 30, 1969

George, Stephen L. B.A., Texas Technological

College, 1965
M.E.S., North Carolina
State University, 1967

Partial Prior Information: Some Empirical Bayes and ϕ -Minimax Decision Functions

Adviser: Associate Professor Richard P. Bland

Doctor of Philosophy degree conferred December 20, 1969

Dissertation completed October 30, 1969

If the prior probability distribution function G of the parameters in a statistical decision theory model is not completely specified then a Bayes decision function cannot be obtained. However, in many cases there may be some partial (i.e., incomplete) prior information concerning G .

The types of partial prior information considered in this paper is of two kinds: (i) N past observations on the compound distribution $F_G(t)$ or (ii) knowledge of a restrictive class G to which the unknown G is assumed to belong.

When past observations are available, empirical Bayes decision functions are found which are: (i) asymptotically optimal (ii) easy to apply and (iii) better than some optimal non-Bayes decision function for reasonably small N . When only knowledge of the class G is assumed, ϕ -minimax decision functions are found.

In all cases, estimators for the parameters of the normal and

Bernoulli processes are found using a quadratic loss function. These

estimators are then compared to each other and to other well-known estimators by means of their Bayes risks and, in the empirical Bayes case, by means of their global risks for small N .

ACKNOWLEDGEMENTS

I wish to thank my major adviser, Dr. Richard P. Bland for his advice and helpful comments given during the writing of this paper. In addition I would like to thank Dr. John T. Webster and Dr. Donald B. Owen for their constructive criticism of certain aspects of the paper.

This paper was supported in part by National Institute of Health grant number 2 TOL GM00951-08

TABLE OF CONTENTS

Page	
iv	ABSTRACT
vi	ACKNOWLEDGMENTS
ix	LIST OF TABLES
x	LISTS OF ILLUSTRATIONS
	Chapter
I	I. Introduction
1	1.1 Bayesian Approach
3	1.2 G-Minimax Decision Functions.
5	1.3 Empirical Bayes Decision Functions.
8	1.4 Problems Considered
11	II. PRELIMINARY RESULTS
11	2.1 General Structural Elements
14	2.2 Bayes Decision Functions.
18	2.3 Empirical Bayes Decision Functions.
24	2.4 G-Minimax Decision Functions.
34	III. NORMAL PROCESS I: EMPIRICAL BAYES DECISION FUNCTIONS.
34	3.1 Definitions and Initial Results
37	3.2 Estimation of λ_{N+1}
51	3.3 Numerical Results
62	IV. NORMAL PROCESS II: G-MINIMAX DECISION FUNCTIONS
62	4.1 Introduction.
63	4.2 Estimation of λ_{N+1} : G Conjugate
72	4.3 Estimation of λ_{N+1} : Moments Known.
76	4.4 Asymptotic G-Minimax Decision Functions

LIST OF REFERENCES	161
APPENDIX	144
6.4 Asymptotic ϵ -Minimax Decision Functions.	140
6.3 Estimation of λ^{N+1} : Moments known	136
6.2 Estimation of λ^{N+1} : G conjugate	129
6.1 Introduction	128
VI. BERNOLLI PROCESS II: ϵ -MINIMAX DECISION FUNCTIONS	128
5.3 Numerical Results.	109
5.2 Estimation of λ^{N+1}	87
5.1 Definitions and Initial results.	80
V. BERNOLLI PROCESS I: EB DECISION FUNCTIONS	80

LIST OF TABLES

Table Page

1. Global Risks of Various Decision Functions Using a $N(0,1)$ Prior.	58
2. Global Risks of Various Decision Functions Using a $N(0,1)$ Prior.	59
3. Global Risks ($N'=20, R'=10$)	115
4. Global Risks ($N'=6, R'=3$)	116
5. Global Risks ($N'=2, R'=1$)	117
6. Global Risks ($N'=24, R'=6$)	118
7. Global Risks ($N'=12, R'=3$)	119
8. Global Risks ($N'=4, R'=1$)	120
9. Values of N_0	121

12.	Preference Regions for δ_1 and δ_T	137
11.	Preference Regions for δ_0 and δ_T	132
10.	$R(\delta)$ as a Function of N ($N'=4$, $R'=1$, $M=5$)	127
9.	$R(\delta)$ as a Function of N ($N'=12$, $R'=3$, $M=5$)	126
8.	$R(\delta)$ as a Function of N ($N'=24$, $R'=1$, $M=5$)	125
7.	$R(\delta)$ as a Function of N ($N'=2$, $R'=1$, $M=5$)	124
6.	$R(\delta)$ as a Function of N ($N'=6$, $R'=3$, $M=5$)	123
5.	$R(\delta)$ as a Function of N ($N'=20$, $R'=10$, $M=5$)	122
4.	Upper and Lower Limits of μ and μ^2	86
3.	Bayes Risk of δ^{G+} and δ_T as a Function of μ^2	68
2.	$R(\delta)$ as a Function of N ($M'=0$, $N'=10$, $M=5$)	61
1.	$R(\delta)$ as a Function of N ($M'=0$, $N'=1$, $M=5$)	60

Page

Figure

LIST OF ILLUSTRATIONS

INTRODUCTION

1.1 The Bayesian Approach

In the usual formulation of problems in statistical decision

theory the probability distribution function of an observable random variable X defined on a sample space $\mathcal{X}$ is assumed to be a member of some specified class. For probability distribution functions. The elements $F_\lambda \in \mathcal{F}$ are indexed by the so-called "parameter space" Λ

with generic element λ . We assume that corresponding to each F_λ in

$\mathcal{F}$ there exists a probability density function (p. d. f.) $f_\lambda = f(x|\lambda)$

with respect to some σ -finite measure μ defined on the space $\mathcal{X}$.

In the non-Bayesian formulation no prior information regarding

the unknown parameter λ is assumed to exist. Optimality criteria

can then be found which make no use of prior information and suitable

decision functions derived therefrom.

However, in some situations it is reasonable to assume that λ

is itself a random variable so that there exists a prior distribution

function G of λ on Λ . In the remainder of this paper we use $\tilde{\lambda}$ to

denote a random variable taking values in Λ and λ for a particular

outcome of the random variable $\tilde{\lambda}$.

$$(1.1.2) \quad R(\delta, G) = \int_V R(\delta, \lambda) dG(\lambda).$$

for any decision function (d.f.) δ for a particular λ and finally minimizing the "average" or "Bayes" risk

$$(1.1.1) \quad R(\delta, \lambda) = \int_X L(\delta(x), \lambda) F(x|\lambda) d\mu(x)$$

the "risk":

action space and $\delta(\cdot)$ is a decision function $\delta: X \rightarrow A$, then finding negative real-valued loss function $L(\delta(X), \lambda)$ on $A \times \Lambda$ where A is the function can be found. In general, this involves defining a non-

is reduced to the standard Bayes decision problem and a Bayes decision in a given decision situation is completely specified then the problem on X according to the p.d.f. $f(\cdot|\lambda_0)$. If the particular G in effect tion function $G(\lambda)$ and then generating an observation (or observations) selecting an element λ_0 from Λ according to the probability distribu-

tion of the random variable X may be thought of as first randomly G on Λ . Conceptually, the process leading to a particular observa-

been randomly selected from F according to the distribution function in effect during a particular experiment is thought of as having space Λ . In this case, the probability distribution function $F(x|\lambda)$ $g = g(\lambda)$ with respect to some σ -finite measure μ^* defined on the

Furthermore we assume that for any G there exists a p.d.f.

$$R(\delta_1, G) \leq R(\delta_0, G) \text{ for all } G \text{ in } \mathcal{G}.$$

information. Hopefully, it will be possible in many cases to have least as good as some optimum decision function δ_0 which ignores such G it may be possible to obtain a decision function δ_1 which is at achieved in this situation but by using the information on the class bution functions G on Λ . Of course the minimum Bayes risk cannot be unrestrictive; that is, G may be the class of all probability distributions of some class $\mathcal{G}$ of distribution functions on Λ where G may be completely specified. In this case, it is assumed only that G is a and hence the Bayes d.f. δ_G is the best attainable but G is not completely specified. Suppose that the Bayesian formulation of the problem is accepted

1.2 $\mathcal{G}$ -Minimax Decision Functions

Bayes envelope function.

is the minimum Bayes risk attainable and $R(G)$ is often called the

$$(1.1.4) \quad R(\delta_G, G) = R(G) = \inf_{\delta} R(\delta, G)$$

for any other d.f. δ . Then

$$(1.1.3) \quad R(\delta_G, G) \leq R(\delta, G)$$

That is, we try to find a d.f. δ_G such that

To this end, the following concept of a G -minimax (or "minimax in G ") decision function will be useful: δ^* is a G -minimax decision function if $\sup_{G \in G} R(\delta^*, G) \leq \sup_{G \in G} R(\delta, G)$ for any other d.f. δ . If such a d.f. δ^* exists for the class G it will be used as the "best" obtainable d.f. when no additional prior information is available. Of course the ease of finding such decision functions depends primarily on the class G .

Cote and Skibinsky [7] solve such a problem when G is assumed to be the family of all distributions on Λ such that $P(\lambda_1 < \lambda_2) > 1 - \alpha$. Then in the binomial and normal case they find an estimator better than the usual estimator in each case. For the binomial case, Weiler [48] assumes G to be the Beta family and also assumes one of three additional bits of information are known: (1) prior observations $x_1, \dots, x_N$ are present (11) $\int \lambda dG(\lambda) = \mu$ is known (mean of prior) with the probability of exceeding 2μ also known, and (111) $P(\lambda_1 < \lambda_2)$ is known. In each case he obtains estimators of the parameters of the prior distribution without reference to G -minimax optimality. Blum and Rosenblatt [3] specify G as $G_{F, P} = \{G: G = pF + (1-p)H\}$ where F and H are distributions on Λ , $0 < p < 1$, F known, H arbitrary. The risks for various G -minimax decision functions derived in this paper are compared with the usual decision function in order to demonstrate the reduction in average risk possible for various amounts of prior knowledge concerning the class G . These cases are described more fully below.

1.3 Empirical Bayes Decision Functions

In addition to knowledge of the class G of distribution functions

on λ we may have N independent "past" observations $x_1, \dots, x_N$ on the

random variable X such that each x_i is a random outcome from the p.d.f. $f(x|\lambda_i)$ with the λ_i 's themselves chosen randomly from a fixed G in $\mathcal{G}$.

That is, we have an independent sequence of bivariate random variables $(x_1, \lambda_1), \dots, (x_N, \lambda_N), (x_{N+1}, \lambda_{N+1})$ where the x_i 's are observable random

variables and the λ_i 's are unobservable random variables. The observations $x_1, \dots, x_N$ are called "past" observations, x_{N+1} is the "current"

observation and the problem consists of making some inference concerning the unknown λ_{N+1} . Note that each x_i has the "compound" or "unconditional"

distribution

$$f_G(x) = \int_{\mathcal{G}} f(x|\lambda) dG(\lambda). \quad (1.3.1)$$

Robbins[33] was the first to explore the possibility of using the

past observations to approximate to certain Bayes procedures and in so doing is credited by Neyman [26] for achieving a "breakthrough" in the

theory of statistical decision making. Robbins coined the phrase "empirical

Bayes" (EB) to refer to any procedure which used the past observations in

a systematic manner to approximate to the optimum Bayes procedures. Some

writers insist that the decision function possess the property of as-

ymptotic optimality (defined below) in order to be called an empirical

Bayes d.f., but the term is used in this paper in its looser sense of a

sequence $\{\delta_N\}$ of d.f.'s of the form

$$\delta_N(\cdot) = \delta_N(x_1, \dots, x_N; \cdot) \quad (1.3.2)$$

with values in A . That is, δ_N is a function of x^{N+1} whose form depends on $x_1, \dots, x_N$. Thus it is possible to speak of asymptotically optimal (a.o.) EB procedures and non-a.o. EB procedures. It is shown later that some non-a.o. EB d.f.'s at least have some desirable properties for small N .

Formally, we say that an EB procedure $\{\delta_N\}$ is asymptotically optimal iff

$$\lim_{N \rightarrow \infty} E[R_N(\delta_N, G)] = R(\delta_G, G) \quad (1.3.3)$$

for any $G \in \mathcal{G}$ where δ_G is the Bayes d.f., E denotes expectation with respect to the N past observations $\{x_1, \dots, x_N\}$ and $R_N(\delta_N, G)$ is the

Bayes risk for the given (fixed) set of past observations. The quantity $E[R_N(\delta_N, G)]$ is called the "global risk" of the d.f. δ_N for a given G and will be denoted by $R(\delta_N, G)$ to distinguish it from the risk function $R_N(\delta_N, G)$.

The empirical Bayes decision situation has received considerable attention from various sources. Prominent contributors to the EB literature and the closely related compound decision theory are Robbins

[32], Samuel [39], and Johns [8]. More recently, articles by Deely and Zimmer [12], Maritz [22] and Krut'chhoff [20] have appeared.

Many of the early EB d.f.'s are "simple" in the sense that they merely attempt to estimate the Bayes d.f. δ_G and although a.o., often converge so slowly that for small N some optimum non-Bayes procedure which completely ignores the previous samples often has a smaller Bayes risk. In fact, Maritz [24] shows that some of these early proposals require an inordinately large number of past observations (e.g. $N = 1000$ in one well-known case) to achieve a global risk as small as the Bayes risk for certain non-Bayes procedures. Also, a substantial amount of attention is given to proving the existence of a.o. EB d.f.'s without actually giving a constructive procedure for explicitly finding the function. Even when constructive procedures are proposed (Deely and Kruse [11] and Maritz [23]), they are often difficult to apply since they involve linear programming and minimization of integrals, problems which have to be submitted to a computer to obtain even the simplest of results. In general, it can be fairly stated that very little attention has been given to "small sample" properties of the decision functions (i.e., small numbers of previous samples) although some recent work appears aimed in that direction.

A major purpose of this paper is to exhibit empirical Bayes procedures in a number of problems that often occur in theory and practice which are:

- (1) asymptotically optimal (if possible)
- (11) easy to calculate and
- (111) better than the "usual" d.f. for reasonably small values of N

(in some cases, $N=1$ suffices).

The two types of prior information discussed above, i.e., knowledge of G and past observations on X , can be thought of loosely as "partial" or "incomplete" prior information. The terms partial and incomplete indicate that we implicitly assume G to contain more than one element and that the past observations do not allow us to specify with certainty the "correct" G in G . The following specifications of the class G will be considered both without prior observations (in which case G -minimax d.f.'s are obtained) and with prior observations (in which case empirical Bayes decision functions are obtained):

(i) G is the conjugate class of priors. That is, $g(\lambda) \propto k(x|\lambda)$ where $k(x|\lambda)$ is the kernel of the likelihood of x given λ and $g(\lambda)$ is the prior density corresponding to G .

1.4 Problems Considered

of these decision functions.

Monte Carlo simulation is used to study the small sample properties of the average risk for small N and either numerical integration or small N . In some cases, it is difficult to determine analytically. Also, a comparison is made of various proposed EB d.f.'s for global risks even for small N .

of restriction on G , then one can obtain a substantial reduction. It is shown below that if one is willing to assume a reasonable amount

above are available then the methods of this paper can be used to obtain a confidence interval that closely approximates a Bayes confidence interval. In the empirical Bayes case, we find "asymptotically

but the types of partial prior information concerning θ discussed confidence interval for λ (given x). If g is not completely specified

having observed x_{N+1} . $I_{\theta}(x_{N+1})$ is then called a 100% Bayesian where $0 \leq \beta \leq 1$ and $g(\lambda|x_{N+1})$ is the posterior distribution of λ

$I_{\theta}(x_{N+1})$ of λ depending on x_{N+1} and β such that $\int_{\beta} g(\lambda|x_{N+1}) d\lambda = \beta$

Zimmer [12]. The problem here simply involves finding an interval

One particular case of this type is considered in Deely and

most cases) than the ordinary non-Bayesian confidence interval.

information can be used to find shorter confidence intervals (in

val for the parameter λ_{N+1} is considered. It is shown that prior

In addition, the problem of finding a Bayesian confidence inter-

estimation of the parameter λ (see section 2.2 for greater detail).

$L(\delta(x), \lambda) = \omega(\lambda)(\delta(x) - \lambda)^2$ where $\omega(\lambda) > 0$ $\forall \lambda \in \Lambda$ corresponding to

to other distributions. The loss function to be considered here is

although the general procedure outlined here can be easily applied

The binomial and normal processes will be of primary concern

(iv) Some combination of the above cases.

$$(1.1) \quad \int_{\lambda} \lambda dG(\lambda) = \mu \text{ is known.}$$

$$(1.2) \quad \int_{\lambda} \lambda^2 dG(\lambda) = \mu^2 \text{ is known.}$$

optimal" confidence intervals, that is, confidence intervals whose end points converge to the Bayes confidence interval end points.

CHAPTER II

PRELIMINARY RESULTS, CONCEPTS, AND STRUCTURAL ELEMENTS

2.1 General Structural Elements

The structural elements of the general Bayesian statistical decision problem are used throughout this paper in the following notation:

(1) Λ : a nonempty parameter space. λ represents the possible "states of nature" or "states of the world" and some $\lambda^0 \in \Lambda$ is the parameter value in effect at the time of experimentation.

(11) A : a nonempty "action" or "decision" space. Each $a \in A$ represents an action or decision available to the statistician which may be chosen as the result of some experiment.

(111) $L: A \times \Lambda \rightarrow R^+$; a non-negative loss function. $L(a, \lambda)$ represents the loss incurred in taking action a when the true state of nature is λ .

In passing, it should be noted that the triplet (Λ, A, L) is the usual definition of a mathematical game. Nature chooses a $\lambda \in \Lambda$, the statistician chooses an $a \in A$ and thereby incurs a loss $L(a, \lambda)$.

(iv) X : an observable random variable (or vector) associated with some experiment. X takes values in a sample space $\mathcal{X}$ on which a σ -finite measure μ is defined. When λ_0 is the true parameter value then X has the probability distribution function $F(x|\lambda_0)$ and has p.d.f. $f(x|\lambda_0)$ with respect to the measure μ .

(v) $\delta: \mathcal{X} \rightarrow \mathcal{A}$: a non-randomized decision function. $\delta(x)$ represents the action taken after observing a particular outcome x of the random variable X . The space of all d.f.'s of this type is denoted by $\mathcal{D}$.

Often, it is necessary to consider randomized d.f.'s, with randomization taken either over the space $\mathcal{D}$ before experimentation or over the action space $\mathcal{A}$ after experimentation (the so-called behavioral decision functions) but in the Bayesian decision problem, as shown below, it is only necessary to consider the space $\mathcal{D}$.

(vi) G : a "prior" probability distribution function of λ on Λ . G may be specified explicitly or only partially.

Of course, the assumption of (vi) is what distinguishes the Bayesian formulation from the non-Bayesian formulation of the problem. Throughout this paper, assumption (vi) is assumed to hold.

The problem is to choose a good (in some sense) decision function δ such that when x is observed action $\delta(x)$ is taken with a consequent loss of $L(\delta(x), \lambda)$. In order to find good decision functions criteria must be established to induce an ordering on the space $\mathcal{D}$. The following notation and terminology for the usual criteria are used throughout this paper:

(1) The risk or average loss in using a d.f. δ when the true parameter value is λ is:

$$R(\delta, \lambda) = E \left[L(\delta(X), \lambda) \right]$$

(2.1.1)

$$= \int_X L(\delta(x), \lambda) f(x|\lambda) dx$$

In this paper, the problems considered are such that $R(\delta, \lambda) < \infty$ A $\delta \in D, \lambda \in \Lambda$

(ii) The Bayes or average risk in using a d.f. δ when G is the prior distribution of λ is:

$$R(\delta, G) = \int_{\Lambda} R(\delta, \lambda) dG(\lambda)$$

(2.1.2)

$$= \int_{\Lambda} \int_X L(\delta(x), \lambda) f(x|\lambda) dx dG(\lambda)$$

(iii) the Bayes envelope function is

$$R(G) = \inf_{\delta} R(\delta, G)$$

(2.1.3)

2.2 Bayes Decision Functions

The desired object in the decision problem outlined above is to find a d.f. δ that has Bayes risk $R(\delta, G)$ as small as possible.

For a given G , a d.f. δ_1 is preferable to δ_2 iff $R(\delta_1, G) < R(\delta_2, G)$ and δ_1 is equivalent to δ_2 iff $R(\delta_1, G) = R(\delta_2, G)$. Recall that the ordinary definition of risk equivalence states that δ_1 and δ_2 are

risk equivalent iff $R(\delta_1, \lambda) = R(\delta_2, \lambda) \forall \lambda \in \Lambda$. Hence risk equivalence implies equivalence but the converse is not true.

In the problems considered in this paper there always exists a best d.f. δ_G ; i.e., there exists a δ_G such that

$$R(\delta_G, G) = R(G) = \inf_{\delta} R(\delta, G) \quad (2.2.1)$$

Such a d.f. δ_G is called a Bayes decision function. Since it is not true in general that δ_G always exists it is sometimes necessary to consider a δ^* such that

$$R(\delta^*, G) \leq \inf_{\delta \in D} R(\delta, G) + \epsilon \quad \text{for } \epsilon > 0. \quad (2.2.2)$$

Such a d.f. is called ϵ -Bayes and if δ^* is ϵ -Bayes for all $\epsilon > 0$ then

δ^* is called an extended Bayes decision function.

An important implication of the existence of a Bayes (or ϵ -Bayes) d.f. with respect to G is the well-known fact (see Ferguson [13], p. 43) that there always exists a non-randomized Bayes d.f. so we may restrict our attention to D .

As described earlier, the process of obtaining an observation x consists of first choosing a λ by the distribution $G(\lambda)$ followed by an x from $F(x|\lambda)$. On the other hand, the joint distribution of λ and x can be determined by first choosing x from

$$F_G(x) = \int_{\Lambda} F(x|\lambda) dG(\lambda) \quad (2.2.3)$$

and then choosing a λ from $G(\lambda|x)$, the conditional or posterior distribution of λ given $X=x$. Hence, the Bayes risk may be written (assuming the validity of the operation of interchanging integrals):

$$R(\delta, G) = \int_{\Lambda} \int_X L(\delta(x), \lambda) dG(\lambda|x) dF_G(x) \quad (2.2.4)$$

The Bayes decision function δ_G may then be found by minimizing

$$\int_{\Lambda} L(\delta(x), \lambda) dG(\lambda|x) \quad (2.2.5)$$

for each x .

$$L(\delta(x), \lambda) = w(\lambda) (\delta(x) - \lambda)^2 \quad (2.2.7)$$

in this paper is the following:

The loss function used for the estimation problems considered

$\lambda_0 \in \Lambda$.

δ such that $R(\delta, \lambda) \leq R(\delta^*, \lambda)$, $\forall \lambda \in \Lambda$, and $R(\delta, \lambda_0) < R(\delta^*, \lambda_0)$ for some

equivalence) then δ^* is admissible. (i.e., there does not exist a

If the Bayes d.f. δ^* with respect to G^* is unique (up to risk

Theorem 1

later chapters:

The following theorem (see Ferguson [13], p. 60) is useful in

$$\begin{aligned} \frac{\int_{\Lambda} f(x|\lambda) g(\lambda) d\lambda}{\int_{\Lambda} L(\delta(x), \lambda) f(x|\lambda) g(\lambda) d\lambda} &= \int_{\Lambda} L(\delta(x), \lambda) g(\lambda) d\lambda \\ &= \int_{\Lambda} L(\delta(x), \lambda) g(\lambda) d\lambda \frac{f_G(x)}{f(x|\lambda) g(\lambda) d\lambda} \end{aligned} \quad (2.2.6)$$

We can write (2.2.5) as:

In particular, if $w(\lambda) = c > 0$, $\forall \lambda \in \Lambda$, then

$$\delta_G(x) = \left\{ \int_{\Lambda} w(\lambda) g(\lambda | x) d\lambda \right\} \left\{ \int_{\Lambda} w(\lambda) g(\lambda | x) d\lambda \right\}^{-1} \quad (2.2.11)$$

It is easily seen that (2.2.10) is minimized by that $\delta_G(x)$ such that

$$\int_{\Lambda} w(\lambda) (\delta(x) - \lambda)^2 g(\lambda | x) d\lambda \quad (2.2.10)$$

or, dividing by $F_G(x)$, by minimizing

$$\int_{\Lambda} w(\lambda) (\delta(x) - \lambda)^2 F(x | \lambda) g(\lambda) d\lambda \quad (2.2.9)$$

Equivalently, we can minimize (for each x):

$$R(\delta, G) = \int_{\Lambda} \int_X w(\lambda) (\delta(x) - \lambda)^2 F(x | \lambda) g(\lambda) dx d\lambda = \int_{\Lambda} \int_X w(\lambda) (\delta(x) - \lambda)^2 F(x | \lambda) g(\lambda) dx d\lambda \quad (2.2.8)$$

this case, it is desired to find a decision function δ that minimizes

where $w(\lambda) > 0$, $\forall \lambda \in \Lambda$ and Λ is some subset of the real line. In

where K is a function of $t(\bar{x}_1)$ and λ_1 , $p(\bar{x}_1)$ is independent of λ , the

$$f(\bar{x}_1|\lambda) = K(t(\bar{x}_1, \lambda)p(\bar{x}_1)) \quad (2.3.1)$$

that is, by the Neyman factorization theorem we can write
exists a sufficient statistic $t(\bar{x}_1)$ of fixed dimensionality s for λ ;
In many cases (in all cases considered in this paper), there
situation some type of empirical Bayes procedure may be applicable.
independent observations $\left\{x_{1j}\right\}_{j=1}^{M_1}$ chosen randomly from $f(x|\lambda_1)$. In this
"past" observations $\bar{x}_1, \dots, \bar{x}_N$, each $\bar{x}_1$ itself consisting of M_1 independent
specified but that there exists a sequence of independent "prior" or
is completely specified. Suppose, however, that G is not completely
the minimum possible Bayes risk $R(\delta^G, G)$ can only be attained if G
ing section can only be found if G is completely specified. Hence
A Bayes decision function δ^G of the type described in the preceding

2.3 Empirical Bayes Decision Functions

this loss function is called an "estimator" (of λ).
to as point estimation problems and a decision function obtained using
(2.2.7) is the loss function used most commonly in problems referred
mean of the posterior distribution of λ given x . The loss function
That is, when $w(\lambda)$ is a positive constant the Bayes d.f. is just the

$$\delta^G(x) = \int \lambda g(\lambda|x) d\lambda = E_{\lambda|x} \quad (2.2.12)$$

The global risk $R(\delta_N, G)$ of δ_N is defined by

which is a random variable since it is a function of $t_1, \dots, t_N$.

$$R_N(\delta_N, G) = R(\delta_N, G) \quad (2.3.4)$$

of δ_N for a given set of past observations is just

ical Bayes d.f. is of the form $\delta_N(t_{N+1}, t_1, \dots, t_N)$ and the Bayes risk

unobservable. The unknown λ_{N+1} is the parameter of interest. An empir-

realization (t_{N+1}, λ_{N+1}) where the t_i 's are observable and the λ_i 's

the random variables $\{(t_i, \lambda_i)\}_{i=1}^N$ in addition to the "current" or "present"

Hence there exist N pairs of past realizations $\{(t_i, \lambda_i)\}_{i=1}^N$ of

$$F_G(t) = \int_0^V F(t|\lambda) dG(\lambda) \quad (2.3.3)$$

from the distribution with p.d.f.

Note that $\{t_i\}_{i=1}^N$ may be considered a random sample of size N

$F(t|\lambda_i)$ as a functional notation for the conditional density function.

distribution of T given $\lambda = \lambda_i$. For convenience, we continue to use

independent observations $t_1, \dots, t_N$ where each t_i is from the conditional

Hence all pertinent past (prior) information is contained in the

$$t(\bar{x}_i) = t(x_{i1}, \dots, x_{iM}) = (t_1, \dots, t_s). \quad (2.3.2)$$

dimensionality s of t does not depend on M ($M \geq s$), and

But

$$(2.3.9) \quad \begin{cases} F_G(1) = \mu \\ F_G(0) = 1 - \mu \end{cases}$$

i.e.,

$$(2.3.8) \quad F_G(t) = \int_0^1 \lambda^t (1-\lambda)^{1-t} dG(\lambda) \quad t = 0, 1$$

i.e., an asymptotically optimal EB d.f. This is not always possible. For example, if $M = 1$ and $P(T = 0 | \lambda) = 1 - \lambda$ and $P(T = 1 | \lambda) = \lambda$ then

$$(2.3.7) \quad \lim_{N \rightarrow \infty} R(\delta_N, G) = R(\delta_G, G) \quad \forall G \in \mathcal{G}$$

since δ_G is the Bayes decision function. However, it is hoped to find a sequence $\{\delta_N\}$ such that

$$(2.3.6) \quad R(\delta_N, G) \geq R(\delta_G, G)$$

where E_N represents expectation with respect to the N random variables $T_1, \dots, T_N$. Note that for all N ,

$$(2.3.5) \quad R(\delta_N, G) = E_N \left\{ R_N(\delta_N, G) \right\}$$

Now, any EB procedure depends on $f_G(t)$ for information concerning G (or any function of G). Hence, as pointed out by Robbins [32], since $f_G(t)$ depends only on n and both $\delta_G(t)$ and $R(\delta_G, G)$ depend on n and n_2 ,

$$R(\delta_G, G) = \frac{n(1-n)}{(n-n_2)(n_2-n_2)} \quad (2.3.13)$$

and

$$\delta_G(t) = \frac{n(1-n)}{t(n_2-n_2) + n(n-n_2)} \quad (2.3.12)$$

Then, for any $G \in \mathcal{G}$,

$$\begin{aligned} R(\delta, G) &= \int_1^0 \left\{ (1-\lambda) (\lambda - \delta(0))^2 + \lambda \left(\lambda - \delta(1) \right)^2 \right\} dG(\lambda) \\ &= \delta^2(0) + n[\delta^2(1) - 2\delta(0) - \delta^2(0)] + n^2[1 - 2\delta(1) + 2\delta(0)] \\ &= \delta^2(0) + n[\delta^2(1) - 2\delta(0) - \delta^2(0)] + n^2[1 - 2\delta(1) + 2\delta(0)] \end{aligned} \quad (2.3.11)$$

and for squared error loss

$$R(\delta, G) = \int_1^0 R(\delta, \lambda) dG(\lambda) \quad (2.3.10)$$

then, in general, no asymptotically optimal EB decision function exists in this case.

A common computation in working with empirical Bayes decision functions when using a squared error loss function is the following (recall that $\delta_G^*(t) = E_{\lambda|t}(\lambda)$)

$$\begin{aligned}
 E_N(R(\delta_N^*, G)) &= E_N \left\{ R_N(\delta_N^*, G) \right\} \\
 &= E_N \left\{ E_{\lambda, t}(\delta_N^*(t) - \lambda)^2 \right\} \\
 &= E_N \left\{ E_{\lambda, t}(\delta_N^*(t) - \delta_G^*(t))^2 + \delta_G^*(t) + \delta_G^*(t) - \lambda)^2 \right\} \\
 &= E_N \left\{ E_{\lambda, t}(\delta_N^*(t) - \delta_G^*(t))^2 + E_{\lambda, t}(\delta_G^*(t) - \lambda)^2 \right\} \\
 &\quad + 2E_{\lambda, t}(\delta_N^*(t) - \delta_G^*(t))(\delta_G^*(t) - \lambda) \Big\} \\
 &= E_{N+1}(\delta_N^*(t) - \delta_G^*(t))^2 + R(\delta_G^*, G)
 \end{aligned}
 \tag{2.3.14}$$

where E_{N+1} represents expectation with respect to the $N+1$ independent random variables $T_1, \dots, T_{N+1}$ each with p.d.f. $f_G(t)$, and $E_{\lambda, t}$ and E_t represent expectation with respect to (respectively) the joint distribution of λ and T , the conditional distribution of λ given $T = t$ and the unconditional distribution of T . Often, when no confusion arises, the symbol "E" (without subscripts) is used for expectations with respect to all random variables involved. Also, as in the above, the subscripts on the random variables T_i (and on the t_i 's) are often omitted when it is clear which variables are being considered.

It will be seen below that the structure of $\delta_N(t)$ often makes evaluation of $R(\delta_N, G)$ difficult. In these cases it will be necessary to use numerical integration or Monte Carlo studies to approximate the value $R(\delta_N, G)$.

Of particular interest is the case in which G is assumed specified except for some unknown parameters. This case is of interest because it is precisely this assumption which is made in many Bayesian-type analyses. However, the usual procedure involved in determining the unknown parameters is of a subjective nature. A typical method of this type is to question the "assessor" in order to determine his prior beliefs (cf. Winkler [49]) with no direct reference to prior observations, if any. Hence a method of directly estimating the parameters of G from past data, assuming reasonable small sample accuracy, is of more than academic interest.

In general, if the form of G is specified except for a finite number of unknown parameters α then it should be possible to devise estimators of α from the fact that $t_1, \dots, t_{N+1}$ may be considered a random sample of size $N+1$ from $f_G(t)$. This could be done, for example, by the well-known method of maximum likelihood or perhaps by the method of moments. In this paper, when the form of G is known, the empirical Bayes procedures proposed can be shown to be essentially equivalent to the method of moments. These procedures are used, even though the method of moments may not have the best asymptotic properties, because they are much simpler than other procedures (such as maximum likelihood) and appear accurate enough for our purposes.

$$\sup_{G \in \mathcal{G}} R(\delta^*, G) = \inf_{\delta} \sup_{G \in \mathcal{G}} R(\delta, G). \quad (2.4.3)$$

as: δ^* is $\mathcal{G}$ -minimax if
 identical to the Bayes d.f. The definition may be equivalently stated
 contains only one element G then the $\mathcal{G}$ -minimax decision function is
 decision function by Blum and Rosenblatt [3]. In particular, if $\mathcal{G}$
 for all $\delta \in \mathcal{D}$. Such a d.f. δ^* is called a $\mathcal{G}$ -minimax (or "minimax in $\mathcal{G}$ ")

$$\sup_{G \in \mathcal{G}} R(\delta^*, G) < \sup_{G \in \mathcal{G}} R(\delta, G) \quad (2.4.2)$$

Hence it would be desirable to find a d.f. δ^* ,

$$\sup_{G \in \mathcal{G}} R(\delta_1, G) \leq \sup_{G \in \mathcal{G}} R(\delta_2, G). \quad (2.4.1)$$

to d.f. δ_2 if
 In this case, it is reasonable (see Robbins [32]) to prefer a d.f. δ_1
 are available but the class $\mathcal{G}$ of which G is a member is specified.
 Assume now that G is unknown and no past observations from $F_G(x)$

2.4 $\mathcal{G}$ -Minimax Decision Functions

sample properties as well.
 shown to be asymptotically optimal while possessing good small (previous)
 In particular, the decision functions based on this method are

We also define the G -minimax value $\underline{V}(G)$ by writing

$$\underline{V}(G) = \inf_{\delta} \sup_G R(\delta, G). \quad (2.4.4)$$

This definition is analogous to the usual definition of the minimax

value $\underline{V}$ since

$$\underline{V} = \inf_{\delta} \sup_{\lambda} R(\delta, \lambda) \quad (2.4.5)$$

In fact, if G contains all degenerate distributions on Λ then $\underline{V} = \underline{V}(G)$. It is useful later to use the fact that a d.f. δ^* is G -minimax

iff

$$R(\delta^*, G) \leq \sup_{G \in G} R(\delta, G) \quad (2.4.6)$$

for all $\delta \in D$ and all $G \in G$. Recall that a d.f. δ_0 is called minimax iff

$$\sup_{\lambda \in \Lambda} R(\delta_0, \lambda) \leq \sup_{\lambda \in \Lambda} R(\delta, \lambda) \quad (2.4.7)$$

or equivalently, iff for all $\delta \in D$,

$$\sup_{\lambda \in \Lambda} R(\delta_0, \lambda) = \inf_{\delta} \sup_{\lambda \in \Lambda} R(\delta, \lambda). \quad (2.4.8)$$

$$R(\delta, g_0) > R(\delta^*, g_0) \text{ for some } g_0 \in G. \quad (2.4.13)$$

and

$$R(\delta, g) < R(\delta^*, g) \text{ for all } g \in G \quad (2.4.12)$$

A d.f. δ^* is G-admissible if there does not exist δ such that

Definition

nary admissibility (see section 2.2).

The concept of G-admissibility is defined analogously to ordi-

so that in this case a d.f. δ^* is G-minimax iff it is minimax.

$$\sup_{g \in G} R(\delta, g) = \sup_{\lambda \in \Lambda} R(\delta, \lambda) \quad (2.4.11)$$

and, of course, if G contains all degenerate distributions on Λ then

$$\sup_{g \in G} R(\delta, g) \leq \sup_{\lambda \in \Lambda} R(\delta, \lambda). \quad (2.4.10)$$

for all $g \in G$. Hence,

$$= \sup_{\lambda} R(\delta, \lambda)$$

$$R(\delta, g) = \int_{\Lambda} R(\delta, \lambda) dG(\lambda) \leq \int_{\Lambda} \sup_{\lambda} R(\delta, \lambda) dG(\lambda) \quad (2.4.9)$$

Now, for any $\delta \in D$ we have

G^* is least favorable in G if

Definition

is called a least favorable distribution.

In game theory terminology, a G -minimax strategy for nature respect to G is unique (up to equivalence) then δ^* is G -admissible. As a corollary to theorem 1, we add: if the Bayes d.f. δ^* with Hence G -admissibility implies admissibility in this case. a contradiction to the assumption of the G -admissibility of δ^* .

$$R(\delta_0, G) \leq R(\delta^*, G) \quad \forall G \in G \quad (2.4.17)$$

and by (2.4.15),

$$R(\delta_0, G_{\lambda_0}) < R(\delta^*, G_{\lambda_0}) \quad (2.4.16)$$

Then there exists $G_{\lambda_0} \in G$ such that

$$R(\delta_0, \lambda) \leq R(\delta^*, \lambda) \quad \forall \lambda \in \Lambda. \quad (2.4.15)$$

and

$$R(\delta_0, \lambda_0) < R(\delta^*, \lambda_0) \quad (2.4.14)$$

Then there exists $\delta_0 \in D$ and $\lambda_0 \in \Lambda$ such that

To see this, assume δ^* is G -admissible but not admissible.

there exists $G_{\lambda} \in G$ such that $G_{\lambda}(\lambda) - G_{\lambda}(\lambda^-) > 0$.

needed for G -admissibility to imply admissibility is that $\forall \lambda \in \Lambda$ admissibility implies admissibility. In fact, all that is really Also, if G contains all degenerate distributions on Λ then G - Obviously, admissibility implies G -admissibility for any class G .

concepts and ordinary minimax concepts it is not surprising that many

In view of the similarities of the definitions of G -minimax

consists of a single member G .

This definition is identical to G equivalence if the class

$$\sup_{G \in G} R(\delta_1, G) = \sup_{G \in G} R(\delta_2, G). \quad (2.4.22)$$

A d.f. δ_1 is said to be G -equivalent to δ_2 if

Definition

The following definition will also be needed:

the ordinary maximin value.

Again, if G is the class of all distributions on Λ then $\bar{V}(G) = \bar{V}$,

$$\bar{V}(G) = \sup_{\delta} \inf_{G \in G} R(\delta, G). \quad (2.4.21)$$

The G -maximin value $\bar{V}(G)$ is defined by:

over G at G^* .

Favorable in G if the Bayes envelope function attains its supremum

Expression (2.4.20) can be expressed in words as: " G^* is least

$$R(G^*) \geq R(G). \quad (2.4.20)$$

or, equivalently, if

$$\inf_{\delta} R(\delta, G^*) = \sup_{G \in G} \inf_{\delta} R(\delta, G) \quad (2.4.19)$$

for all $G \in G$ or, equivalently, if

$$\inf_{\delta} R(\delta, G^*) \geq \inf_{\delta} R(\delta, G) \quad (2.4.18)$$

$$R(\delta_0, G) \leq R(\delta_0, G_0) \quad (2.4.24)$$

If δ_0 is Bayes with respect to $G_0 \in G$ and for all $G \in G$

Theorem 3

and are invaluable in finding G -minimax decision functions in practice.

of minimax theorems 1, 2 and 3 found in Ferguson [13], pages 90-91,

The following theorems (3, 4 and 5) are exact G -minimax analogues

q.e.d.

equivalent to δ^* this contradicts the uniqueness of δ^* .

That is, δ_0 is also G -minimax. But since δ_0 is not

$$\sup_G R(\delta_0, G) \leq \sup_G R(\delta^*, G). \quad (2.4.23)$$

then

$$R(\delta_0, G_0) < R(\delta^*, G_0) \text{ for some } G_0 \in G$$

$$R(\delta_0, G) \leq R(\delta^*, G) \quad \forall G \in G$$

that

Suppose δ^* is not G -admissible. Then there exists δ_0 such

Proof:

function then δ^* is G -admissible.

If δ^* is a unique (up to equivalence) G -minimax decision

Theorem 2

prove and are useful in the later sections:

Following G -minimax analogues to certain minimax results are easy to

modification to G -minimax decision functions. In particular, the

well-known results on minimax decision functions carry over with little

then

$$\underline{V}(G) = \overline{V}(G), \quad (2.4.25)$$

G_0 is G -minimax and G_0 is least favorable.

Proof

Obviously $\overline{V}(G) \leq \underline{V}(G)$. Also,

$$\underline{V}(G) = \inf_{\delta} \sup_{G_0} R(\delta, G) \quad (2.4.26)$$

$$\leq \sup_{G_0} R(\delta_0, G)$$

$$\leq R(\delta_0, G_0)$$

$$= \inf_{\delta} R(\delta, G_0)$$

$$\leq \sup_{\delta} \inf_{G_0} R(\delta, G)$$

$$= \overline{V}(G).$$

Hence, $\underline{V}(G) = \overline{V}(G)$. Next, from (2.4.24),

$$\sup_{G_0} R(\delta_0, G) \leq \sup_{\delta} \inf_{G_0} R(\delta, G) \quad (2.4.27)$$

$$\leq \sup_{G_0} R(\delta, G) \text{ wded}$$

$$\Rightarrow \delta_0 \text{ is } G\text{-minimax}$$

Finally, also from (2.4.24),

$$\inf_{\delta} R(\delta, G_0) \geq \inf_{\delta} \sup_{G_0} R(\delta, G) \quad (2.4.28)$$

$$\geq \inf_{\delta} R(\delta, G) \text{ wgeg}$$

$\Rightarrow G_0$ is least favorable.
q.e.d.

Theorem 4

If δ_N is Bayes with respect to $G_N \in G$ and

$$\lim_{N \rightarrow \infty} R(\delta_N, G_N) = C \quad (2.4.29)$$

and there exists δ_0 such that

$$R(\delta_0, G) < C \quad \forall G \in G \quad (2.4.30)$$

then

$$\underline{V}(G) = \bar{V}(G)$$

and δ_0 is G -minimax.

Proof $\forall \delta$,

$$\sup_{G \in G} R(\delta, G) \geq R(\delta, G_N) \geq R(\delta_N, G_N) \quad \forall N \quad (2.4.31)$$

$$\Rightarrow \sup_{G \in G} R(\delta, G) \geq C \geq R(\delta_0, G) \quad \forall G$$

$$\Rightarrow \sup_{G \in G} R(\delta, G) \geq R(\delta_0, G) \quad \forall G$$

$$\Rightarrow \sup_{G \in G} R(\delta, G) \geq \sup_{G \in G} R(\delta_0, G)$$

therefore,

$$\delta_0 \text{ is } G\text{-minimax.} \\ \underline{\text{q.e.d.}}$$

Theorem 5

If δ^* is Bayes (or extended Bayes) with respect to $G^* \in G$ and

where δ^* is a G -minimax decision function. Robbins [31] calls such an estimator asymptotically subminimax but to be consistent with the

$$\lim_{N \rightarrow \infty} R(\delta_N^*, G) = R(\delta^*, G) \quad (2.4.33)$$

with the property:

least for small N) to find a decision function $\delta_N^*(t_{N+1}, t_1, \dots, t_N)$ $R(\delta_N^*, G)$ converges too slowly to $R(\delta^*, G)$ then it may be desirable (at as defined in section 2.3 does not exist or an a.o. $EB \delta_N^*$ exists but If N prior observations on $F_G(t)$ exist but an a.o. $EB \delta_N^*$ exists but some $G_0 \in G$.

$G \in G$ and check to see if δ_0^* is Bayes (or extended Bayes) for (ii) Using theorem 5: Find a d.f. δ_0^* such that $R(\delta_0^*, G) = K$ for all $R(\delta_0^*, G_0) \geq R(\delta_0^*, G) \quad \forall G \in G$.

G_0 , obtain the Bayes d.f. δ_0^* for G_0 and check to see if (i) Using theorem 3: Find a supposed least favorable distribution find G -minimax decision functions in this paper. Namely:

The above theorems provide the two principal methods used to

q.e.d.

By theorem 3, δ^* is G -minimax and G^* least favorable.

$$R(\delta^*, G) = K = R(\delta^*, G^*) \quad \forall G \in G$$

Proof:

favorable in G .

where K is some constant then δ^* is G -minimax and G^* is least

$$R(\delta^*, G) = K \quad \forall G \in G \quad (2.4.32)$$

present terminology a decision function with property (2.4.33) is called asymptotically G -minimax or asymptotically minimax in G in this paper.

Chapter 4). The current observations are denoted by $\left\{x_{N+1,j} \mid j=1, \dots, M\right\}$ and the past observations by $\left\{x_{1,j} \mid j=1, \dots, M\right\}$ where

which the prior distribution G is a member (as will be discussed in observations on $f_G(x)$ or some knowledge concerning the class G of with the aid of the prior information at hand - whether it be past at the time of drawing the sample. The inference is always to be made inference concerning the unknown parameter value or values in effect the distribution with p.d.f. (3.1.1) and it is desired to make some where $-\infty < \lambda < \infty$ and $\sigma^2 > 0$. A random sample of size M is drawn from

$$(3.1.1) \quad f(x|\lambda, \sigma^2) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(x-\lambda)^2}{2\sigma^2}}$$

with p.d.f. as a generator of independent real-valued variables $x_1, x_2, \dots$, each The independent normal process as used in this paper is defined

3.1 Introduction

NORMAL PROCESS I: EMPIRICAL BAYES DECISION FUNCTIONS

CHAPTER III

$$F_G(t) = \frac{\sqrt{2\pi\sigma^2\left(\frac{M}{1} + \frac{n}{1}\right)}}{e^{-\frac{2\sigma^2\left(\frac{M}{1} + \frac{n}{1}\right)}{(t-\mu)^2}}}$$

(3.1.11)

When the prior G is conjugate, the compound density (3.1.6) becomes

$$\left. \begin{aligned} \mu &= m' \\ \mu_2 &= \frac{n'}{\sigma^2} \\ \mu_1^2 &= \frac{n'}{\sigma^2} + (m')^2 \end{aligned} \right\}$$

(3.1.10)

and

$$\tilde{\lambda} \sim N\left(m', \frac{n'}{\sigma^2}\right)$$

(3.1.9)

that is,

$$g(\lambda | m', n') = \frac{\sqrt{2\pi\sigma^2}}{\sqrt{n'}} e^{-\frac{n'}{2\sigma^2}(\lambda - m')^2}$$

(3.1.8)

The conjugate prior density in this case is:

$$\left. \begin{aligned} \int \lambda dg(\lambda) &= \mu \\ \int \lambda^2 dg(\lambda) &= \mu_2 \\ \mu_2 &= \mu_1^2 - \mu^2 \end{aligned} \right\}$$

(3.1.7)

and the following notation for the first few moments of G is used:

$$\delta_{\pi}(t_{N+1}) = t_{N+1} \quad (3.2.2)$$

denoted here by:

on λ . Now the usual (ML, UMVU) estimator of λ_{N+1} is simply t_{N+1} and is Let G^* denote the class of all probability distribution functions

$$L(\delta(t), \lambda) = (\delta(t) - \lambda)^2 \quad (3.2.1)$$

The loss function to be used when estimating λ_{N+1} is:

3.2 Estimation of λ_{N+1}

$$\lambda | t_{N+1} \sim N\left(m'', \sigma_2^2\right) \quad (3.1.14)$$

that is,

$$m'' = \frac{1}{n''} (n' \mu + M t_{N+1}) = \frac{\sigma_2^2 \mu + M \mu}{\sigma_2^2 + M \mu^2}$$

where $n'' = n' + M$

$$g(\lambda | t_{N+1}) = \frac{\sqrt{\frac{1}{2\pi\sigma_2^2}}}{\sigma_2^2} e^{-\frac{n''(\lambda - m'')^2}{2\sigma_2^2}} \quad (3.1.13)$$

Also, the posterior density of λ given t_{N+1} becomes:

$$\pi \sim N\left(\mu, \sigma_2^2 \left[\frac{1}{M} + \frac{1}{n''}\right]\right) \quad (3.1.12)$$

that is,

limiting distribution as a member. When we speak of the class of conjugate priors we will include this distribution (see (3.1.13)) and since the Bayes d.f. for the loss function (3.2.1) is the mean of the posterior, then $\delta_{\mathbb{T}}(t_{N+1})$ is Bayes

$$(3.2.5) \quad \int_F g(\lambda|F) d\lambda = [u_{\mathbb{T}}(F)]^{-1} \int_F d\lambda = 1.$$

where $u_{\mathbb{T}}(F)$ is the Lebesgue measure of F . Hence,

$$(3.2.4) \quad g(\lambda|F) = [u_{\mathbb{T}}(F)]^{-1}$$

i.e., $\delta_{\mathbb{T}}$ has a constant risk for all $G \in \mathcal{G}^*$. Now, if G is in the class of conjugate priors we note that as $n' \rightarrow 0$ we have increasingly vague prior information (see (3.1.9)). In fact, we can consider the limiting distribution ($n' = 0$) to be a "uniform" distribution in the sense of Lindley [21]. That is, if F is any set of λ 's of finite measure then the distribution (conditional on $\lambda \in F$) of λ has density:

$$(3.2.3) \quad R(\delta_{\mathbb{T}}, G) = \int_{-\infty}^V \int_{-\infty}^V f(t|\lambda) {}_2F_1(t|\lambda) d\lambda dt = \int_{\mathbb{T}} \text{Var} \left| \frac{V}{\lambda} \right| dG(\lambda) = \frac{M}{\sigma^2} = \text{AGEG}^*$$

The Bayes risk of $\delta_{\mathbb{T}}$ is:

$$\bar{u} = \sum_{i=1}^{N+1} t_i / (N+1) \sim N\left(u, \frac{\sigma^2}{M(N+1)}\right) \sim N\left(u, \frac{\sigma^2}{M} \left[\frac{1}{N+1} + \frac{1}{M}\right]\right) \quad (3.2.10)$$

And since $\{t_i\}_{i=1}^N$ is a random sample from the distribution with p.d.f. (3.1.11).

$$\bar{u}^* = \sum_{i=1}^N t_i / N \sim N\left(u, \frac{\sigma^2}{MN}\right) \sim N\left(u, \frac{\sigma^2}{M} \left[\frac{1}{N} + \frac{1}{M}\right]\right) \quad (3.2.9)$$

The following facts are useful in devising an EB d.f.:

$$\frac{R(\delta_{G'}, G)}{R(\delta_{T'}, G)} = \frac{M}{M + n'} < 1. \quad (3.2.8)$$

The ratio of $R(\delta_{G'}, G)$ to $R(\delta_{T'}, G)$ is:

$$\begin{aligned} R(\delta_{G'}, G) &= \int_0^V t^\lambda |g(\delta_{G'}(t) - \lambda)^2 dg(\lambda) \\ &= \int_0^V \int_0^\infty t^\lambda |g(\delta_{G'}(t) - \lambda)^2 g(\lambda) | t F_G(t) d\lambda dt \\ &= \int_0^\infty \frac{M + n'}{\sigma^2} F_G(t) dt \\ &= \frac{M + n'}{\sigma^2}. \end{aligned} \quad (3.2.7)$$

The minimum Bayes risk in this case is:

variances. That is, δ_G is the weighted average of the prior mean u and the sample mean t_{N+1} with weights proportional to the inverses of their respective variances.

Define:

$$(3.2.11) \quad s_2^t = \sum_{i=1}^{N+1} (t_i - \bar{t})^2 / N.$$

Now,

$$(3.2.12.) \quad \begin{aligned} \bar{t} - \bar{t} &= t_i - \sum_{j=1}^{N+1} t_j / (N+1) \\ &= \left\{ \sum_{j=1, j \neq i}^{N+1} t_j - \sum_{j=1}^{N+1} t_j \right\} / (N+1) \end{aligned}$$

and since the t_i are all independently distributed $N(\bar{t}, \frac{M}{\sigma^2} + \mu_2)$,

$$(3.2.13) \quad \bar{t} - \bar{t} \sim N(0, \frac{N+1}{N} \left[\frac{M}{\sigma^2} + \mu_2 \right])$$

and hence

$$(3.2.14) \quad E(s_2^t) = \frac{M}{\sigma^2} + \mu_2.$$

Therefore we have shown that

$$(3.2.15) \quad E(s_2^t) - \frac{M}{\sigma^2} = \mu_2$$

and by the strong law of large numbers, as $N \rightarrow \infty$:

$$(3.2.16) \quad \begin{aligned} \bar{t} - \bar{t} &\xrightarrow{a.s.} \mu \\ s_2^t - \frac{M}{\sigma^2} &\xrightarrow{a.s.} \mu_2. \end{aligned}$$

$$\begin{aligned} \hat{\mu} &= \bar{t} \\ \hat{\sigma}_2^2 &= s_2^2 - \frac{t^2}{M} \end{aligned} \quad (3.2.20)$$

i.e.,

$$\begin{aligned} \hat{\mu} &= \bar{t} \\ \hat{\sigma}_2^2 + \frac{M}{\hat{\sigma}_2^2} &= s_2^2 \end{aligned} \quad (3.2.19)$$

and setting

$$\begin{aligned} E(t_G) &= \mu \\ \text{Var}(t_G) &= \sigma_2^2 \left(\frac{M}{1} + \frac{N}{1} \right) = \sigma_2^2 + \frac{M}{\mu_2} \end{aligned} \quad (3.2.18)$$

identical to the moment estimators obtained by noting
As pointed out in section 2.3, the estimators $\hat{\mu}$ and $\hat{\mu}_2$ are essentially

$$\hat{\mu}_2 = \max \left\{ 0, s_2^2 - \sigma_2^2/n \right\}. \quad (3.2.17)$$

$\hat{\mu}$ and variance = μ_2 where

As our estimator G_N of G we take a normal distribution with mean =

this last property is obviously true for large N .
have $R(\delta_N^*, G) \leq R(\delta_N^*, G)$ for reasonably small N . If δ_N^* is a.o. then
apply and (11) good for small N . In particular it is desirable to
d.f. with prior G_N , is: (1) asymptotically optimal (11) easy to
require G_N to be a member of G . We show that the d.f. δ_N^* , the Bayes
is to obtain an explicit estimator G_N of the unknown G where we
The general approach used here and in the following sections

If the random variables $\xi_N, \dots, \psi_N$ converge with probability one to the constants $x, y, \dots, r$ respectively then any rational function $R(\xi_N, \dots, \psi_N)$ converges with probability one provided the latter is finite.

Theorem 6

we obtain the following theorem:
20.6, Cramer [9], for the case of convergence with probability one
Now, by modifying the theorem and its proof given in section

$$R(\delta_N, G) = R(\delta_N, G) + E_{N+1} \{ \delta_N(t_{N+1}) - \delta_G(t_{N+1}) \}^2. \quad (3.2.25)$$

and, from (2.3.14), we know that the global risk is:

$$\delta_N(t_{N+1}) = \frac{\sigma^2 + M\hat{\mu}_2}{\sigma^2 + M\hat{\mu}_2 + M\hat{\mu}_2 t_{N+1}} \quad (3.2.24)$$

In any case, our empirical Bayes estimator is:

$$s^2_t > \frac{t}{M}. \quad (3.2.23)$$

then for large N we have with probability one

$$s^2_t - \frac{t}{M} \xrightarrow{\text{a.s.}} \mu_2 > 0 \quad (3.2.22)$$

But since

$$\delta_N(t_{N+1}) = \mu = \sum_{i=1}^{N+1} t_i / (N+1). \quad (3.2.21)$$

empirical Bayes d.f. is just

That is, G_N is a degenerate distribution at μ in this case and our
Note that from (3.2.17) we are estimating μ_2 to be 0 if $s^2_t < \frac{t}{M}$.

Now, from (3.2.16) and theorem 6 we have that

$$(3.2.26) \quad 0 \leftarrow \left(\frac{\sigma_2^2 + M_2}{\sigma_2^2 + M_2 + M_2 t_{N+1}} - \frac{\sigma_2^2 + M_2}{\sigma_2^2 + M_2 + M_2 t_{N+1}} \right) \leftarrow \text{a.s.}$$

i.e.,

$$(3.2.27) \quad \left(\delta_N(t_{N+1}) - \delta_G(t_{N+1}) \right) \leftarrow \text{a.s.} \quad 0.$$

Defining

$$(3.2.28) \quad \left\{ \begin{aligned} \alpha_1 &= \frac{\sigma_2^2 + M_2}{\sigma_2^2} \\ \alpha_2 &= \frac{M_2}{\sigma_2^2 + M_2} \\ \alpha_3 &= \frac{\sigma_2^2}{\sigma_2^2 (N+1) (\sigma_2^2 + M_2)} \\ \alpha_4 &= \frac{M_2^2 + \sigma_2^2}{M_2^2 + \sigma_2^2 (N+1)} \end{aligned} \right.$$

we see that

$$(3.2.29) \quad \left\{ \begin{aligned} 0 &\leq \alpha_i \leq 1 \quad i=1, \dots, 4 \\ \alpha_1 + \alpha_2 &= 1 \\ \alpha_3 + \alpha_4 &= 1 \end{aligned} \right.$$

and

$$(3.2.30) \quad \delta_G(t_{N+1}) = \alpha_1 u + \alpha_2 t_{N+1}$$

In particular, suppose that

Suppose next that additional information concerning G is known.

$G \in G$ and a comparison is made with the global risks for other estimators.

In section 3.3, $R(\delta_N^*, G)$ is evaluated for small N and for various

i.e., δ_N^* is asymptotically optimal in G .

$$\lim_{N \rightarrow \infty} E \left[\delta_N^*(t) - \delta_G(t) \right]^2 = 0 \quad (3.2.36)$$

Then, (3.2.27) and (3.2.35) imply

$$E \left[\delta_N^*(t) - \delta_G(t) \right]^2 < \infty. \quad (3.2.35)$$

Hence,

$$\left\{ \delta_N^*(t) - \delta_G(t) \right\}^2 \leq (\mu^*)^2 + \mu^2 + t^2 + |\mu^* \mu| + |\mu t| + |\mu^* t|. \quad (3.2.34)$$

$$-1 \leq \beta_3 \leq 1$$

$$-1 \leq \beta_2 \leq 0$$

$$\text{where } 0 \leq \beta_1 \leq 1$$

$$\sum_{i=1}^3 \beta_i = 0 \quad \text{and} \quad \text{so that}$$

$$\delta_N^*(t_{N+1}) - \delta_G(t_{N+1}) = \beta_1 \mu^* + \beta_2 \mu + \beta_3 t_{N+1} \quad (3.2.33)$$

we see that

$$\begin{cases} \beta_1 = \alpha_3 \\ \beta_2 = -\alpha_1 \\ \beta_3 = \alpha_4 - \alpha_2 \end{cases} \quad (3.2.32)$$

$$\text{where } \mu^* = \sum_{i=1}^N t_i / N. \quad \text{Also, defining}$$

$$\delta_N^*(t_{N+1}) = \alpha_3 \mu^* + \alpha_4 t_{N+1} \quad (3.2.31)$$

$$\lim_{N \rightarrow \infty} R(\delta_N', G) = R(\delta_{G'}, G). \quad (3.2.41)$$

Hence, obviously

$$\begin{aligned} &= \frac{\sigma^2}{2} \left[1 + \frac{M(N+1)}{n'} \right] \\ &= R(\delta_{G'}, G) + \frac{n' \sigma^2}{(N+1)(n'+M)M} \\ &R(\delta_N', G) = R(\delta_{G'}, G) + \left(\frac{n'+M}{n'} \right)^2 \frac{\sigma^2}{2} \left(\frac{M}{1} + \frac{n'}{1} \right) \end{aligned} \quad (3.2.40)$$

$$\begin{aligned} &\text{and since } \hat{\mu} \sim N \left(\mu, \frac{\sigma^2}{2} \left[\frac{M}{1} + \frac{n'}{1} \right] \right) \text{ by (3.2.10) then} \\ &= R(\delta_{G'}, G) + \left(\frac{n'+M}{n'} \right)^2 E_{N+1} (\hat{\mu} - \mu)^2 \\ &= R(\delta_{G'}, G) + E_{N+1} \left\{ \frac{n' \mu + M \hat{\mu}_{N+1}}{n' \hat{\mu} + M \hat{\mu}_{N+1}} - \frac{n' \mu + M}{n' \hat{\mu} + M \hat{\mu}_{N+1}} \right\}^2 \end{aligned}$$

$$R(\delta_N', G) = E_N R_N(\delta_N', G) \quad (3.2.39)$$

The global risk of (3.2.38) is (by (2.3.14))

$$\delta_N'(t_{N+1}) = \sigma^2 \frac{\mu + M \mu_2}{n' + M} = n' \frac{\mu + M \hat{\mu}_{N+1}}{\hat{\mu} + M \hat{\mu}_{N+1}} \quad (3.2.38)$$

is known and denote the class of all conjugate priors with variance μ_2 by G' . In this case, our estimator G_N of G will be a $N(\hat{\mu}, \mu_2)$ distribution and our EB d.f. will be

$$\mu_2 = \int_V (\lambda - \mu)^2 dG(\lambda) \quad (3.2.37)$$

$$\frac{R(\delta_N', G)}{N+4} = \frac{R(\delta_{\Pi}', G)}{4N+4} \quad (3.2.44)$$

this case,

possible improvement when μ_2 is known, suppose $M = 3$, $\mu_2 = \sigma^2/9$. In that δ_N' is asymptotically optimal. As a particular example of the can be deduced directly from (3.2.42) or simply from the knowledge

$$\lim_{N \rightarrow \infty} \frac{R(\delta_N', G)}{R(\delta_{\Pi}', G)} = \frac{M+n'}{M} = \frac{R(\delta_{\Pi}', G)}{R(\delta_G', G)} \quad (3.2.43)$$

of how large that variance may be. Of course, the fact that the prior variance enables us to improve on the estimator δ_{Π}' regardless than the usual non-Bayes estimator δ_{Π}' . In other words, knowledge of That is, for any $N \geq 1$, the empirical Bayes estimator δ_N' is better

$$\leq 1.$$

$$= \frac{M(N+1) + n'}{M(N+1) + n' + n'(N+1)}$$

$$= \frac{M(N+1) + n'}{M(N+1) + n' + n'(N+1)}$$

$$\frac{R(\delta_N', G)}{R(\delta_{\Pi}', G)} = \frac{M}{n' + M} \left[1 + \frac{M(N+1)}{n'} \right] \quad (3.2.42)$$

An important characteristic of δ_N' is deduced from the following:

That is, δ_N' is asymptotically optimal in G' .

where

$$\alpha'_1 \bar{z}' \sim N(\alpha'_1 \bar{z}, \alpha'_1 \bar{z} \alpha'_1) \quad (3.2.48)$$

and hence

$$\bar{z} \sim N(\mu, \bar{z}) \quad (3.2.47)$$

We also have:

$$\begin{aligned} \alpha_3 &= -1 \\ z_3 &= \lambda \\ \alpha_2 &= \frac{u_1 + M}{M} \\ z_2 &= t_{N+1} \\ \alpha_1 &= \frac{u_1}{u_1 + M} \\ z_1 &= \mu \end{aligned} \quad (3.2.46)$$

where

$$\begin{aligned} \bar{z}' &= \begin{pmatrix} z_3 \\ z_2 \\ z_1 \end{pmatrix} = (\alpha_1, \alpha_2, \alpha_3) \\ \delta_N^1(t_{N+1}) - \lambda &= \left(\frac{u_1}{u_1 + M} \right) \mu + \left(\frac{u_1 + M}{M} \right) t_{N+1} - \lambda \end{aligned} \quad (3.2.45)$$

We can write:

$$\{\delta_N^1(t) - \lambda\}^2 \text{ it is useful to know the exact distribution of } \delta_N^1(t) - \lambda.$$

In addition to computing the expectation of the random variable

obtained. See section 3.3 for more complete results.

so that even when $N=1$ a substantial reduction in the global risk is

and after some straight - forward manipulation it is easy to see that

$$\bar{\alpha}'_n \bar{z} = 0 \quad (3.2.51)$$

obviously,

$$\sigma_2^{n, \lambda} t_{N+1}^{n, \lambda} = \left(\frac{1}{\sigma_2^{N+1}} \right) \sigma_2^{n, \lambda} t_{N+1}^{n, \lambda} = \frac{\sigma_2^{N+1}}{\sigma_2^{M+n}} \frac{M!}{n!}$$

$$\sigma_2^{n, \lambda} = \sigma_2^{n, \lambda} = \frac{n!}{\sigma_2^2}$$

$$\sigma_2^{\lambda} = \sigma_2^{\lambda} = \frac{n!}{\sigma_2^2}$$

$$\sigma_2^{t_{N+1}^{n, \lambda}} = \sigma_2^{\left(\frac{1}{\sigma_2^{N+1}} + \frac{n!}{\sigma_2^{M+n}} \right)} = \frac{M!}{\sigma_2^{M+n}} \frac{n!}{\sigma_2^{M+n}}$$

$$\sigma_2^{\frac{n}{\sigma_2}} = \sigma_2^{\left(\frac{1}{\sigma_2^{N+1}} + \frac{n!}{\sigma_2^{M+n}} \right)} = \frac{M!}{\sigma_2^{M+n}} \frac{n!}{\sigma_2^{M+n}} \quad (3.2.50)$$

and the elements of $\bar{\Sigma}_z$ are simply

$$\bar{\Sigma}_z = \begin{pmatrix} \sigma_2^{n, \lambda} & \sigma_2^{t_{N+1}^{n, \lambda}} & \sigma_2^{\lambda} \\ \sigma_2^{t_{N+1}^{n, \lambda}} & \sigma_2^{t_{N+1}^{n, \lambda}} & \sigma_2^{t_{N+1}^{n, \lambda}} \\ \sigma_2^{\lambda} & \sigma_2^{t_{N+1}^{n, \lambda}} & \sigma_2^{\lambda} \end{pmatrix}$$

$$\bar{\Sigma}_z = (n, n, n) \quad (3.2.49)$$

interval on λ_{N+1} , namely:

which is obviously longer than the usual Bayesian $(1-\alpha)$ 100 % confidence

$$\delta_N^I(t) \pm z_{1-\alpha/2} \frac{\sqrt{M}}{\sigma} \quad (3.2.56)$$

λ_{N+1} is simply:

The usual non-Bayesian $1-\alpha$ level confidence interval on the unknown

$$\delta_N^I(t) - \lambda \sim N(0, 1) \cdot \frac{\sqrt{M}}{\sigma} \quad (3.2.55)$$

so that

$$\delta_N^I(t) - \lambda = t - \lambda \sim N(0, \frac{\sigma^2}{M}) \quad (3.2.54)$$

In a similar fashion we can easily see that

also, the random variable $\delta_N^I(t) - \lambda$ converges in distribution to a $N(0, \frac{\sigma^2}{n+m})$ which is the distribution of $\delta_G(t) - \lambda$. This should be expected of course, since $\delta_N^I(t)$ is asymptotically optimal.

$$\delta_N^I(t) - \lambda \sim N\left(0, \sigma^2 \frac{M(N+1) + n}{M(N+1)(M+n)}\right) \quad (3.2.53)$$

that

which is, of course, identical with (3.2.40). That is, we have shown

$$\bar{\sigma}_\alpha^2 = \sigma^2 \frac{M(N+1) + n}{M(N+1)(M+n)} \quad (3.2.52)$$

given.

In this section the results of several Monte Carlo type simulations relating to the decision functions described in the previous section are

3.3 Numerical Results

approach.

(3.2.59) was also found by Deely and Zimmer [12] using a different to the Bayes confidence interval (3.2.57). The confidence interval usual confidence interval (3.2.56) and also, as $N \rightarrow \infty$, it converges The confidence interval represented by (3.2.59) is shorter than the

$$\delta_N^0(t) + z_{1-\alpha/2} \sqrt{\frac{\sigma^2}{1 + \frac{M(N+1)}{n}}} \quad (3.2.59)$$

(3.2.58). The result is:

and hence an empirical Bayes confidence interval can be formed using

$$\frac{\delta_N^0(t) - \lambda}{\sqrt{\frac{\sigma^2}{M(N+1) + n}}} \sim N(0,1) \quad (3.2.58)$$

Now, from (3.2.53) we have

$$\delta_N^0(t) + z_{1-\alpha/2} \frac{\sqrt{M+n}}{\sigma} \quad (3.2.57)$$

Since any asymptotically optimal empirical Bayes decision function δ_N is nearly as good as the Bayes d.f. δ_G for large N , the primary concern here is the case where N is small. (e.g., $N=1, 10, 20$). Although asymptotically optimality is a desirable property for any decision function depending on past observations to possess, it is a minimal requirement in that it implies nothing about the behavior of δ_N for small N . If however, we could find a d.f. δ_N^* such that

$$R(\delta_N^*, G) \leq R(\delta_N, G) \quad \forall N \text{ and } \forall G \quad (3.3.1)$$

where δ_N^* is any other a.o. EB decision function, then one would feel justified in calling δ_N^* the "best" EB d.f. available. However, we are unlikely to be able to find such a δ_N^* except in very restrictive cases.

In practice, one may be willing to use δ_N^* if it can be shown that

$$R(\delta_N^*, G) \leq R(\delta_N, G) \quad \forall N \geq N_0 \text{ and } \forall G \quad (3.3.2)$$

where, as usual, δ_N^* is some optimal non-Bayes decision function. In any case, if a d.f. δ_N^* satisfies (3.3.2) for reasonably small N_0 then δ_N^* is at least worthy of further consideration.

The main purpose in this section is to closely approximate the global risk $R(\delta_N, G)$ for each proposed decision function given in section 3.2 where an exact determination is difficult or impossible by analytical means. Then these decision functions are compared (by means of $R(\delta_N, G)$ to each other and to the Bayes decision function δ_G).

The method used is first to generate a sample $t_1, \dots, t_N$ of size N , each t_i having the p.d.f. $f_G(t)$, then to generate an observation λ_{N+1} from the distribution with p.d.f. $G(\lambda)$ and finally to generate a t_{N+1} from the distribution with p.d.f. $f(t|\lambda_{N+1})$. Equivalently, we could have generated a random sample of size $N+1$ from the distribution with p.d.f. $G(\lambda)$, say $\left\{ \lambda_i \right\}_{i=1}^{N+1}$ and then generate an observation t_i from the distribution with p.d.f. $f(t|\lambda_i)$ for $i=1, 2, \dots, N+1$. The decision function value $\delta_N(t_{N+1})$ resulting from this process is then compared with the actual λ_{N+1} (unknown in practice) by means of the squared error loss function:

$$L(\delta_N(t), \lambda) = (\delta_N(t) - \lambda)^2 \quad (3.3.3)$$

and the entire process is then repeated. In this section, all runs involve 100 replications unless otherwise stated. The actual calculations were performed on an IBM 360 model 44. The programs used and more details can be found in the appendix. In addition to the decision functions considered in section 3.2, one other decision function is studied here since it exhibited an explicit method for construction of a sequence of distributions G_N which estimate the unknown G where G is assumed only to be the class of all continuous distribution functions on λ . This procedure is described by Deely and Kruse [11] and is outlined briefly below. Define G_N to be the class of discrete distributions on λ with

prior distributions:

For convenience we let $\sigma^2 = 1$ and consider the following two
The computer program used for this purpose is found in the appendix.
This optimal strategy may be found by linear programming techniques.

$$(3.3.8) \quad \begin{pmatrix} (a_{1j})_{2N \times N} \\ (-a_{1j})_{2N \times N} \end{pmatrix} \quad 4N \times N$$

strategy for the second player in a game with payoff matrix
Finding $G_N^*(\lambda)$ then can be shown to be equivalent to finding an optimal

$$(3.3.7) \quad a_{1j} = \begin{cases} F(t_{1j} | \lambda_{1N}^{jN}) - \frac{N}{(1-N)}, & 1 \leq j \leq N \\ F(t_{1-N} | \lambda_{1N}^{jN}) - \frac{N}{(1-N)}, & N+1 \leq j \leq 2N \end{cases}$$

where

$$(3.3.6) \quad ||F_H - F_N|| = \max_{1 \leq j \leq 2N} |\sum_j a_{1j} h_j|$$

tion in G_N with weights $h_1, \dots, h_N$ at $\lambda_{1N}, \dots, \lambda_{NN}$. Then,

The actual method of construction is to denote by $H(\lambda)$ the distribu-

$$(3.3.5) \quad P \left\{ \lim_{N \rightarrow \infty} G_N^*(\lambda) = G(\lambda); \lambda \text{ any continuity point of } G \right\} = 1.$$

certain fairly general conditions, Deely and Kruse show:

where $F_N^N(x)$ is the usual empirical distribution function. Then under

$$(3.3.4) \quad ||F_H - F_N|| = \sup_x |F_H^N(x) - F_N^N(x)| \quad \text{for } H \in G_N$$

weights at $\lambda_N, \dots, \lambda_{NN}$. We want to find a G_N^* which minimizes

Hence, for simplification and ease of presentation we set $N_1=3$ for very little difference observed between using $N_1=3$ and using $N_1=N$. much less than N). In fact, in every case considered here there was ably good approximations are obtained even when N_1 is very small (and in the d.f. δ_D^N is an additional parameter, it was found that reason- Although the number of λ 's involved (N_1) in finding the G_N used is just δ_G with G replaced by G_N as described earlier. Deely and Kruse's method used, yielding this EB estimator. δ_D^N

(vi) δ_D^N ; $R(\delta_D^N, G)$ determined in Monte Carlo study.

empirical Bayes estimator, μ_2 known.

$$(v) \quad \delta_N^N; R(\delta_N^N, G) = \frac{\sigma^2}{\sigma^2} \frac{M(N+1)+N_1}{M(N+1)} \quad \text{Empirical Bayes estimator, } \mu \text{ known.}$$

Empirical Bayes estimator, μ known.

(iv) δ_N^N ; $R(\delta_N^N, G)$ determined in Monte Carlo study.

Empirical Bayes estimator, G assumed conjugate.

(iii) δ_N^N ; $R(\delta_N^N, G)$ determined in Monte Carlo study.

The "usual" estimator

$$(ii) \quad \delta_N^N = \frac{t_{N+1}}{\sigma^2}; R(\delta_N^N, G) = \frac{M}{\sigma^2}.$$

The Bayes estimator.

$$(i) \quad \delta_G^G; R(\delta_G^G, G) = \frac{M+N_1}{\sigma^2}.$$

The estimators considered are:

i.e., the prior variance is $1/10$ the variance of T .

$$(ii) \quad \mu = M' = 0; \mu_2 = \frac{N_1}{\sigma^2} = .10 \quad (N_1=10).$$

That is, the prior variance and the variance of T are equal.

$$(i) \quad \mu = M' = 0; \mu_2 = \frac{N_1}{\sigma^2} = 1 \quad (N_1=1).$$

all computations considered here. This simplification also resulted in much shorter computer run times.

The following tables (1 and 2) present the main results on risks in a tabular form for $N=1, 10, 20$ and $M = 1, 5, 10$. The numbers in parentheses are standard errors. If absent, then the results are exact. Besides the value $R(\delta, G)$ for each d.f. δ , the tables give the ratio

$$R(\delta) = R(\delta, G) / R(\delta, G).$$

That is, the ratio of the minimum possible Bayes risk to that obtained by using the d.f. δ . Note that

$$0 \leq R(\delta) \leq 1$$

and obviously the larger $R(\delta)$, the better the decision function δ . However, it must be borne in mind that the estimators δ'_N and δ''_N represent decision situations in which one is prepared to assume more than the usual assumptions concerning G and hence cannot fairly be compared with estimators that do not assume this knowledge.

The figures following the tables (figures 1 and 2) show $R(\delta)$ for each estimator as a function of N . The curves representing $R(\delta'_N)$ and $R(\delta''_N)$ are not given since they obviously lie above $R(\delta_N)$ and, as explained above, have an unfair advantage over the other estimators given. These figures not only allow an easy comparison among the

$$(3.3.10) \quad 0 \leq R(\delta) \leq 1$$

(3.3.9)

various estimators for small N but enable one to check visually the rapidity of convergence of $R(\delta_N)$ to 1.

Note that in the first case (table 1 and figure 1) the relatively large prior variance ($N'=1$) results in a $R(\delta_N)$ value of .833 so there is not much room for improvement in this case. However, even for $N'=1$ it is seen that $R(\delta_N) > R(\delta_N^T)$ for all N . Also, for the small N indicated we see that $R(\delta_D^N) > R(\delta_N^T)$ even though we know that

$$\lim_{N \rightarrow \infty} R(\delta_D^N) = 1.$$

In the second case ($N'=10$) we have $R(\delta_N^T) = .333$ so there is considerable room for improvement over δ_N^T when the prior variance is small. Here we also have $R(\delta_N^T) > R(\delta_N)$ and, in addition,

$$R(\delta_D^N) > R(\delta_N^T) \text{ for the indicated small } N \text{ (even for } N=1).$$

TABLE 1

GLOBAL RISKS OF VARIOUS ESTIMATORS USING
A $N(0, 1)$ PRIOR

Estimator	N=1			N=10			N=20		
	M=1	M=5	M=10	M=1	M=5	M=10	M=1	M=5	M=10
1) Bayes	.500	.167	.091	.500	.167	.091	.500	.167	.091
2) Usual (mean)	1.000	.200	.100	1.000	.200	.100	1.000	.200	.100
3) G conjugate	.983 (.106)	.195 (.047)	.106 (.011)	.741 (.103)	.183 (.029)	.095 (.018)	.661 (.053)	.178 (.011)	.094 (.009)
4) G conjugate, mean known	.885 (.103)	.183 (.034)	.095 (.013)	.618 (.107)	.172 (.082)	.093 (.010)	.505 (.044)	.166 (.015)	.092 (.007)
5) G conjugate, variance known	.750	.183	.095	.545	.170	.092	.524	.168	.091
6) Deely and Kruse's estimator	1.549 (.170)	.304 (.038)	.196 (.027)	1.249 (.164)	.256 (.026)	.185 (.018)	1.106 (.107)	.245 (.020)	.106 (.016)

* Numbers in parentheses are standard errors. If absent, results are exact.

TABLE 2

GLOBAL RISKS OF VARIOUS ESTIMATORS USING
A $N(0, .1)$ PRIOR

Estimator	N=1			N=10			N=20		
	M=1	M=5	M=10	M=1	M=5	M=10	M=1	M=5	M=10
1) Bayes	.091	.067	.050	.091	.067	.050	.091	.067	.050
2) Usual (mean)	1.000	.200	.100	1.000	.200	.100	1.000	.200	.100
3) G conjugate	.647 (.055)	.152 (.018)	.097 (.008)	.198 (.011)	.103 (.009)	.088 (.007)	.155 (.015)	.074 (.008)	.058 (.006)
4) G conjugate, mean known	.538 (.052)	.148 (.010)	.078 (.007)	.175 (.015)	.091 (.021)	.069 (.006)	.148 (.019)	.007 (.008)	.055 (.006)
5) G conjugate, variance known	.545	.133	.075	.174	.079	.055	.134	.073	.052
6) Deely & Kruse's estimator	.772 (.069)	.176 (.012)	.092 (.006)	.377 (.034)	.151 (.015)	.087 (.004)	.293 (.024)	.124 (.017)	.079 (.004)

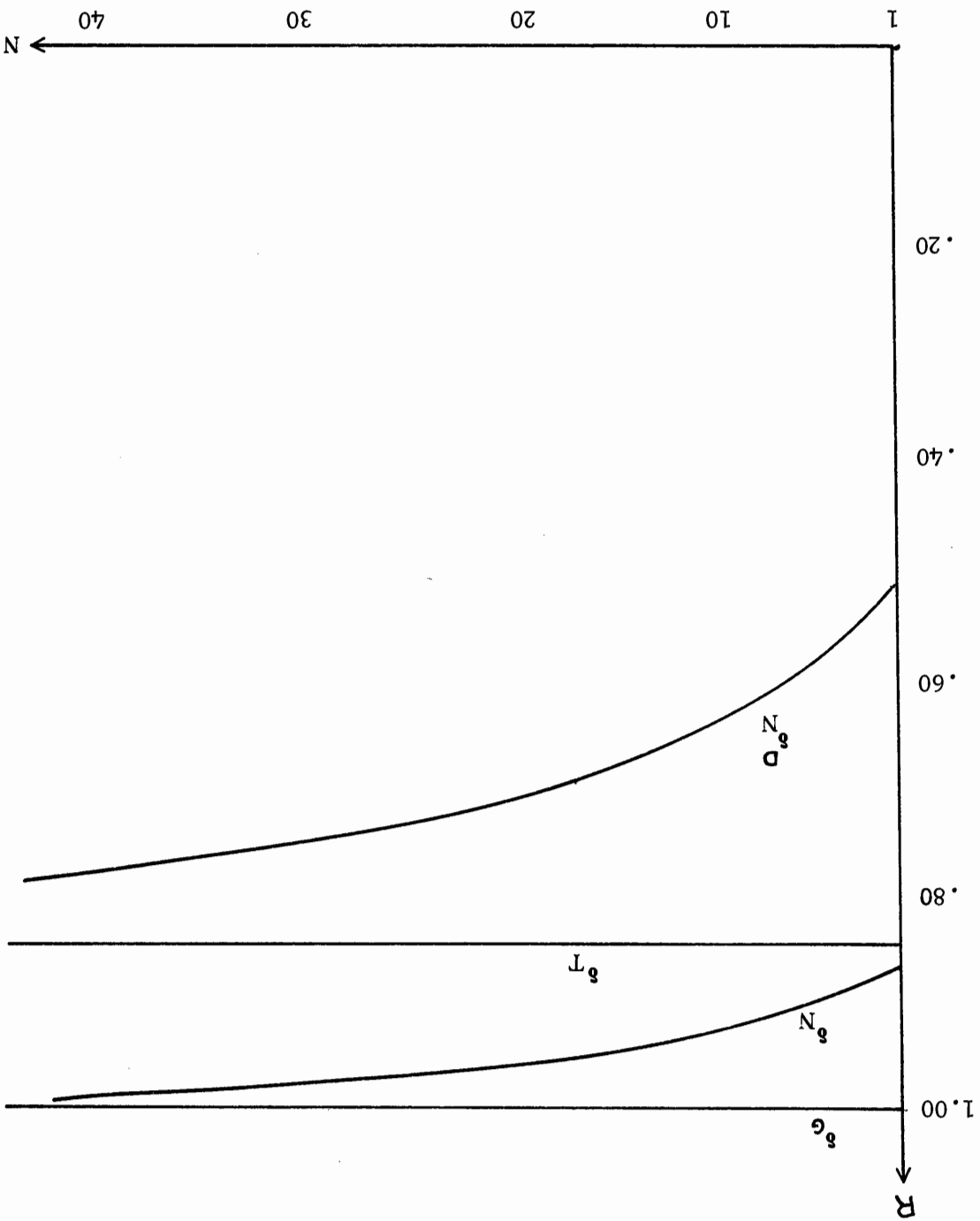


Figure 1. $R(\delta)$ as a function of $N(M'=0, N'=1, M=5)$.

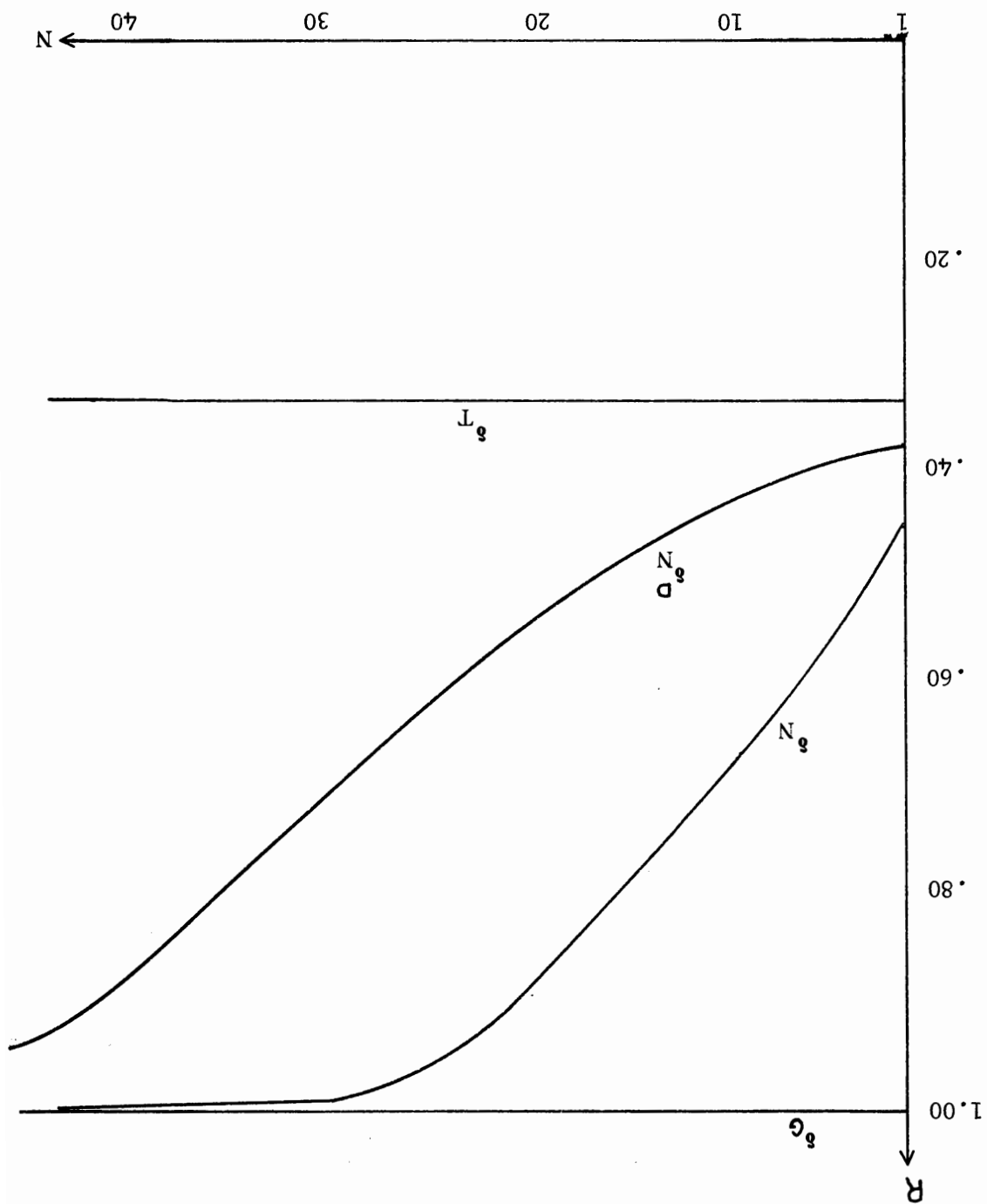


Figure 2. $R(\delta)$ as a function of N ($M'=0$, $N'=10$, $M=5$).

4.1 Introduction

Suppose now that instead of N past observations on $f_G(t)$ it is known only that the distribution function G is a member of some class G . The best we can do in this situation, in the spirit of Chapter II, is to find a G -minimax decision function. The different specifications of the class G to be considered here are:

- (i) G : the conjugate class of priors.
- (ii) G_1 : conjugate class, μ known.
- (iii) G_1^+ : conjugate class, μ known, $\mu_2 \leq \mu_2^+$.
- (iv) G_2 : conjugate class, μ_2 known.
- (v) G_2^+ : conjugate class, μ_2 known, $\mu \in [\mu_L, \mu_U]$.
- (vi) G_3 : μ known, $\mu_2 \leq \mu_2^+$.
- (vii) G_4 : $\mu_2 > \infty$.

It is shown below that cases (i), (ii), and (iv) represent situations in which the G -minimax d.f. is the same in each case. A similar result is true for cases (iii) and (vi). Hence, the above seven categories yield only four distinct decision functions in this G -minimax

context.

In each case, the G -minimax d.f. is obtained, its Bayes risk computed and compared with other decision functions. In section 4.4, asymptotic G -minimax d.f.'s are obtained for the same cases. The general results of Chapter III, relating to the independent normal process and related topics, will not be repeated here but it will be necessary to refer to them throughout this chapter.

4.2 Estimation of λ_{N+1} : G Conjugate

Since the conjugate class G is the class of all normal distributions on λ and any degenerate distribution on λ may be thought of as a normal distribution with zero variance, then the G -minimax estimator is just the ordinary minimax estimator. That is,

$$\delta_G(t_{N+1}) = \delta_T(t_{N+1}) = t_{N+1}. \quad (4.2.1)$$

The Bayes risk for δ_G has already been shown to be:

$$R(\delta_G, G) = \frac{\sigma^2}{M}. \quad (4.2.2)$$

Next, suppose that the mean μ of G is assumed known and denote

the class of all normal distributions with mean μ by G_1 . Now, of course, G_1 does not contain all degenerate distributions on λ .

To find a G_1 -minimax estimator we will attempt to find an estimator δ_0 such that $R(\delta_0, G) = k \geq 0$ for all $G \in G$, and such that δ_0 is Bayes

This strange sounding result occurs because the conjugate class as defined contains, as a limiting distribution, a normal distribution with infinite variance. This distribution is our least favorable

G is simply known to be a conjugate prior distribution function. of the prior does not yield an improved estimator over the case when That is, when no prior observations exist, knowledge of the mean

$$\delta_0(t_{N+1}) = 0 \cdot u + 1 \cdot t_{N+1} = \delta_0(t_{N+1}). \quad (4.2.7)$$

i.e., $\beta_0 = 1$ and $\alpha_0 = 0$. Hence,

Hence, for $R(\delta_0, G)$ to be constant for all $G \in G_1$ we must have

$$\beta_0^2 - 2\beta_0 + 1 = 0 \quad (4.2.6)$$

$$\begin{aligned} R(\delta_0, G) &= E_{\lambda} \left\{ \delta_0(t) - \lambda \right\}^2 \\ &= E_{\lambda} \left\{ \alpha_0^2 u^2 + \beta_0^2 t^2 + \lambda^2 + 2\alpha_0 \beta_0 u t - 2\beta_0 \lambda t \right\} \\ &= E_{\lambda} \left\{ \alpha_0^2 u^2 + \beta_0^2 t^2 + \lambda^2 + 2\alpha_0 \beta_0 u t - 2\beta_0 \lambda t \right\} \\ &= \alpha_0^2 u^2 + \beta_0^2 t^2 + \lambda^2 + 2\alpha_0 \beta_0 u t - 2\beta_0 \lambda t \\ &= \alpha_0^2 u^2 + \beta_0^2 t^2 + \lambda^2 + 2\alpha_0 \beta_0 u t - 2\beta_0 \lambda t \end{aligned} \quad (4.2.5)$$

The Bayes risk for such a δ_0 is

$$\alpha_0, \beta_0 \geq 0 \text{ and } \alpha_0 + \beta_0 = 1. \quad (4.2.4)$$

where

$$\delta_0(t_{N+1}) = \alpha_0 u + \beta_0 t_{N+1} \quad (4.2.3)$$

one) we require δ_0 to be such that: to any $G \in G_1$ is a weighted average of u and t_{N+1} (with weights totalling with respect to some $G \in G_1$. Since the Bayes estimator with respect

distribution and knowledge of the mean of a distribution with infinite variance is of no value in the present context.

However, this problem can be easily avoided by simply restricting μ_2 in such a way that:

$$(4.2.8) \quad 0 \leq \mu_2 \leq \mu_2^+ < \infty$$

That is, μ_2^+ is some fixed upper bound on the variance μ_2 . In practice, (4.2.8) need not be much of a restriction since μ_2^+ may be extremely large.

This additional knowledge about μ_2 would not enable us to do better than δ_T in the case where G is known only to belong to the conjugate family since this class still contains all degenerate distributions.

But if, in addition, the mean is known to be μ_0 , say, then one may suspect that the least favorable distribution in this new class, say G_1^+ , is a normal distribution with mean μ_0 and variance μ_2^+ . If the prior distribution is a $N(\mu_0, \mu_2^+)$, denoted by G_0 , then the Bayes estimator is:

$$(4.2.9) \quad \delta_0(t_{N+1}^+) = \frac{\sigma_2^2 \mu_0 + \mu_2^+ t_{N+1}^+}{\sigma_2^2 + \mu_2^+}$$

and the Bayes risk is:

$$(4.2.10) \quad R(\delta_0, G_0) = \frac{\mu_2^+ \sigma_2^2}{\sigma_2^2 + \mu_2^+}$$

and for any other $G \in G_1^+$, we have:

$$\frac{R(\delta^I, G)}{R(\delta, G)} = \frac{\frac{\sigma_2^{M+Mu_2} + \sigma_4^{u_2}}{\sigma_2^{\frac{M}{2}}}}{\frac{\sigma_2^{M(u_2)} + \sigma_4^{u_2}}{\sigma_2^{\frac{M}{2}}}} \quad (4.2.14)$$

For any GeG_+^I we have

$$\sup_{GeG_+^I} \frac{R(\delta^I, G)}{R(\delta, G)} = \frac{R(\delta^I, G)}{R(\delta, G_0)} = \frac{M + \sigma/u_2}{M} > 1. \quad (4.2.13)$$

By using δ_+^I , the worst we can do, relative to δ^I is:

$$\delta_+^I(t_{N+1}) = \delta_+^I(t_{N+1}) \quad (4.2.12)$$

For convenience and consistency we denote δ_0 by

Hence by theorem 3, G_0 is least favorable and δ_0 is G_+^I -minimax.

$$\begin{aligned} &= R(\delta_0, G_0) \text{ for all } GeG_+^I \\ &= \frac{\sigma_2^{u_2} + \sigma_2^{M+Mu_2}}{\sigma_2^{M(u_2)} + \sigma_4^{u_2}} \\ &\leq \frac{\sigma_2^{M(u_2)} + \sigma_2^{u_2}}{\sigma_2^{M(u_2)} + \sigma_4^{u_2}} \\ &= \frac{\sigma_2^{M(u_2)} + \sigma_4^{u_2}}{\sigma_2^{M(u_2)} + \sigma_4^{u_2}} \\ &R(\delta_0, G) = E^{t/\lambda} (\delta_0(t) - \lambda)^2 \end{aligned} \quad (4.2.11)$$

That is, regardless of which $G \in G_1$ is in effect, we find that $\delta_{G_1}^+$ is preferred to $\delta_{G_1}^+$.
 As a particular example, suppose $M=2$, $\sigma_2^2=4$ and $\mu_2^+ = 10.00$. Then figure 3 gives $R(\delta_{G_1}^+, G)$ and $R(\delta_{G_1}^+, G)$ as a function of μ_2 .
 Suppose next that μ_2 , but not μ , is known. We denote the class of conjugate priors with variance μ_2 by G_2 .
 Now, a Bayes estimator δ_0 for any $G \in G_2$ is

$$\delta_0(t_{N+1}) = \frac{\sigma_2^2 \mu_2^+ + \sigma_2^2 \mu_2}{\sigma_2^2 + \sigma_2^2 \mu_2^+} \quad (4.2.15)$$

where α is the (unknown) mean of G . If $\alpha = t_{N+1}$ then

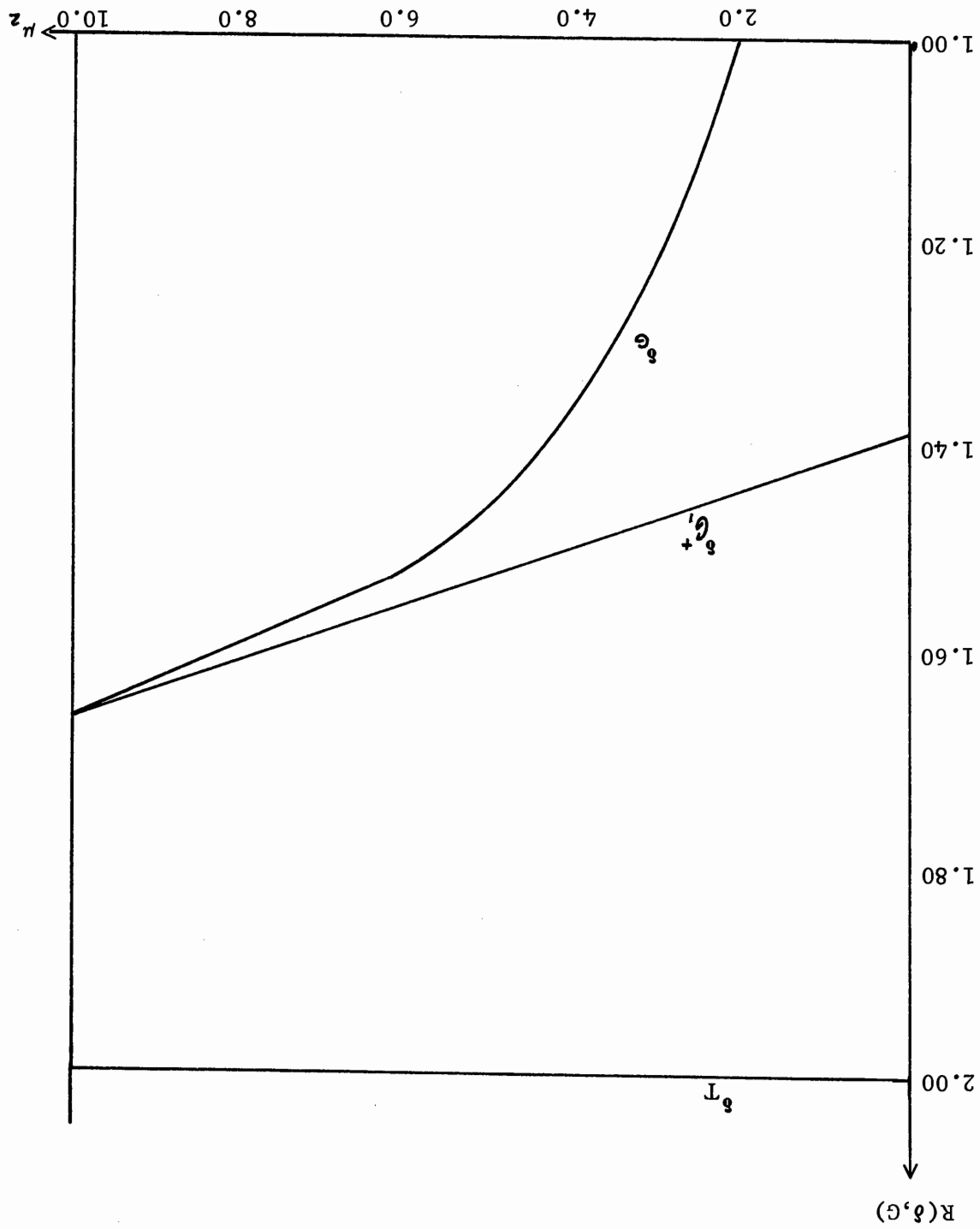
$$\delta_0(t_{N+1}) = t_{N+1} = \delta_{G_1}^+(t_{N+1}) \quad (4.2.16)$$

and

$$R(\delta_0, G) = \frac{M}{\sigma_2^2} \text{ for all } G \in G_2. \quad (4.2.17)$$

Since t_{N+1} is Bayes with respect to a $N(t_{N+1}, \mu_2)$ distribution then, by theorem 5, $\delta_{G_1}^+$ is G_2 -minimax.

Another way to see this is to assume $\alpha \neq t_{N+1}$. Then,

Figure 3. Bayes risk of δ'_+ and δ'_T as a function of μ_z 

defined by (4.2.16) if we are willing to assume

Just as in the previous case (μ known) we can improve upon δ_0

better than $\delta_{\mathbb{T}}$ and, in fact, δ_0 may be much worse.

That is, if we choose $\alpha = t_{N+1}$ then we can never be sure that δ_0 is

$$(4.2.21) \quad \sup_{G \in G_2} R(\delta_0, G) \geq \sup_{G \in G_2} R(\delta_{\mathbb{T}}, G).$$

true so that

Hence, there obviously always exists a $G \in G_2$ such that (4.2.20) is

$$(4.2.20) \quad R(\delta_0, G) \geq \frac{M}{\sigma^2} = R(\delta_{\mathbb{T}}, G).$$

we have

$$(4.2.19) \quad (\mu - \alpha)^2 \geq \frac{M}{\sigma^2} + \mu^2$$

and for any $G \in G_2$ such that

$$(4.2.18) \quad R(\delta_0, G) = E_{t, \lambda} \left\{ \frac{\sigma^2 + \mu^2}{\mu^2} \cdot \frac{M}{\sigma^2} + \frac{\sigma^2 + \mu^2}{\sigma^2 (\alpha - \lambda)^2} \right\} = \frac{(\sigma^2 + \mu^2)^2}{\sigma^2 M (\mu^2)^2 + \sigma^4 [\mu^2 + (\mu - \alpha)^2]}$$

$$\sup_{G_2^+} (u - \alpha^*)^2 = (u^n - u_L)^2 / 4. \quad (4.2.26)$$

then we see that

$$\alpha^* = u^n - \frac{|u^n - u_L|^2}{u^n + u_L} \quad (4.2.25)$$

If we let

$$\sup_{G_2^+} (u - \alpha)^2 = \max \left\{ (u_L - \alpha)^2, (u^n - \alpha)^2 \right\} < (u^n - u_L)^2. \quad (4.2.24)$$

where

$$\sup_{G_2^+} R(\delta_0, g) = M(u^2) \frac{\sigma^2 + \sigma_+^2 [u^2 + \sup (u - \alpha)^2]}{g} \quad (4.2.23)$$

it is easy to see that

consider are of the form (4.2.15) where $u_L < \alpha < u^n$. From (4.2.18)

Under the restriction (4.2.22), the only estimators we need

with mean in the interval $[u_L, u^n]$ by G_2^+ .

We denote the class of conjugate priors with variance u^2 and

strong restriction in practice since $|u_L|$ and $|u^n|$ can be quite large.

which the unknown μ is assumed to lie. Again, this need not be a

That is, if we are willing to place upper and lower limits within

$$-\infty < u_L < \mu < u^n < \infty. \quad (4.2.22)$$

Defining δ_0^* by:

$$\delta_0^* (t_{N+1}^*) = \sigma^2 \alpha^* + M u_2 t_{N+1}^* \quad (4.2.27)$$

then we have:

$$\sup_{\delta \in G_+^2} R(\delta_0^*, G) = \inf_{\delta \in G_+^2} \sup_{R(\delta, G)} \quad (4.2.28)$$

that is, δ_0^* is G_+^2 -minimax. As a comparison with the usual estimator

we note:

$$\frac{\sup_{\delta \in G_+^2} R(\delta_0^*, G)}{R(\delta_0^*, G)} = \frac{R(\delta_0^*, G)}{R(\delta_0^*, G)} = \frac{R(\delta_0^*, G)}{R(\delta_0^*, G)} \quad (4.2.29)$$

and the above "worst possible" or maximum ratio is ≤ 1 iff:

$$(M u_2)^2 + \sigma^2 M u_2 + \sigma^2 M \left[\frac{4}{(u_{L-u}^n)^2} \right] \leq (\sigma^2 + M u_2)^2 \quad (4.2.30)$$

$$\Leftrightarrow \frac{4}{M} (u_{L-u}^n)^2 \leq M u_2 + \sigma^2$$

$$\Leftrightarrow (u_{L-u}^n)^2 \leq 4 \left(\frac{M}{\sigma^2} + u_2 \right) = 4 \text{ Var } (T)$$

i.e., if $(u_{L-u}^n)^2$ is less than or equal to four times the (unconditional) variance of T then the G_+^2 -minimax estimator δ_0^* is to be preferred to δ_T for all $G \in G_+^2$. Note, however, that even if the inequality (4.2.30) is

not satisfied δ_0^* may still be preferred to δ_T . This depends on just how close α^* is to the true μ .

As a particular example, suppose $\mu_2=4$, $\sigma^2=1$ and $M=2$. Then, $(\mu_T - \mu_u)^2$ must be less than or equal to 14. If it is true that $\mu_T = -3$, $\mu_u=0$ then $(\mu_T - \mu_u)^2 = 9 < 16$ and regardless of the true μ ,

$$\delta_0^*(t_{N+1}) = [\sigma^2(-3/2) + \mu_2 t_{N+1}^2] / (\sigma^2 + \mu_2)^2 \quad (4.2.31)$$

is a better estimator than δ_T . In fact,

$$R(\delta_T, G) = \frac{1}{2} = .50 \neq G \sigma^2 \quad (4.2.32)$$

and

$$\sup_{*} R(\delta_0^*, G) = \frac{51}{108} = .47 \quad (4.2.33)$$

while the Bayes envelope function is: (by (3.2.7)):

$$R(G) = \frac{\sigma^2}{M+n} = \frac{1}{2+1} = \frac{1}{3} = .44 \quad (4.2.34)$$

regardless of the particular $G \sigma^2$.

4.3 Estimation Of λ_{N+1} : Moments Known

In this section we do not assume that G is conjugate but only that certain prior moments are known (or at least known to lie in some interval of possible values).

$$(4.3.3) \quad \int (\lambda - \mu)^2 dG_{\lambda_0}(\lambda) = (\lambda_0 - \mu)^2 \varepsilon + (\lambda_0^* - \mu)^2 (1 - \varepsilon)$$

is true since

so that the variance is less than μ_2^+ as required by (4.3.1). This as required of any distribution in G_3 . Also, we can always choose ε

$$(4.3.2) \quad \int \lambda dG_{\lambda_0}(\lambda) = \lambda_0 \varepsilon + (\lambda_0^*) (1 - \varepsilon) = (\mu - \alpha) \varepsilon + \mu - \frac{1 - \varepsilon}{\varepsilon} \cdot \alpha (1 - \varepsilon) = \mu$$

$\beta = \left(\frac{1 - \varepsilon}{\varepsilon} \right) \alpha$. Then,

function that puts weight ε at λ_0 and weight $1 - \varepsilon$ at $\lambda_0^* = \mu + \beta$ where To see this, we let $\lambda_0 = \mu - \alpha$ and choose G_{λ_0} as that distribution

$G_{\lambda_0}(\lambda_0^-) > 0$. since for every $\lambda_0 \in \Lambda$ there exists a $G_{\lambda_0} \in G_3$ such that $G_{\lambda_0}(\lambda_0^-) > 0$. Note that although G_3 does not contain all degenerate distributions on Λ , it is true that G_3 -admissibility implies admissibility

by μ_2^+ by G_3 . We denote the class of all priors with mean μ and variance bounded

$$(4.3.1) \quad 0 \leq \mu_2 \leq \mu_2^+ < \infty$$

that

variance μ_2 . However, as in the last section, suppose it is known mainly because we have no way to estimate other moments, e.g., the For example, if μ is known then we can do no better than t^{N+1}

so that G_0 is least favorable in G_3 (as well as G_1^+) and δ_0 is G_3 -minimax (as well as G_1^+ -minimax). Hence, although $G_1^+ \subset G_3$, the restriction G_1^+ does not enable us to lower the Bayes risk from the case G_3 . That is, knowledge that G_1^+ does not enable us to find a better estimator (in the current context) than that attainable when G_3 . This situation arises because restriction of G to the conjugate class still entails consideration of the least favorable distribution in G_3 . Hence we have one more example of a situation in which restriction of consideration to the conjugate class is really not a stringent restraint.

$$R(\delta_0, G) \leq R(\delta_0, G_0) \quad (4.3.4)$$

obvious (by (4.2.11)) that for all $G \in G_3$:
Also, just as for the class G_1^+ in the previous section, it is is given by (4.2.10).
is the Bayes estimator with respect to G_0 and the Bayes risk for δ_0 consider a $N(\mu, \mu_+^2)$ distribution, say G_0 , then δ_0 (given by (4.2.9)) distribution with mean μ and variance μ_+^2 . For example, if we
A good guess at a least favorable distribution in G_3 is any $\lambda_0 \in \Lambda$.

so that the variance is less than μ_+^2 . That is, $G_{\lambda_0} \in G_3$ for any
Hence, regardless of the size of $|\lambda - \lambda_0|$ we can always choose
and the above variance approaches zero as ϵ approaches zero.

$$= \alpha \left(\epsilon^2 + \frac{1-\epsilon}{\epsilon} \right)$$

$$R(\delta_0^*, G) = \frac{\sigma_2^2 \mu_2}{\sigma_2^2 + M \mu_2^2} \quad \forall G \in G_4 \quad (4.3.7)$$

has a Bayes risk of

$$\delta_0^*(t_{N+1}^*) = \frac{\sigma_2^2 \mu_2}{\sigma_2^2 + M \mu_2^2} \quad (4.3.6)$$

Now, the estimator

distributions on λ with mean $= \mu$ and variance $= \mu^2$ will be denoted by G_4 .

We next assume that μ and μ_2 are both known. The class of all

one as $M \rightarrow \infty$ or $\alpha \rightarrow 1$.

α is large. In fact, as would be expected, $R(\delta_0^*, G^*)/R(\delta_{G^*}^*, G^*)$ approaches

where $\alpha = \mu_2^2/\mu_+^2$. The above ratio is always ≥ 1 but is small if M or

$$\begin{aligned} & \frac{R(\delta_0^*, G^*)}{R(\delta_{G^*}^*, G^*)} = \frac{\frac{\sigma_2^2 + M \mu_+^2}{\mu_+^2}}{\frac{\sigma_2^2 + M \mu_2^2}{\mu_2^2}} \cdot \frac{\alpha}{1 + \frac{\alpha}{(\alpha-1)^2} \frac{\sigma_2^2 \mu_2^2}{\sigma_2^2 + M \mu_2^2}} \\ & = \frac{\frac{\sigma_2^2 + M \mu_+^2}{\mu_+^2}}{\frac{\sigma_2^2 + M \mu_2^2}{\mu_2^2}} \cdot \frac{\mu_2^2}{M(\mu_+^2 + \sigma_2^2 \mu_+^2)} \\ & = \frac{R(\delta_0^*, G^*)}{R(\delta_{G^*}^*, G^*)} = \frac{\frac{\sigma_2^2 + M \mu_+^2}{\mu_+^2}}{\frac{\sigma_2^2 + M \mu_2^2}{\mu_2^2}} \cdot \frac{\mu_2^2}{\sigma_2^2 + M \mu_+^2} \end{aligned} \quad (4.3.5)$$

we have:

We note that if the true G in effect is a $N(\mu, \mu_+^2)$, say G^* , then

as in previous sections and consider the estimator:

$$\hat{u} = \sum_{i=1}^{N+1} t_i / (N+1) \quad (4.4.3)$$

where u_0^+ is some known finite upper bound for u_2 . Next we define

$$0 \leq u_2 \leq u_2^+ < \infty \quad (4.4.2)$$

Assume, in the first case, that the unknown u_2 is such that

where δ_G is the G -minimax estimator for the class G .

$$\lim_{N \rightarrow \infty} R(\delta_N^N, G) \rightarrow R(\delta_G, G) \quad \forall G \in \mathcal{G} \quad (4.4.1)$$

past observations to form an estimator δ_N^N , in such a way that

Here, we have N past observations on $f_G(t)$ and we wish to use these

In this section we derive various asymptotic G -minimax estimators.

4.4 Asymptotic G -Minimax Decision Functions

so that regardless of the true G in effect, δ_0^* is preferred to δ_T .

$$R(\delta_0^*, G) \leq R(\delta_T, G) \quad \forall G \in \mathcal{G}_T \quad (4.3.8)$$

again, we note that:

G_T -minimax and a $N(u, u_2)$ distribution is least favorable. Here, δ_0^* is Bayes with respect to a $N(u, u_2)$ and hence by theorem 5, δ_0^* is i.e., the Bayes risk for δ_0^* is constant over the class G_T . Also, δ_0^*

$$R_N(\delta_N, G) = E_{t, \lambda} \left\{ \delta_N(t_{N+1}) - \lambda \right\} = E_{\lambda} \left\{ \frac{\left(M_{\mu_2} + \frac{\sigma_2^2}{M} \right) \frac{(\sigma_2^2 + M_{\mu_2})}{\sigma_2^2}}{\left(M_{\mu_2} + \frac{\sigma_2^2}{M} \right) \frac{(\sigma_2^2 + M_{\mu_2})}{\sigma_2^2}} + \left[\frac{\sigma_2^2 t_{N+1}}{(\sigma_2^2 + M_{\mu_2})} - \lambda \right] \right\} \quad (4.4.8)$$

The global risk of δ_N is computed by first finding:

$$\delta_N(t_{N+1}) = \frac{\sum_{i=1}^N \sigma_2^2 t_i}{\sigma_2^2 + M_{\mu_2}} + t_{N+1} \left(\frac{M_{\mu_2}}{\sigma_2^2} + \frac{N+1}{\sigma_2^2} \right) \quad (4.4.7)$$

as:

where $\delta_0(t_{N+1})$ is defined by (4.2.9). Now, we can rewrite (4.4.4)

$$\delta_N(t_{N+1}) \xleftarrow{\text{a.s.}} \delta_0(t_{N+1}) \quad (4.4.6)$$

and hence that

$$\hat{\mu} \xleftarrow{\text{a.s.}} \mu \quad (4.4.5)$$

We have seen that

$$\delta_N(t_{N+1}) = \frac{\sigma_2^2 \mu + M_{\mu_2}}{\sigma_2^2 + M_{\mu_2}} \quad (4.4.4)$$

i.e., δ_N is asymptotically G_3 -minimax. As a matter of fact, δ_N is

$$\lim_{N \rightarrow \infty} R(\delta_N, G) = \frac{M(\sigma_2^2 + M\mu_+^2)^2}{M(\mu_+^2)^2 \sigma_2^2 + M\sigma_4^2 \mu_+^2} = \frac{M(\sigma_2^2 + M\mu_+^2)^2}{M(\mu_+^2)^2 \sigma_2^2 + M\sigma_4^2 \mu_+^2} \quad (4.4.10)$$

and therefore we have:

$$R(\delta_N, G) = E[R_N(\delta_N, G)] = \frac{M(N+1)^2 (\sigma_2^2 + M\mu_+^2)^2}{(M(N+1)\mu_+^2 + \sigma_2^2)^2 \sigma_2^2 + M\sigma_4^2 N^2} \left\{ \mu_+^2 + \frac{\sigma_2^2 + M\mu_+^2}{MN} \right\} \quad (4.4.9)$$

$$\frac{\Sigma t_i}{N} \sim N \left(\mu, \frac{\sigma_2^2 + M\mu_+^2}{MN} \right). \quad \text{Hence,}$$

$$\begin{aligned} & \frac{M(N+1)^2 (\sigma_2^2 + M\mu_+^2)^2}{\left[M\mu_+^2 (N+1) + \sigma_2^2 \right]^2 \sigma_2^2 + M\sigma_4^2 N^2} \left\{ \mu_+^2 + \left[\frac{\Sigma t_i}{N} - \mu \right]^2 \right\} \\ &= \left(M\mu_+^2 + \frac{\sigma_2^2}{M} \right) \frac{N+1}{\sigma_2^2} + \frac{N^2 \sigma_4^2 E_1 \left\{ \left[\frac{\Sigma t_i}{N} - \mu \right]^2 \right\}}{M(N+1)^2 (\sigma_2^2 + M\mu_+^2)^2} \end{aligned}$$

That is, δ_N^* is an asymptotic G_4 -minimax estimator.

$$(4.4.14) \quad \frac{\delta_0^*(t_{N+1})}{\sigma_{2+M\mu_2}^2} = \mu_{\sigma_{2+M\mu_2}^2} t_{N+1}.$$

where δ_0^* is defined by

$$(4.4.13) \quad \lim_{N \rightarrow \infty} R(\delta_N^*, G) = R(\delta_0^*, G) \quad \text{for all } G \in G_4$$

that

jugate priors are independent of the normality assumption, we have asymptotic optimality of the estimator δ_N^* for the class of conjugate priors. Now, since the steps following (3.2.25) establishing the

as in (3.2.20).

$$\mu_2 = \max \left\{ 0, s - \frac{\sigma_2^2}{M} \right\},$$

$$(4.4.12) \quad \mu = \sum_{i=1}^{N+1} t_i / (N+1)$$

where μ and μ_2 are defined by:

$$(4.4.11) \quad \delta_N^*(t_{N+1}) = \frac{\sigma_{2+M\mu_2}^2}{\sigma_{2+M\mu}^2} t_{N+1}$$

used in this case is:

functions on Λ with μ_2 finite (but unknown) by G_4 . The estimator known bounds for μ or μ_2 . Denote the class of all distribution Next, consider the case where we are unwilling to assume

also asymptotically G_1^+ -minimax (see (4.2.11)).

BERNOULLI PROCESS I: EMPIRICAL BAYES DECISION FUNCTIONS

5.1 Definitions And Preliminary Results

The Bernoulli process as used in this chapter is defined as a generator of X 's with probability mass function:

$$(5.1.1) \quad \lambda^x (1-\lambda)^{1-x} \quad x=0,1 \quad . \quad 0 < \lambda < 1$$

The Bernoulli process could also be defined as a generator of N 's

with probability mass function (p.m.f.)

$$(5.1.2) \quad \lambda (1-\lambda)^{N-1} \quad N=1,2,\dots \quad . \quad 0 < \lambda < 1$$

Here, for example, N may represent the number of trials required to obtain the first success with a probability of success λ at each trial. If we let t equal the number of N 's observed and $M=\sum_{i=1}^t N_i$ then the likelihood of the sample is proportional to:

That is, each T_i has a binomial distribution with (known) parameter λ_i and unknown parameter λ_i . The p.m.f. for each T_i is simply

$$(5.1.6) \quad T_i / \lambda_i \sim B(M, \lambda_i) .$$

It is well-known that

some inference concerning the unknown λ_{N+1} .

is a sufficient statistic for λ_i . The problem, as usual, is to make

$$(5.1.5) \quad t_i = \sum_{j=1}^M x_{ij}$$

where in each sample, for $i \in \{1, \dots, N+1\}$,

$$(5.1.4) \quad \left\{ x_{1j} \right\}_{j=1}^M, \dots, \left\{ x_{Nj} \right\}_{j=1}^M$$

are denoted, as in previous chapters, by

$x_{N+1,j}$ having the p.m.f. given by (5.1.1). The past observations
A current random sample $\left\{ x_{N+1,j} \right\}_{j=1}^M$ of size M is chosen, each

of the Bernoulli process.

z_{x_i} . Hence, for convenience, we consider only the first definition
about the parameter λ using t should be the same as that made using
by the likelihood principle (Lindley [21], p. 59) any statement made
will be the same in either definition of the Bernoulli process then
which in turn is proportional to (5.1.1). Since the prior densities

$$(5.1.3) \quad \lambda^t (1-\lambda)^{M-t}$$

$$\begin{aligned}
 \mu_2 &= \frac{R'(N, N'+1)}{R'(N, -R')} = \frac{N(N'+1)}{\mu(1-\mu)} \\
 \mu_2' &= \frac{R'(R'+1)}{R'(N'+1)} = \frac{N'+1}{\mu(N'+1)} \\
 \mu &= \frac{R'}{N'}
 \end{aligned}
 \tag{5.1.10}$$

distribution are:

and is denoted by $Be(R', N'-R')$. The first few moments of the beta

$$0 < \lambda < 1, \quad N' > R' > 0$$

$$g(\lambda | R', N') = \frac{\Gamma(R')\Gamma(N'-R')}{\Gamma(N')} \lambda^{R'-1} (1-\lambda)^{N'-R'-1} \tag{5.1.9}$$

a beta distribution with p.d.f.

parameter λ . The conjugate prior distribution in this case is just

and $G(\lambda)$ is the prior probability distribution function of the

$$\text{where } \lambda = [0, 1]$$

$$F_G(t) = \int_0^t \binom{t}{m} \lambda^{m-t} (1-\lambda)^{M-t} dg(\lambda) \tag{5.1.8}$$

distribution with p.m.f.:

be considered a random sample from the compound or unconditional
with $E(T) = M\lambda$ and $E(T-M\lambda)^2 = M\lambda(1-\lambda)$. Also, the set $\left\{t_i \mid i=1, \dots, N+1\right\}$ can

$$0 < \lambda < 1$$

$$t \in \{0, 1, \dots, M\}$$

$$F(t|\lambda) = \binom{t}{m} \lambda^{m-t} (1-\lambda)^{M-t} \tag{5.1.7}$$

and the parameters R' and N' in terms of the first two moments are:

$$(5.1.11) \quad R' = \frac{\mu_2}{\mu(\mu-\mu_1)} \quad N' = \frac{\mu_2}{\mu-\mu_1} \quad .$$

Some useful facts concerning distributions on $[0,1]$ in general and beta distributions in particular are derived below.

For any distribution on $[0,1]$ it is obvious that

$$(5.1.12) \quad \int_0^1 \lambda^j dG(\lambda) > \int_0^1 \lambda^{j+1} dG(\lambda)$$

$$j=0,1,2,\dots$$

since for $0 < \lambda < 1$,

$$(5.1.13) \quad \lambda^j > \lambda^{j+1} \quad j=0,1,2,\dots$$

Also, since the range of λ is bounded it is true that

$$(5.1.14) \quad \mu_j^j < \infty \quad j=1,2,\dots$$

In fact, an equivalent way to write (5.1.12) is:

$$(5.1.15) \quad \mu_j^j > \mu^{j+1} \quad j=0,1,2,\dots$$

and Skibinsky [40] shows that the maximum range for μ_j^j is 2^{-2j+2} .

From (5.1.15) we have that

$$(5.1.16) \quad \mu_1^2 \leq \mu \quad \text{i.e., } \mu - \mu_1^2 \geq 0$$

and if $N' \geq 1$

$$(5.1.20) \quad \mu_2 \leq \mu - \mu_2' \leq \mu$$

$$\begin{aligned} \mu_2 &= \\ &= \frac{N}{N'} \frac{1+N}{\mu_2} \\ &= \frac{N}{N'} \frac{1+N}{\mu(N', \mu+1)} = \mu - \mu_2' = \frac{N}{R'(N', -R')} \frac{N}{R'(N', -R')} \\ &= \frac{N}{\mu - \mu_2'} \frac{1+N}{\mu(N', \mu+1)} \\ &= \mu_2' - \mu_2 = \frac{N}{R'(N', -R')} \frac{N}{R'(N', -R')} \end{aligned} \quad (5.1.19)$$

For a Be ($R', N'' - R''$) distribution (see (5.1.10)):

$$\begin{aligned} \mu_2 &\leq \mu_2' \leq \mu \\ \mu - \mu_2' &\leq \mu - \mu_2 \leq \frac{1}{4} \\ \mu_2 &= \mu_2' - \mu_2 \leq \mu - \mu_2 \leq \frac{1}{4} \end{aligned} \quad (5.1.18)$$

For any distribution on $\Lambda = [0, 1]$:

The following inequalities are also useful in the later sections.

as is true for any distribution.

$$(5.1.17) \quad \mu_2 = \mu_2' - \mu_2 \geq 0$$

although, of course,

$Be(R', N'-R')$ is:

The posterior density of λ given t_{N+1} when the prior density is a

$$\begin{aligned} E(T^2) &= M(M-1)\mu_2 + M\mu \\ \text{Var}(T) &= M(M-1)\mu_2 + M\mu(1-\mu) \\ &= \frac{M(M+N')R'(N'-R')}{(N'+1)^2} \\ E(T) &= \frac{MR'}{N'} = M\mu \end{aligned} \quad (5.1.23)$$

and M . This distribution has the following first few moments:

That is, T has a beta - binomial distribution with parameters R' , N'

$$F_G(t) = \frac{\Gamma(t+1)\Gamma(M-t+1)\Gamma(N'-R')\Gamma(M+N')}{\Gamma(t+R')\Gamma(M+N'-t-R')\Gamma(M+1)\Gamma(N')} \quad (5.1.22)$$

density of t becomes

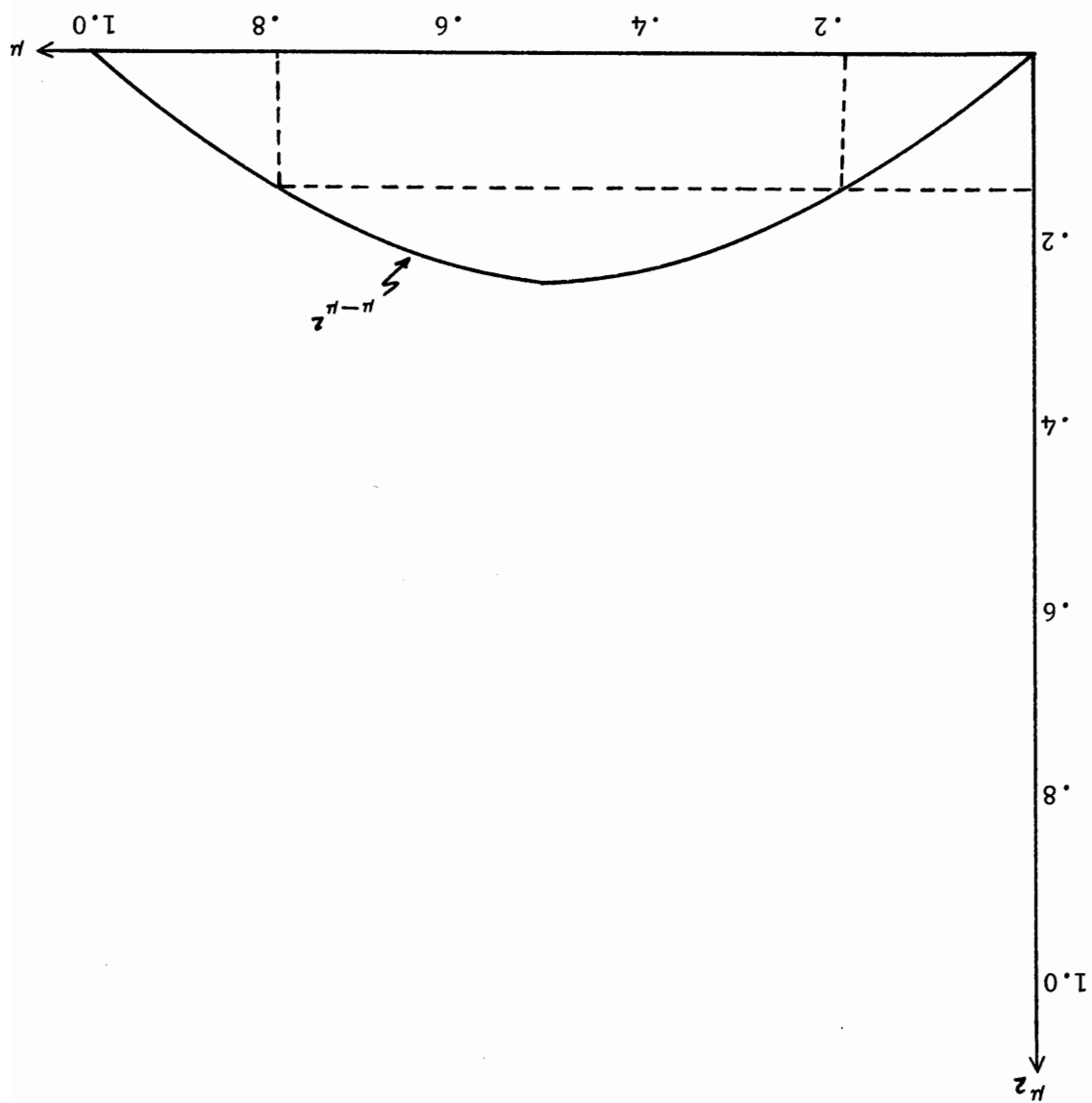
When $G(\lambda)$ is a $Be(R', N'-R')$ distribution the unconditional

$$\mu_2 \leq \frac{4(N'+1)}{1} \quad \mu - \mu_2 \leq \left(\frac{N'}{N'+1} \right)^{\frac{1}{4}} \quad (5.1.21)$$

fixed interval. See figure 4. Hence, for any beta distribution interval and for a given μ , μ_2 (or μ_2) can take values only in a That is, for a given μ_2 , μ can take values only in a strictly defined

$$\mu_2 \leq \mu - \mu_2; \text{ i.e., } 1 - \sqrt{1-4\mu_2} \leq \mu \leq 1 + \sqrt{1-4\mu_2} \quad \text{and } \mu_2 \leq \mu_2 \leq \mu.$$

Note that for any distribution on $[0,1]$ we have (by (5.1.18))

Figure 4. Upper and lower limits for μ and μ^2 .

When attempting to derive empirical Bayes decision functions in

5.2 Estimation of λ_{N+1}

where $\text{Var}_{M_1}(\cdot|t)$ represents the posterior variance of λ for a sample size M_1 .

$$\text{Var}_{M_1}(\lambda|t_{N+1}) > \text{Var}_{M_2}(\lambda|t_{N+1}) \quad (5.1.27)$$

Notice that for $M_1 < M_2$

$$= \frac{(R' + t_{N+1})^{N'+M-1} (N'+M-1)!}{(N'+M)!} \cdot$$

$$E(\lambda) = \frac{R''}{R'' + t_{N+1}} = \frac{R''}{N'' + M}$$

(5.1.26)

and this posterior distribution has a mean and variance of:

$$R'' = R' + t_{N+1} \quad N'' = N' + M \quad (5.1.25)$$

that is, $\lambda|t_{N+1} \sim \text{Be}(R'', N''-R'')$ where

$$g(\lambda|t_{N+1}) = \frac{\Gamma(N'')}{\Gamma(R'')\Gamma(N''-R'')} \lambda^{R''-1} (1-\lambda)^{N''-R''-1} \quad (5.1.24)$$

$$L(\delta(t), \lambda) = (\delta(t) - \lambda)^2 \quad (5.2.2)$$

The loss function to be used in estimating λ_{N+1} is:

an a.o. EB d.f. for some restrictive class G .

the EB procedure depends on $f_{G_1}(t)$. However, it may be possible to find to find an asymptotically optimal EB decision function for this case since imply that $\delta_{G_1}(t) \equiv \delta_{G_2}(t)$ then if G is unrestricted it is impossible identifiable in G . Now, in general since $f_{G_1}(t) \equiv f_{G_2}(t)$ does not $= f_{G_2}(t)$, $G_1, G_2 \in G \Rightarrow G_1 \equiv G_2$ then the mixture $f_G(t)$ is said to be For a given class G of distribution functions on Λ , if $f_{G_1}(t)$

Definition

of all distribution functions on Λ . then $f_{G_1}(t) \equiv f_{G_2}(t)$ so that $f_G(t)$ is not "identifiable" in the class the first M moments of G_1 are identical with the first M moments of G_2 and hence if any two distribution functions G_1 and G_2 are such that

$$\begin{aligned} &= \sum_{j=0}^{M-t} \binom{M}{j} (-1)^j \binom{j}{t+j} \lambda^{t+j} dg(\lambda) \\ &= \sum_{j=0}^{M-t} \binom{M}{j} (-1)^j \binom{j}{t+j} \lambda^{t+j} dg(\lambda) \\ &= \sum_{j=0}^{M-t} \binom{M}{j} (-1)^j \lambda^j \left[\sum_{j=0}^{M-t} \binom{j}{t+j} \lambda^{t+j} dg(\lambda) \right] \\ &= \int_{M-t}^V \lambda^{t(1-\lambda)} dg(\lambda) \end{aligned} \quad (5.2.1)$$

the binomial case, it is important to note that

The usual (minimum variance unbiased) estimator of λ_{N+1} is denoted

by $\delta_{\mathbb{T}}$:

$$(5.2.3) \quad \delta_{\mathbb{T}}(t_{N+1}) = t_{N+1} / M$$

and has a Bayes risk of

$$(5.2.4) \quad R(\delta_{\mathbb{T}}, G) = E_{\mathbb{T}} \lambda_{\mathbb{T}} / (t/M - \lambda)^2$$

$$= E_{\mathbb{T}} \text{Var}(\mathbb{T}/M)$$

$$= E_{\mathbb{T}} \frac{\lambda}{\lambda(1-\lambda)} \frac{1}{M}$$

$$= \frac{M}{\mu - \mu^2}$$

for any prior G with mean μ and second moment μ^2 . (we are always

assuming $\int_0^1 dg(\lambda) = 1$). From (5.1.17) and (5.1.18) we note that

for any distribution G it is true that

$$(5.2.5) \quad 0 \leq R(\delta_{\mathbb{T}}, G) \leq \frac{M}{\mu - \mu^2} \leq \frac{1}{4M}$$

and for a Beta $(R', N'-R')$ prior

$$(5.2.6) \quad 0 \leq R(\delta_{\mathbb{T}}, G) = \frac{1}{M} \cdot \frac{N'+1}{N'} \cdot \frac{M}{(\mu - \mu^2)} \leq \frac{1}{4M} \cdot \frac{N'+1}{N'}$$

The Bayes estimator δ_G with respect to a Be $(R', N'-R')$ prior is:

$$(5.2.7) \quad \delta_G(t_{N+1}) = \frac{t_{N+1} + R'}{M + N'}$$

In the above sense then, a beta prior with $N'=R'=0$ is defined and this prior corresponds to the most vague information in the beta class of priors. The posterior distribution here is a $\text{Be}(t_{N+1}, M-t_{N+1})$ and is itself an improper density unless $M > t_{N+1} > 0$; i.e., unless both a

$$\int_{\lambda_1}^{\lambda_0} g(\lambda | \lambda_0, \lambda_1) d\lambda = 1 \quad (5.2.14)$$

and hence

$$g(\lambda | \lambda_0, \lambda_1) = \frac{\lambda^{1-\lambda} \int_{\lambda_1}^{\lambda_0} \lambda^{1-\lambda} d\lambda}{\lambda^{1-\lambda} \int_{\lambda_1}^{\lambda_0} \lambda^{1-\lambda} d\lambda} = \frac{\lambda^{1-\lambda} \left(\frac{\lambda_0^{1-\lambda} - \lambda_1^{1-\lambda}}{1-\lambda} \right)}{\lambda^{1-\lambda} \left(\frac{\lambda_0^{1-\lambda} - \lambda_1^{1-\lambda}}{1-\lambda} \right)} \quad (5.2.13)$$

we can define, for $\lambda \in [\lambda_0, \lambda_1]$,

$$0 < \lambda_0 < \lambda_1 < 1 \quad (5.2.12)$$

Now, for any λ_0 and λ_1 such that

$$\lambda^{1-\lambda} (1-\lambda)^{-1} > 0 > \lambda > 1 \quad (5.2.11)$$

$R' = 0$ can be defined. By (5.1.9) this density is proportional to as in the normal case a "conditional" density corresponding to $N' =$ in a proper beta distribution as given by (5.1.9). However, just increasingly vague prior information. Of course, $N'=0$ is not allowed and $\mu_2 \rightarrow 0$ as $N' \rightarrow \infty$, $\mu_2 \rightarrow \mu - \mu^2$ as $N' \rightarrow 0$. Hence as $N' \rightarrow 0$ we have

$$0 \leq \mu_2 < \mu - \mu^2 \quad (5.2.10)$$

This is the usual minimax estimator as given by Ferguson [13] (pages 93-94) and in Lehmann's "Notes on Estimation". For any distribution function G on λ (not necessarily beta) the Bayes risk is

$$\delta_0(t_{N+1}) = \frac{t_{N+1} + \sqrt{M}/2}{M + \sqrt{M}} \quad (5.2.16)$$

Next, consider the estimator "natural" parameter for the binomial distribution. so that ϕ , rather than the usual λ , could easily be thought of as the

$$f(t|\lambda) \propto \binom{M}{t} e^{-t\lambda} \quad (5.2.15)$$

$$t[\ln \lambda - \ln(1-\lambda)]$$

of the exponential family, we find
 tion on the real line and that writing the density for T as a member
 random variable $\phi = \ln \lambda - \ln(1-\lambda)$ has a ("conditional") uniform distribu-
 are intuitively appealing it should be noted that when $N'=R'=0$ the
 prior is slight (see Lindley [21], p. 145). Also, since uniform priors
 in posterior distributions resulting from using a $Be(0,0)$ and a $Be(1,1)$
 a posterior distribution of $Be(t_{N+1}^{N+1}, M-t_{N+1}^{N+1}+1)$ so that the difference
 information, namely a uniform or $Be(1,1)$ distribution, results in
 in passing that a commonly used distribution representing vague prior
 "success" and a "failure" have been observed. It should be noted

$$\lim_{N \rightarrow \infty} \frac{M}{4(\sqrt{M+1})^2} = \frac{1}{4} \quad (5.2.20)$$

at least, unlikely to be known) since

Of course, for large M it is unlikely that (5.2.19) holds (or,

$$\text{preferred if } u-u_1^2 > \frac{144}{25}$$

For example, if $M=1$, δ_0 is preferred if $u-u_1^2 > \frac{16}{1}$ and if $M=25$, δ_0 is

$$\frac{4(\sqrt{M+1})^2}{M} < u-u_1^2 \quad (5.2.19)$$

Now the above is < 1 (i.e., δ_0 preferred to δ_{π}) iff

$$\frac{R(\delta_0, G)}{R(\delta_{\pi}, G)} = \frac{4(\sqrt{M+1})^2 (u-u_1^2)}{M} \quad (5.2.18)$$

To see this, consider

especially for small M one should be aware of this fact.

It should be noted that in some cases δ_0 is preferred to δ_{π} and

$$= \frac{4(\sqrt{M+1})^2}{1}$$

$$= E_{\lambda} \left\{ \frac{(M+\sqrt{M})^2}{M[\lambda(1-\lambda)] + M(\frac{1}{2}-\lambda)^2} \right\}$$

$$R(\delta_0, G) = E_{\lambda} [\delta_0(t) - \lambda]^2 \quad (5.2.17)$$

and we know that $\mu - \mu^1 \leq \frac{1}{4}$. However, for large M even δ^1 has a Bayes risk close to that of δ^G (actually $R(\delta^1, G)$ is only $\frac{M}{M+N^1}$ times as large as $R(\delta^G, G)$ in any case) so that our main concern is with small M . (see section 6.2 for a more complete discussion of δ^0 .)

Since δ^0 has a constant Bayes risk for all G and since δ^0 is Bayes with respect to a $Be(\sqrt{M}/2, \sqrt{M}/2)$ then, by theorem 5, the least favorable distribution in the class of all distribution functions is a $Be(\sqrt{M}/2, \sqrt{M}/2)$.

Therefore, just as in the normal case, the least favorable distribution among all distribution functions is a member of the conjugate class. That is, the conjugate class is not only rich (many different possible distributions) and tractable (e.g., the posterior distribution is also a member of the conjugate class) but is also conservative - in the sense that is we restrict attention to the conjugate class we must have the "worst possible" (i.e., least favorable) distribution still under consideration.

Also, by restricting our attention to the class G of beta priors then any two distribution functions with the same two first moments are identical. Hence, if $M > 1$ there is no problem of identifiability in the class G . (i.e., $f_{G_1} \equiv f_{G_2} \Rightarrow G_1 \equiv G_2$). Throughout the remainder of this chapter then, we will restrict our attention to the class G of beta priors.

The following inequality on the Bayes risk of a Bayes d.f. (assuming a beta prior) is useful:

$$= R(\delta_0, G) \cdot$$

$$(5.2.23) \quad R(\delta_G, G) \leq \frac{1}{1} \cdot \frac{1}{\sqrt{M}} \cdot \frac{1}{\sqrt{M+1}} \cdot \frac{1}{\sqrt{M+1}} \left[\frac{1}{1} \right]^{\frac{1}{2}} = \frac{1}{1}$$

we have

positive and approaches zero as $N' \rightarrow 0$ and as $N' \rightarrow \infty$ then for all N'

implies $N' = \sqrt{M}$ and since the upper bound in (5.2.21) is always

$$(5.2.22) \quad \frac{d}{dN'} \left\{ \frac{1}{1} \cdot \frac{1}{N'} \cdot \frac{1}{N'+1} \cdot \frac{1}{M+N'} \right\} = 0$$

Now, setting

where $N' > 0$, $M \geq 1$.

$$\leq \frac{1}{1} \cdot \frac{1}{N'} \cdot \frac{1}{N'+1} \cdot \frac{1}{M+N'}$$

$$(5.1.10) \quad \text{by } \frac{N'}{N'^2} \cdot \frac{M+N'+1}{N'} =$$

$$(5.2.9) \quad \text{by } \frac{M+N'}{N'^2}$$

$$(5.2.21) \quad R(\delta_G, G) = \frac{M + \frac{1}{N'^2}}{\frac{1}{N'^2}}$$

Then, by the strong law of large numbers:

$$E^{\lambda} = \left(\frac{M(M-1)}{T(T-1)} \right)^{\lambda} E^{\lambda} \quad (5.2.25)$$

$$E \left(\frac{M}{T_1} \right) = E \left(\frac{M}{T_1} \right)^\lambda = 1 \text{ for } 1 \in 1, 2, \dots, N+1 \quad (5.2.24)$$

joint distributions of $\mathbb{T}$ and λ):

the mean of the posterior distribution with prior g and hence $h(g)^N$ is the mean of the posterior with prior g^N . Since we are restricting g to be the beta class, we also require $g^N \in \mathcal{G}$. To find a consistent estimator g^N we need the following expectations (with respect to the

To form an empirical Bayes estimator we will first obtain an estimator G_N of the unknown G from the observations $t_1, \dots, t_{N+1}$. Then, since the Bayes estimator is, in general, some function $h(G)$ of the prior G our empirical Bayes estimator will be $h(G_N)$. For example, in the estimation situation as considered here, $h(G)$ is

(5.2.23) .

Robbins [32] showed that for $M=1$, $R(\delta^G, G) < \frac{1}{16}$ which agrees with

where $M > 1$ and Σ refers to summation over all $N+1$ variables. A definition for $\hat{\mu}_1^2$ that is equivalent to (5.2.28) is:

$$\hat{\mu}_1^2 = \begin{cases} \frac{1}{N+1} \cdot \frac{M(M-1)}{\Sigma t_i(t_i-1)} & \text{if } \hat{\mu}_1 \geq \hat{\mu}_2 \\ \frac{1}{N+1} \cdot \frac{M(M-1)}{\Sigma t_i(t_i-1)} & \text{if } \hat{\mu}_1 < \hat{\mu}_2 \end{cases} \quad \text{otherwise}$$

(5.2.28)

$$\begin{aligned} \hat{\mu}_1 &= \frac{\Sigma t_i}{N+1} \\ \hat{\mu}_1^2 &= \frac{\hat{\mu}_1 - \hat{\mu}_2^2}{\hat{\mu}_1 - \hat{\mu}_2} \\ \hat{R}' &= \frac{\hat{\mu}_1^2 - \hat{\mu}_2^2}{\hat{\mu}_1(\hat{\mu}_1 - \hat{\mu}_2)} \end{aligned}$$

(5.2.27)

with

Hence, as our estimate G_N of G we use a beta $(\hat{R}', \hat{N}' - \hat{R}')$ where

$$\begin{aligned} \hat{N}' &= \frac{1}{N+1} \sum_{i=1}^{N+1} \frac{1}{t_i(t_i-1)} \cdot \frac{M(M-1)}{a.s.} \leftarrow \hat{\mu}_1^2 \\ \hat{R}' &= \frac{1}{N+1} \sum_{i=1}^{N+1} \frac{1}{t_i} \cdot \frac{M}{a.s.} \leftarrow \hat{\mu}_1 \end{aligned}$$

(5.2.26)

which is always true since $\tau_i \in \{0, 1, 2, \dots\}$, $\forall i$. That is, for large N we can be nearly certain that $\hat{\mu}_2 = \frac{1}{N+1} \sum_{i=1}^{N+1} \frac{\tau_i}{\tau_i - 1}$. In any

$$\sum_{i=1}^{N+1} \tau_i > \sum_{i=1}^{N+1} \tau_i^2 \quad (5.2.33)$$

and, as $N \rightarrow \infty$ the above inequality becomes

$$\begin{aligned} & \sum_{i=1}^{N+1} \tau_i > \sum_{i=1}^{N+1} \tau_i^2 \\ & \Rightarrow \sum_{i=1}^{N+1} \tau_i > \sum_{i=1}^{N+1} \tau_i^2 \\ & \Rightarrow \sum_{i=1}^{N+1} \tau_i > \sum_{i=1}^{N+1} \tau_i^2 \\ & \Rightarrow \sum_{i=1}^{N+1} \tau_i > \sum_{i=1}^{N+1} \tau_i^2 \\ & \Rightarrow \sum_{i=1}^{N+1} \tau_i > \sum_{i=1}^{N+1} \tau_i^2 \end{aligned} \quad (5.2.32)$$

only if N is small since

of G will be a degenerate beta at μ . However, this will usually happen

If it turns out that $\hat{\mu}_2 > \frac{1}{N+1} \cdot \frac{M(M-1)}{M(M-1)}$ then our estimate G_N

and the last restriction in (5.2.28) is superfluous.

$$\sum_{i=1}^{N+1} \tau_i > \frac{(N+1)M(M-1)}{(M-1)M(N+1)} - \frac{(N+1)M(M-1)}{(M-1)M(N+1)} \quad (5.2.31)$$

and hence

$$\sum_{i=1}^{N+1} \tau_i > \frac{(N+1)M(M-1)}{(M-1)M(N+1)} + \frac{(N+1)M(M-1)}{(M-1)M(N+1)} \quad (5.2.30)$$

since obviously (for $M > 1$)

$$\hat{\mu}_2 = \max \left\{ \frac{1}{N+1}, \frac{N+1}{M(M-1)} \right\} \quad (5.2.29)$$

Now, by theorem 6 and (5.2.26), it follows that

5.3 for further details).

better than the ordinary estimator δ^T even for small N (see section 5.3. There it appears that δ^N is not only a.o. but usually a Monte Carlo study was undertaken, the results of which are given in $\lim_{N \rightarrow \infty} W(\delta^N, G) = 0$ as is seen below). To evaluate $W(\delta^N, G)$ for small N , difficult to evaluate analytically for a given finite N (although As might be expected, the term $W(\delta^N, G)$ in (5.2.35) is extremely the "penalty" for not knowing G exactly.

with the second term above $W(\delta^N, G)$ being the addition to the risk or

$$\begin{aligned} R(\delta^N, G) &= E_{N,E}^{\lambda, t} \left\{ \delta^N(t) - \delta^G(t) \right\}^2 \\ &= R(\delta^G, G) + E_{N,E}^{\lambda, t} \left\{ \delta^N(t) - \delta^G(t) \right\}^2 \\ &= R(\delta^G, G) + W(\delta^N, G) \end{aligned} \quad (5.2.35)$$

and the global risk is:

$$\begin{aligned} \delta^N(t_{N+1}) &= \frac{M_{\mu_2} + (\mu - \mu_1^2)}{t_{N+1}(\mu_2) + \mu(\mu - \mu_1^2)} \\ &= \frac{M + N}{t_{N+1} + R} \end{aligned} \quad (5.2.34)$$

event, the empirical Bayes estimator to be used is:

as $N \rightarrow \infty$. That is, δ_N is asymptotically optimal. It may not be immediately obvious that, as claimed in section 2.3, the estimators for μ and μ_2 given earlier are essentially the same as those obtained by the method of moments. However, by recalling that

$$E\{\delta_N(t) - \delta_G(t)\}^2 \xrightarrow{0} 0 \quad (5.2.41)$$

it is true that

for all $t \in \{0, 1, \dots, M\}$ then by theorem 3A, Parzen [27], p. 424,

$$\{\delta_N(t) - \delta_G(t)\}^2 \leq 1 \quad (5.2.40)$$

Finally, since

$$\delta_N(t) - \delta_G(t) \xrightarrow{d} 0 \quad (5.2.39)$$

Also, (5.2.38) implies (see Tucker [46], p. 105) that

$$\delta_N(t_{N+1}) - \delta_G(t_{N+1}) \xrightarrow{a.s.} 0 \quad (5.2.38)$$

or, equivalently,

$$\left[\frac{t_{N+1} + R_1}{M + N_1} - \frac{t_{N+1} + R_1}{M + N_1} \right] \xrightarrow{a.s.} 0 \quad (5.2.37)$$

and hence

$$\begin{aligned} \hat{N}_1 &\xrightarrow{a.s.} N_1 \\ \hat{R}_1 &\xrightarrow{a.s.} R_1 \end{aligned} \quad (5.2.36)$$

and R' to be:

After some manipulation we could also show the estimators for N'

essentially the same.

as estimators of μ and μ_2 , respectively. Hence the two methods are

$$\begin{aligned} \hat{\mu}_2 &= \sum_{i=1}^M t_i (t_i - 1) / (N+1)(M-1) \\ \hat{\mu} &= \sum_{i=1}^M t_i / M(N+1) \end{aligned} \quad (5.2.46)$$

or, equivalently,

$$\begin{cases} \hat{\mu}_1 = \frac{M-1}{1} \left[s_2^2 \frac{M}{t} + \frac{M}{t^2} - \frac{M}{t} \right] \\ \hat{\mu}_2 = t/M \end{cases} \quad (5.2.45)$$

we obtain

$$\begin{cases} \hat{\mu}_1 = t \\ \hat{\mu}_2 = s_2^2 \end{cases} \quad (5.2.44)$$

and finally setting

$$\begin{aligned} s_2^2 &= \sum_{i=1}^{N+1} (t_i - \bar{t})^2 / (N+1) \\ \bar{t} &= \sum_{i=1}^{N+1} t_i / (N+1) \end{aligned} \quad (5.2.43)$$

and defining

$$\begin{cases} E(T) = \mu \\ E(T - \mu)^2 = M(M-1)\mu_1^2 - M^2\mu_2 + M\mu \end{cases} \quad (5.2.42)$$

$$R_N(\delta_N', G) = E_{t, \lambda}(\delta_N'(t) - \lambda)_2^2 = E_{\lambda} \left\{ \text{Var}_{t, \lambda}(\delta_N'(t)) + [E_{t, \lambda}(\delta_N'(t) - \lambda)_2]^2 \right\} = \frac{1}{n - \mu_2^2} + \frac{\binom{N+1}{2} \left[\mu_2^2 + \frac{\sum_{i=1}^N t_i^2}{N} \right]}{\binom{M(N+1)}{2}}. \quad (5.2.49)$$

by first finding and apply the usual estimator δ_N' . The global risk for δ_N' is computed That is, we just pool the $N+1$ samples into one sample of size $M(N+1)$ as if the $N+1$ samples had all been chosen from the same distribution.

$$\delta_N'(t_{N+1}) = \sum_{i=1}^{N+1} t_i / (N+1)(M) \quad (5.2.48)$$

to use the estimator global risk smaller than that of δ_N' . To see this, suppose we decide hold even a naive use of the past observations can result in a It is interesting to note that when the Bayesian assumptions

to be such that $\hat{\mu}_2^2 \leq \hat{\mu}_2^2 \leq \mu$. defined above must be modified so that $\hat{N}' \geq 0$ just as we require $\hat{\mu}_2^2$ where we note that $\sum t_i / M(N+1) \leq 1$ so that $\hat{R}' \leq \hat{N}'$. However, $\hat{N}'$ as

$$\hat{N}' = \frac{M(N+1)}{\sum_{i=1}^N (t_i^2) (N')} = \frac{M(N+1) \left[\sum_{i=1}^N (t_i^2 - 1) - (M-1) \sum_{i=1}^N t_i^2 \right]}{M(N+1) (M \sum_{i=1}^N t_i^2 - \sum_{i=1}^N t_i^2)} \quad (5.2.47)$$

Hence the global risk is:

$$(5.2.50) \quad R(\delta_N^1, G) = E_{R_N}(\delta_N^1, G) = \frac{\mu_1 - \mu_2}{2} \left[\frac{M(N+1)}{2} + \frac{N}{2} \right] + \frac{N}{2} \left(\frac{N+1}{N} \right) \frac{(MN)}{(M+N)^2 \mu_2}$$

$$= \frac{R'(N'-R')}{[N'+MN^2+NM+NN']^2} \frac{M(N+1)^2 (N')^2 (N'+1)^2}{(M+N')^2 (N'+MN^2+NM+NN')^2}$$

$$R(\delta_T^1, G) = \frac{[N+1]^2_{2N'}}{N'+MN^2+NM+NN'} \quad R(\delta_G^1, G) = \frac{N'M(N+1)^2}{(M+N')^2 (N'+MN^2+NM+NN')^2}$$

The estimator δ_N^1 is not asymptotically optimal since

$$(5.2.51) \quad \lim_{N \rightarrow \infty} R(\delta_N^1, G) = \frac{N'}{M+N'} \quad R(\delta_G^1, G) > R(\delta_G^1, G) \quad .$$

It should also be noted that δ_N^1 is not, in general, unbiased for λ_{N+1}

since

$$(5.2.52) \quad E_N(\delta_N^1(t_{N+1}^N)) = \mu$$

although this is no real drawback since also

$$E(\delta_G^1(t_{N+1}^N)) = \frac{\lambda_{N+1} \cdot M+R'}{M+N'} \neq \lambda_{N+1} \quad .$$

mean μ is defined by:

Now, since $\mu \in \left[\frac{1 - \sqrt{1 - 4\mu_2}}{2}, \frac{1 + \sqrt{1 - 4\mu_2}}{2} \right]$, then it is reasonable to require $\hat{\mu}$ to be in the same interval. Hence, our estimator $\hat{\mu}$ of the prior

$$\frac{1}{2} - \frac{\sqrt{1 - 4\mu_2}}{2} \leq \mu \leq \frac{1}{2} + \frac{\sqrt{1 - 4\mu_2}}{2} \quad (5.2.56)$$

or, equivalently,

$$\mu_2 - \mu + \mu^2 \leq 0 \quad (5.2.55)$$

we must be careful since from (5.1.18) we have with mean $E t_1 / (N+1)M$ and variance μ_2 as our estimate G_N of G . However, be denoted by G' . In this case, one would be tempted to use a beta variance is μ_2 . The class of all beta priors with variance μ_2 will Next, suppose we have the additional information that the prior section 5.3 for some numerical comparison of the Bayes risks. is better than δ_N^T for all N regardless of how small N may be. See Hence, if the prior $Be(R', N' - R')$ is such that $N' > M$ then δ_N^T

$$\begin{aligned} & \Leftrightarrow M(N^2 + N) < N'(N^2 + N) \Rightarrow M < N' \\ & \Rightarrow N' + MN^2 + NM + NN' < N^2N' + 2NN' + N' \\ & \frac{N' + MN^2 + NM + NN'}{(N+1)^2 N'} < 1 \end{aligned} \quad (5.2.54)$$

δ_N^T is better than δ_N^T if and only if estimator than δ_N^T . In fact we have $R(\delta_N^T, G) < R(\delta_N^T, G)$, that is, However, even this simple-minded approach often yields a better

i.e., the distribution that assigns all probability to the points $\lambda=0$ and $\lambda=1$ with equal weights at 0 and 1. And in the case that μ_2 is

$$G_0(\lambda) = \begin{cases} 1/2 & \text{if } 0 < \lambda < 1 \\ 1 & \text{if } \lambda = 1 \end{cases} \quad (5.2.59)$$

$G_0(\lambda)$ where

possible variance, we know that $G(\lambda)$ is approaching the distribution short length for large μ_2 . This is so since as $\mu_2 \rightarrow 1/4$, the maximum

interval $\left[1 - \sqrt{1 - 4\mu_2}, 1 + \sqrt{1 - 4\mu_2}\right]$, from which μ is chosen, has a estimator $\hat{\mu}$ of μ . However, from (5.2.58) and figure 4 see that the

knowledge that μ_2 is small usually gives one more confidence in the

$$\left| \frac{1}{2} + \sqrt{1 - 4\mu_2} - \left(\frac{1}{2} - \sqrt{1 - 4\mu_2} \right) \right| = \sqrt{1 - 4\mu_2} \quad (5.2.58)$$

between $\hat{\mu}$ and the true μ is:

Note that when using the above $\hat{\mu}$ the maximum possible difference

$$\hat{\mu} = \frac{\sum t_i}{(N+1)M} \quad \text{if} \quad \frac{\sum t_i}{(N+1)M} < \frac{1 - \sqrt{1 - 4\mu_2}}{2} \\ \text{if} \quad \frac{\sum t_i}{(N+1)M} > \frac{1 + \sqrt{1 - 4\mu_2}}{2} \\ \left[\frac{\sum t_i}{(N+1)M}, \frac{1 - \sqrt{1 - 4\mu_2}}{2} \right], \left[\frac{1 + \sqrt{1 - 4\mu_2}}{2}, \frac{\sum t_i}{(N+1)M} \right] \quad (5.2.57)$$

Next assume μ is known but μ_2 is unknown. Denote this class of priors by G'' . In this case, our G_N is chosen to be a beta with mean μ and second moment: (assuming $M > 1$)

See (5.2.36) to (5.2.41) for details. can be shown in exactly the same manner used to show δ_N is a.o. in G . various values of N and different G . The fact that δ_N is a.o. in G' A Monte Carlo evaluation of $W(\delta_N, G)$ is given in section 5.3 for $T_1, \dots, T_{N+1}$ each with p.m.f. $f_G(t)$.

where E represents expectation with respect to the random variables

$$W(\delta_N, G) = E \left\{ \delta_N''(t) - \delta(t) \right\}^2 \quad (5.2.62)$$

where

$$R(\delta_N, G) = R(\delta, G) + W(\delta_N, G) \quad (5.2.61)$$

and the global risk is

$$\delta_N''(t_{N+1}) = \frac{M\mu_2 + \mu - \mu_1}{t_{N+1}\mu_2 + \mu(\mu - \mu_1)} \quad \text{where } \hat{\mu}_1 = \mu_2 + \mu_2$$

(5.2.60)

and variance μ_2 we find an empirical Bayes estimator of Using the estimator G_N , i.e., a beta with mean defined by (5.2.57)

risk is attainable.

known to be $= 1/4$ then of course we also know $\mu = 1/2$ and the Bayes

Hence if μ is either very small or very large then the estimator $\delta_{N'}^{T, G}$ should have a value very close to that of the Bayes estimator δ_G .

$$(5.2.65) \quad |\hat{\mu}_2 - \mu_2| = \mu - \mu_2.$$

maximum possible difference between $\hat{\mu}_2$ and the true μ_2 is:

$$\text{and } \frac{R(\delta_{N'}^{T, G}, G)}{R(\delta_G^{T, G}, G)} = N'. \quad \text{Note that by defining } \hat{\mu}_2 \text{ by (5.2.63) the}$$

If $M=1$ then we find that $R(\delta_{N'}^{T, G}, G) = \mu_2$ so that $R(\delta_G^{T, G}, G) = \frac{1+N'}{N'}$

$$(5.2.64) \quad \delta_{N'}^{T, G} = \frac{\hat{\mu}_2 + \mu - \mu_2}{\hat{\mu}_2 - \mu_2} M + \mu - \mu_2$$

If $M > 1$

μ if $M = 1$

The empirical Bayes estimator in this case is then just where $\hat{\mu}_2$ is required to lie in the interval $[\mu_2, \mu]$ since $\mu_2 \in [\mu_2, \mu]$.

$$(5.2.63) \quad \hat{\mu}_2 = \mu_2 \text{ if } \frac{1}{\sum t_i(t_i-1)} \frac{(N+1)}{M(M-1)} \leq \mu$$

$$\text{if } \mu_2 > \mu \text{ then } \frac{1}{\sum t_i(t_i-1)} \frac{(N+1)}{M(M-1)} > \mu$$

$$\text{if } \mu_2 < \mu \text{ then } \frac{1}{\sum t_i(t_i-1)} \frac{(N+1)}{M(M-1)} < \mu$$

where $\alpha=1-\beta$ and, using Lindley's notation,

$$(5.2.68) \quad \frac{\bar{F}(R' + t)}{\bar{F}(R' + t)} > \lambda > \frac{(N' - R' + M - t) + \bar{F}(R' + t)}{(N' - R' + M - t) + \bar{F}(R' + t)}$$

interval for λ can be seen to be:

is $F \sim F[2(R' + t), 2(N' - R' + M - t)]$. Hence an exact β Bayesian confidence interval for λ is an F -distribution with parameters $2(R' + t)$ and $2(N' - R' + M - t)$. That

$$(5.2.67) \quad \tilde{F} = \frac{N' - R' + M - t}{\lambda} \left(\frac{1 - \lambda}{\lambda} \right)$$

Lindley [21], p. 141, the posterior distribution of

Bayes confidence intervals on the parameter λ_{N+1} . By corollary 1,

just as in the normal case it is possible to obtain empirical

in equations (5.2.36) to (5.2.41).

be shown to be asymptotically optimal by following the steps outlined

various small values of N and different prior G . Also, $\delta_{N+1}^G(t)$ can
Again, section 5.3 contains Monte Carlo evaluations of $R(\delta_{N+1}^G, G)$ for

$$(5.2.66) \quad R(\delta_{N+1}^G, G) = R(\delta_{N+1}^G, G) + W(\delta_{N+1}^G, G) \cdot$$

The global risk of δ_{N+1}^G is:

$$\begin{aligned} \sum_{t=1}^T (t-1) \frac{M(M-1)}{t} &> \mu \text{ then } \delta_{N+1}^G(t_{N+1}) = \frac{M}{t_{N+1}} = \delta_{N+1}^G(t_{N+1}) \cdot \\ \text{Also, if } \sum_{t=1}^T (t-1) \frac{M(M-1)}{t} &< \mu \text{ then } \delta_{N+1}^G(t_{N+1}) = \text{and if } \frac{N+1}{1} \cdot \end{aligned}$$

The main purposes of this study are:
relating to the estimators described in the previous section are given.
In this section the results of several Monte Carlo experiments

5.3 Numerical Results

val is asymptotically optimal.
converge to the end points of (5.2.59) and hence the confidence interval
be β . However, the end points of the EB confidence interval will
confidence level of the EB confidence interval will not, in general,
replacing R' and N' in (5.2.69) with their estimates $\hat{R}'$ and $\hat{N}'$. The
and an empirical Bayes confidence interval for λ can be formed by

$$\bar{F}' = \left\{ \bar{F} [2(R'+t), 2(N'-R'+M-t)] \right\}^{-1}$$

$$= \bar{F} [2(N'-R'+M-t), 2(R'+t)]$$

where

$$\bar{F} = \bar{F}_{\alpha/2} [2(R'+t), 2(N'-R'+M-t)]$$

$$\bar{F}' = \bar{F}_{\alpha/2} [2(R'+t), 2(N'-R'+M-t)]$$

or, since $\bar{F}_{\alpha}(x_1, x_2) = \left\{ \bar{F}_{\alpha}(x_2, x_1) \right\}^{-1}$, (5.2.68) can be written:

$$\frac{1}{\bar{F}} = \frac{1}{\bar{F}'} > \lambda > \frac{1}{\bar{F}''}$$

$$\bar{F}'' = \frac{\bar{F}'(N'-R'+M-t)}{(R'+t)} + 1$$

$$\bar{F}' = \frac{\bar{F}(R'+t)}{N'-R'+M-t} + 1 \quad (5.2.69)$$

$$\begin{aligned}
 & \frac{P_{N+1,M-1}^{N+1}(t) = \text{no. of terms of } t_1, \dots, t_{N+1}^{N+1}}{N+1} \\
 & \frac{P_{N+1,M}^{N+1}(t) = \text{no. of terms of } t_1, \dots, t_{N+1}^{N+1}}{N+1} \quad (5.3.2)
 \end{aligned}$$

where

$$\delta_S^N(t_{N+1}) = t_{N+1}^M \cdot \frac{P_{N+1,M-1}^{N+1}(t_{N+1}^{N+1})}{P_{N+1,M}^{N+1}(t_{N+1}^{N+1})} \quad (5.3.1)$$

is defined by:

$\delta_G^N(t)$ directly. This estimator was first given by Robbins [33] and Bayes estimator because it makes no attempt to form an explicit estimate G_N of G . Rather, δ_S^N bypasses this point and attempts to estimate The second one, denoted by δ_S^N , is called the "simple" empirical in section 3.3.

Found using Deely and Kruse's method of estimating G by G_N as described

additional ones are studied here. The first one, denoted by δ_D^N , is

In addition to the estimators considered in section 5.2, two given in section 3.3 are applicable here and will not be repeated. The general comments concerning methods used and other items

priors and different values of M and N .

(iii) to compare the various proposed estimators for different

(ii) to study the behavior of the estimators for small N

estimator when an exact determination is difficult

(i) to closely approximate the global risk for each proposed

and t_i is the number of successes in the first $M-1$ out of M trials

in the i th sample ($i=1, \dots, N+1$).

Robbins does not point out that his proposed estimator δ_N^S can be greater than 1. For example suppose, $N=2$, $M=5$ and

$$(5.3.3) \quad \left\{ \begin{array}{l} t_1 = 3, \quad t_2 = 3, \quad t_3 = 2 \\ t_1 = 3, \quad t_2 = 3, \quad t_3 = 2 \end{array} \right.$$

then,

$$(5.3.4) \quad \left\{ \begin{array}{l} p_{3,5}(t_3+1) = p_{3,5}(3) = 2 \\ p_{3,4}(t_3) = p_{3,4}(2) = 1. \end{array} \right.$$

Hence,

$$(5.3.5) \quad \delta_N^S(t_{N+1}) = \delta_N^S(2) = \frac{5}{3} \cdot \frac{1}{2} = \frac{5}{6} > 1.$$

That is, we are estimating the binomial parameter λ to be greater than 1. The obvious modification of δ_N^S to correct this defect (and the one used in this study) is to simply restrict the values of δ_N^S

to the interval $[0, 1]$.

The prior distributions considered are all beta distributions

with parameters:

- (1) $N' = 20, R' = 10$
- (11) $N' = 6, R' = 3$
- (111) $N' = 2, R' = 1$
- (1V) $N' = 24, R' = 6$

in section 3.3.

comparison. The tables also give $R(\delta) = R(\delta_G, G) / R(\delta, G)$ as defined note the numbers given are actually 10 times the risk for ease of and if these numbers are absent then the risk given is exact. Also $M=1, 5, 10$. Again, the numbers in parentheses are standard errors risks for the different prior distributions and for $N=1, 10, 20$ and The following tables (tables 5-10) present the main results on Kruse's method (again, $N_1=3$ in this study).

(ix) δ_D - empirical Bayes estimator obtained by using Deely and

$$R(\delta_N, G) = \frac{R'(N'-R)(N'+MN^2+NM+NN')}{M(N+1)2(N'+1)}$$

(viii) δ_N - a non-a.o. EB estimator (pooled mean)

(vii) δ_S - Robbins simple EB estimator

(vi) δ_N - empirical Bayes estimator, variance known

(v) δ_N - empirical Bayes estimator, mean known

(iv) δ_N - empirical Bayes estimator, G assumed conjugate

$$(iii) \delta_0 - \text{the minimax estimator; } R(\delta_0, G) = \frac{4(\sqrt{M+1})^2}{1}$$

$$(ii) \delta_I - \text{the ordinary mean with } R(\delta_I, G) = \frac{MN'(N'+1)}{R'(N'-R')}$$

$$(i) \delta_G - \text{the Bayes estimator with } R(\delta_G, G) = \frac{R'(M+N')(N'+1)}{R'(N'-R')}$$

The estimators under consideration are:

$$(vi) N' = 4, R' = 1$$

$$(v) N' = 12, R' = 3$$

The figures (3-8) following the tables show $R(\delta)$ as a function of N and allow us simultaneously to compare two estimators for a small N and to see the rate of convergence of $R(\delta_N, G)$ to $R(\delta, G)$, assuming the estimator is asymptotically optimal). Just as in chapter III, however, the figures will not contain the curves of $R(\delta_N)$ and $R(\delta_N')$ since, as is obvious from the tables, they lie above $R(\delta)$ and they are not applicable unless we assume the prior mean or variance to be known.

Some points to keep in mind when reading the tables and figures are:

(1) $\delta_N = \delta_N' = \delta_N''$ for $M = 1$. This implies that δ_N is not asymptotically optimal in this case but this is not surprising since it was shown earlier that no asymptotically optimal empirical Bayes estimator exists in this situation ($M=1$).

(11) δ_N is Bayes with respect to a $B_e(\sqrt{M}/2, \sqrt{M}/2)$. Hence, for certain $M, R(\delta_N)$ may be very close to 1 although it should be remembered that it is also possible for $R(\delta_N)$ to be quite small.

(111) $\delta_N'' = \mu$ for $M = 1$. This fact is implicit in the definition of δ_N'' but is pointed out here since $\delta_N'' = \mu$ for all t implies $R(\delta_N''), G) = \mu^2$.

(iv) $\lim_{N \rightarrow \infty} R(\delta_N) = \frac{M+N}{N}$

The number N_0 is then the number of prior observations necessary for us to "break-even" using δ instead of the usual δ_{π} . If N_0 is very large for a given decision function δ , the fact that δ is asymptotically optimal may mean very little in practice. Approximations to the value of N_0 were obtained from this study by simply reading the point of intersection of the curves $R(\delta)$ and $R(\delta_{\pi})$ from the following figures. The results are given in table 9.

$$(5.3.7) \quad R(\delta) \geq R(\delta_{\pi}).$$

or equivalently, the smallest N such that

$$(5.3.6) \quad R(\delta, G) \leq R(\delta_{\pi}, G)$$

smallest integer N such that

it is interesting to know the value N_0 where N_0 is defined to be the

For any given empirical Bayes estimator δ and a given prior G

tag disappears.

δ'_N is better than δ_{π} for small M although as M increases this advan-

distribution considered. As a matter of fact, by (v) above, even

studies is that $R(\delta'_N, G) \leq R(\delta_{\pi}, G)$ for all N and for every prior

An obvious, but nonetheless important fact concerning these

$$(v) \quad \delta'_N \text{ is preferred over } \delta_{\pi} \text{ if } M \leq N'. \\ (vi) \quad \delta_0 \text{ is preferred over } \delta_{\pi} \text{ if } \frac{4(\sqrt{M+1})^2}{M} \leq \frac{4(\sqrt{M+1})^2}{M}.$$

TABLE 3

GLOBAL RISKS* ($N'=20$, $R'=10$)

Estimator	N=1			N=10			N=20		
	M=1	M=5	M=10	M=1	M=5	M=10	M=1	M=5	M=10
1) Bayes	.113	.095	.079	.113	.095	.079	.113	.095	.079
2) Usual (mean)	2.381	.476	.238	2.381	.476	.238	2.381	.476	.238
3) Minimax	.625	.239	.144	.625	.239	.144	.625	.239	.144
4) G conjugate	1.250	.297	.169	.325	.150	.112	.227	.129	.105
5) G conjugate, mean known	.119	.236 (.048)	.151 (.040)	.119	.138 (.029)	.108 (.028)	.119	.118 (.031)	.093 (.029)
6) G conjugate, variance known	.674 (.087)	.268 (.045)	.158 (.037)	.312 (.034)	.146 (.033)	.106 (.027)	.205 (.032)	.121 (.029)	.097 (.028)
7) Pooled mean	1.250	.297	.178	.324	.151	.129	.226	.136	.124
8) Robbins' simple EB	1.250	1.338 (.101)	.954 (.098)	.324	.653 (.072)	.277 (.039)	.226	.423 (.044)	.239 (.032)
9) Deely and Krusse	1.307 (.111)	.588 (.049)	.218 (.027)	.629 (.052)	.296 (.038)	.193 (.028)	.413 (.038)	.183 (.035)	.177 (.021)

* Entries are actually 10 times risk for ease of comparison

TABLE 4

GLOBAL RISKS ($N' = 6$, $R' = 3$)

Estimator	N=1			N=10			N=20		
	M=1	M=5	M=10	M=1	M=5	M=10	M=1	M=5	M=10
1) Bayes	.306	.199	.133	.306	.194	.133	.306	.194	.133
2) Usual (mean)	2.142	.428	.214	2.142	.428	.214	2.142	.428	.214
3) Minimax	.625	.238	.144	.625	.238	.144	.625	.238	.144
4) G conjugate	1.250	.389 (.031)	.200 (.031)	.519	.314 (.043)	.189 (.042)	.442	.281 (.037)	.172 (.036)
5) G conjugate, mean known	.357	.329 (.029)	.186 (.030)	.357	.306 (.031)	.174 (.028)	.357	.274 (.036)	.161 (.035)
6) G conjugate, variance known	.675 (.089)	.278 (.024)	.179 (.020)	.487 (.046)	.279 (.035)	.177 (.031)	.429 (.037)	.260 (.034)	.160 (.036)
7) Pooled mean	1.250	.392	.285	.519	.363	.344	.442	.360	.350
8) Robbins' simple EB	1.250	1.215 (.100)	.906 (.085)	.519	.762 (.095)	.343 (.034)	.442	.558 (.076)	.276 (.029)
9) Deely and Krusse	1.329 (.101)	.590 (.051)	.247 (.022)	.657 (.051)	.322 (.042)	.207 (.025)	.459 (.041)	.296 (.043)	.193 (.017)

TABLE 5

GLOBAL RISKS ($N' = 2$, $R' = 1$)

Estimator	N+1			N=10			N=20		
	M=1	M=5	M=10	M=1	M=5	M=10	M=1	M=5	M=10
1) Bayes	.555	.238	.138	.555	.238	.138	.555	.238	.138
2) Usual (mean)	1.666	.333	.166	1.666	.333	.166	1.66	.333	.166
3) Minimax	.625	.238	.144	.625	.238	.144	.625	.238	.144
4) G conjugate	1.250	.323 (.042)	.165 (.032)	.909	.290 (.039)	.158 (.035)	.873	.259 (.038)	.150 (.032)
5) G conjugate, mean known	.833	.308 (.050)	.164 (.032)	.833	.287 (.038)	.155 (.029)	.833	.273 (.041)	.145 (.029)
6) G conjugate, variance known	1.123 (.108)	.254 (.033)	.158 (.028)	.846 (.073)	.243 (.032)	.153 (.027)	.810 (.058)	.244 (.038)	.143 (.021)
7) Pooled mean	1.250	.583	.500	.909	.787	.772	.873	.809	.801
8) Robbins' simple EB	1.250	1.041 (.133)	.926 (.079)	.909	.761 (.113)	.356 (.042)	.873	.659 (.100)	.300 (.033)
9) Deely and Kruse	1.357 (.112)	.609 (.052)	.273 (.021)	.713 (.068)	.282 (.052)	.218 (.011)	.484 (.030)	.336 (.047)	.205 (.011)

TABLE 6

GLOBAL RISKS ($N' = 24$, $R' = 6$)

Estimator	N=1			N=10			N=20		
	M=1	M=5	M=10	M=1	M=5	M=10	M=1	M=5	M=10
1) Bayes	.072	.062	.052	.072	.062	.052	.072	.062	.052
2) Usual (mean)	1.800	.360	.180	1.800	.360	.180	1.800	.360	.180
3) Minimax	.625	.238	.144	.625	.238	.144	.625	.238	.144
4) G conjugate	.937 (.027)	.210 (.018)	.105 (.018)	.231 (.013)	.100 (.007)	.064 (.007)	.157 (.011)	.088 (.011)	.062 (.015)
5) G conjugate, mean known	.075	.138 (.024)	.102 (.016)	.075	.094 (.081)	.063 (.011)	.075	.084 (.071)	.060 (.014)
6) G conjugate, variance known	.508 (.070)	.204 (.026)	.103 (.017)	.201 (.024)	.079 (.009)	.056 (.007)	.151 (.025)	.063 (.008)	.054 (.012)
7) Pooled mean	.937	.217	.127	.231	.100	.084	.157	.088	.080
8) Robbins' simple FB	.937	.586 (.077)	.516 (.052)	.231	.333 (.073)	.332 (.043)	.157	.205 (.034)	.307 (.031)
9) Deely and Krusse	.958 (.070)	.281 (.025)	.142 (.018)	.303 (.023)	.187 (.032)	.138 (.013)	.219 (.026)	.126 (.018)	.120 (.013)

TABLE 7

GLOBAL RISKS ($N' = 12$, $R' = 3$)

Estimator	N=1			N=10			N=20		
	M=1	M=5	M=10	M=1	M=5	M=10	M=1	M=5	M=10
1) Bayes	.133	.101	.078	.133	.101	.078	.133	.101	.078
2) Usual (mean)	1.730	.346	.173	1.730	.346	.173	1.730	.346	.173
3) Minimax	.625	.238	.144	.625	.238	.144	.625	.238	.144
4) G conjugate	.937	.209	.154	.288	.151	.109	.219	.139	.104
		(.028)	(.026)		(.025)	(.017)		(.728)	(.016)
5) G conjugate, mean known	.144	.198	.144	.144	.147	.101	.144	.136	.092
		(.026)	(.022)		(.019)	(.018)		(.023)	(.012)
6) G conjugate, variance known	.683	.193	.143	.294	.148	.093	.202	.129	.089
	(.063)	(.025)	(.022)	(.038)	(.023)	(.019)	(.021)	(.024)	(.081)
7) Pooled mean	.937	.245	.158	.288	.162	.146	.219	.153	.145
8) Robbins' simple EB	.937	.903	.847	.288	.604	.251	.219	.404	.214
		(.085)	(.083)		(.056)	(.028)		(.030)	(.016)
9) Deely and Kruse	1.011	.321	.162	.351	.305	.159	.281	.200	.133
	(.088)	(.025)	(.017)	(.023)	(.031)	(.013)	(.027)	(.021)	(.011)

TABLE 8

GLOBAL RISKS ($N' = 4$, $R' = 1$)

	N=1			N=10			N=20		
	M=1	M=5	M=10	M=1	M=5	M=10	M=1	M=5	M=10
1) Bayes	.300	.166	.107	.300	.166	.107	.300	.166	.107
2) Usual (mean)	1.500	.300	.150	1.500	.300	.150	1.500	.300	.150
3) Minimax	.625	.238	.144	.625	.238	.144	.625	.238	.144
4) G conjugate	.937	.257 (.024)	.148 (.024)	.477	.214 (.020)	.139 (.021)	.428	.178 (.029)	.120 (.020)
5) G conjugate, mean known	.652	.237 (.023)	.141 (.021)	.456	.204 (.025)	.131 (.019)	.412	.168 (.023)	.118 (.013)
6) G conjugate, variance known	.652 (.080)	.237 (.019)	.141 (.022)	.456 (.045)	.204 (.021)	.131 (.021)	.412 (.038)	.168 (.019)	.118 (.010)
7) Pooled mean	.937	.337	.262	.477	.368	.354	.428	.371	.364
8) Robbins' simple EB	.937	.909 (.080)	.856 (.077)	.477	.721 (.062)	.310 (.053)	.428	.502 (.048)	.249 (.020)
9) Deely and Krusse	1.084 (.076)	.349 (.031)	.171 (.019)	.377 (.027)	.278 (.037)	.167 (.012)	.310 (.021)	.237 (.030)	.152 (.012)

TABLE 9

VALUES OF N (M = 5)

Prior	Estimator			
	Empirical Bayes	Pooled mean*	Robbins	Deely & Kruse
1) $N'=20$, $R'=10$	1	1	16	3
2) $N'=6$, $R'=3$	1	1	35	5
3) $N'=2$, $R'=1$	1		250	7
4) $N'=24$, $R'=6$	1	1	9	1
5) $N'=12$, $R'=3$	1	1	28	1
6) $N'=4$, $R'=1$	1		165	8

* Recall that the pooled mean is better than the usual estimator if and only if $M > N'$.

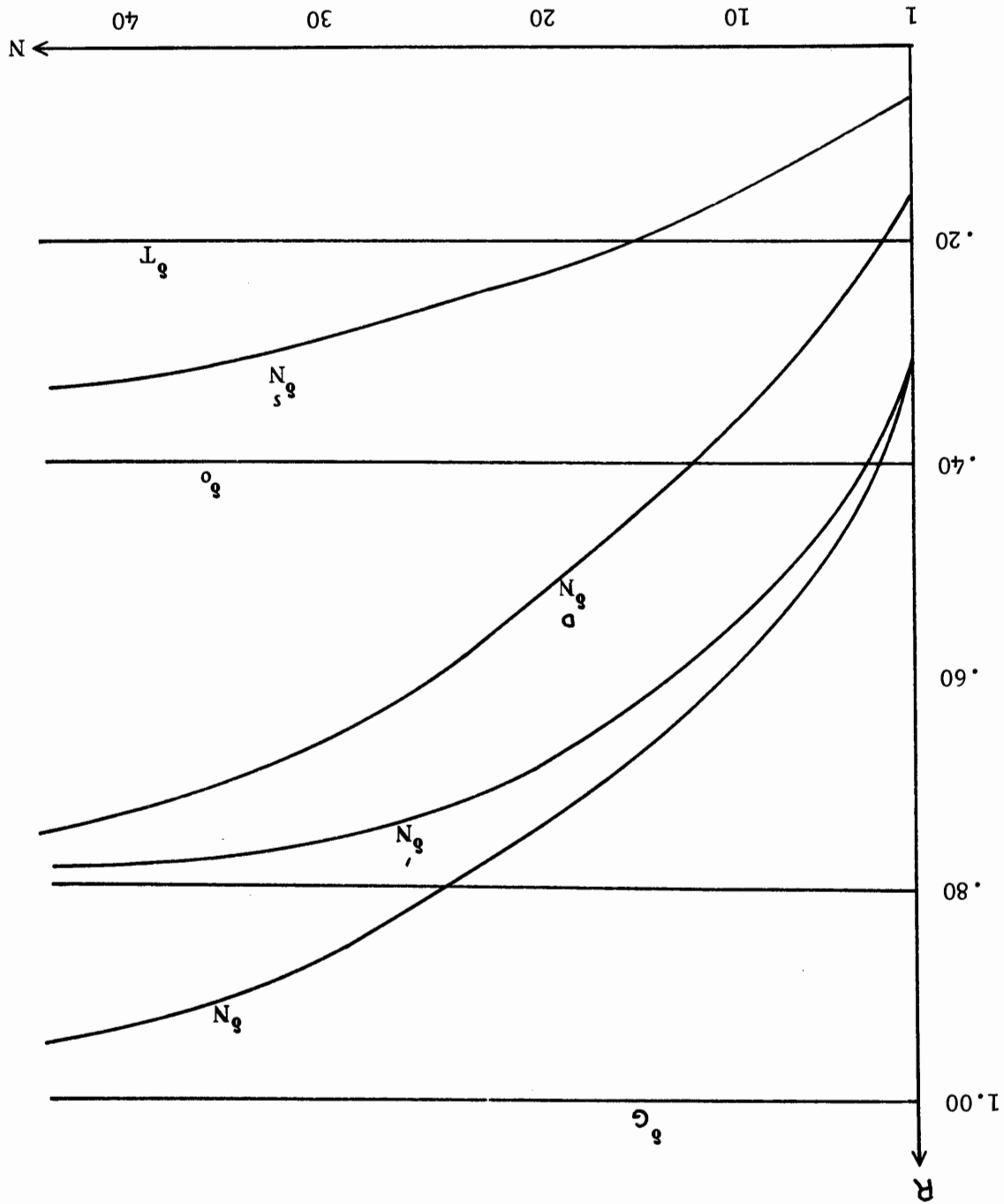
Figure 5. $R(\delta)$ as a function of $N(N'=20, R'=10, M=5)$.

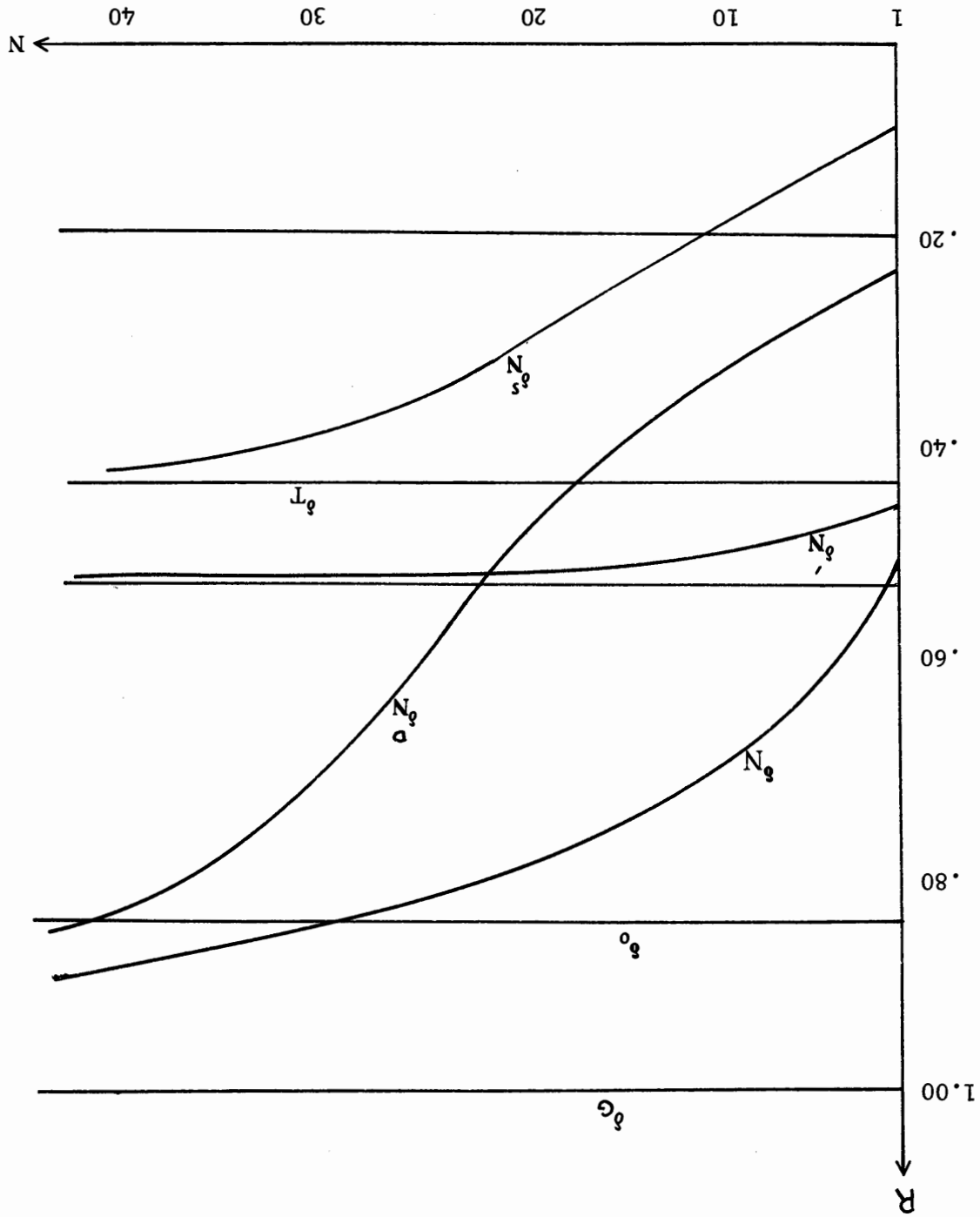
Figure 6. $R(\delta)$ as a function of N ($N'=6$, $R'=3$, $M=5$).

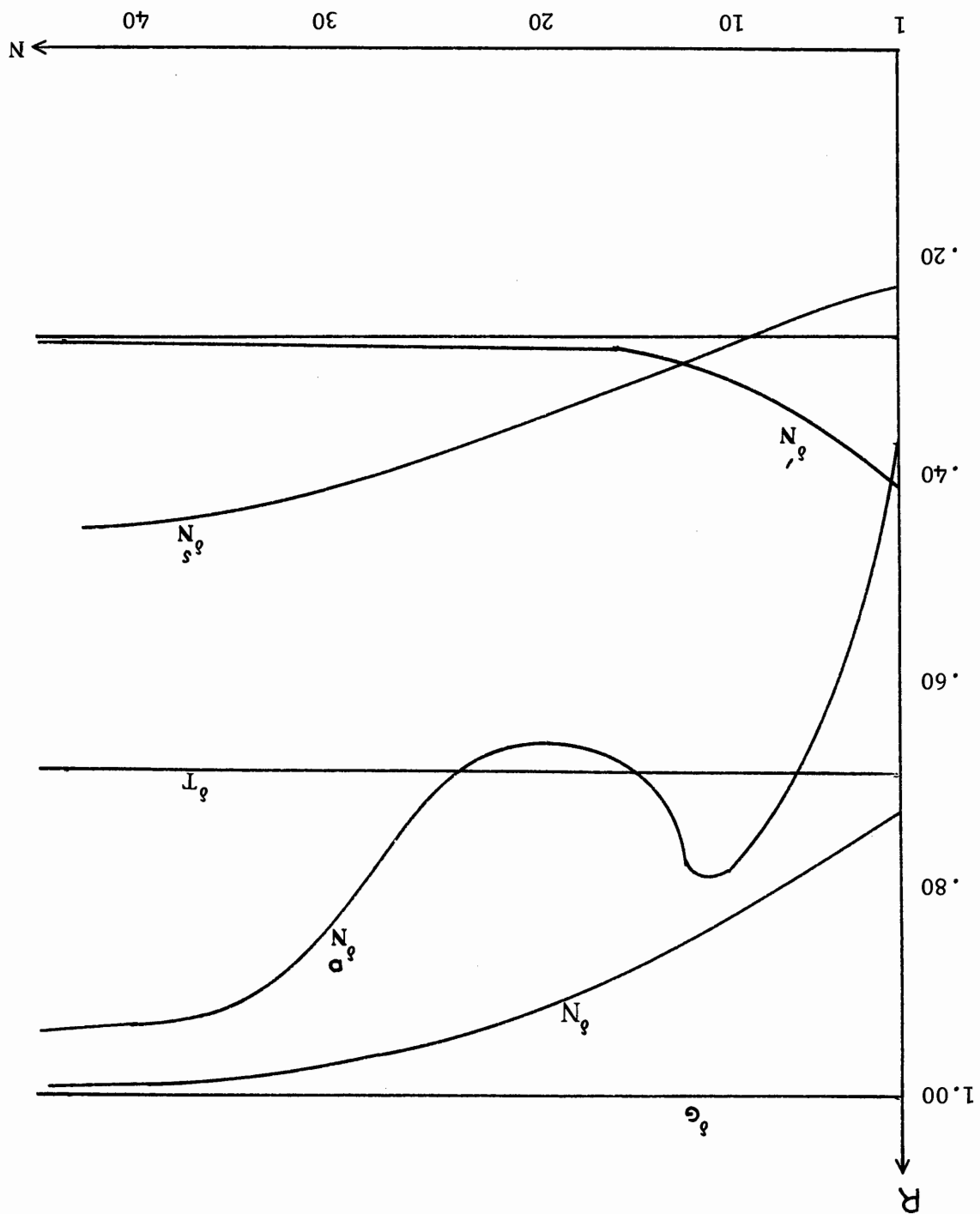
Figure 7. $R(\delta)$ as a function of $N(N'=2, R'=1, M=5)$.

Figure 8. $R(\delta)$ as a function of $N(N'=24, R'=6, M=5)$.

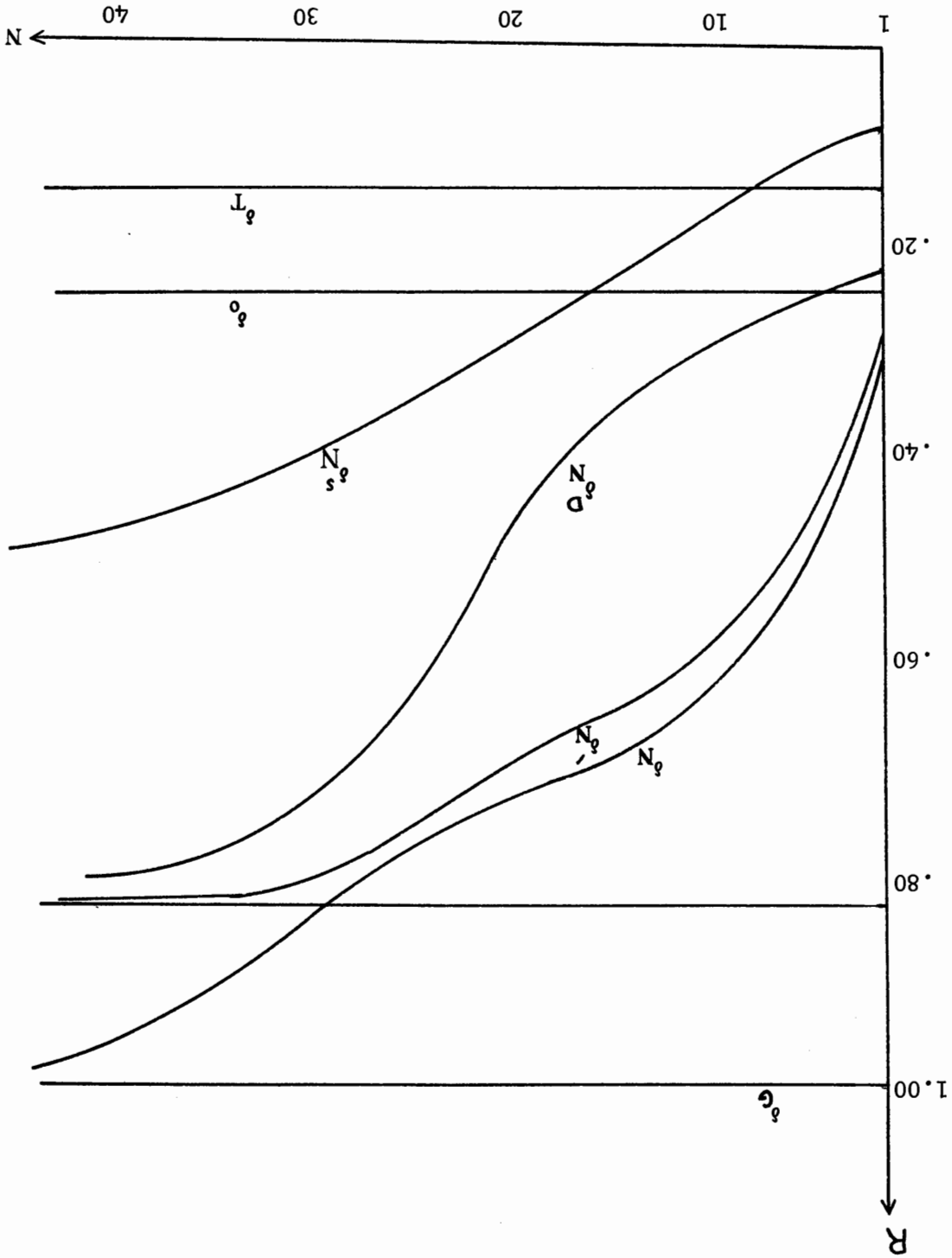


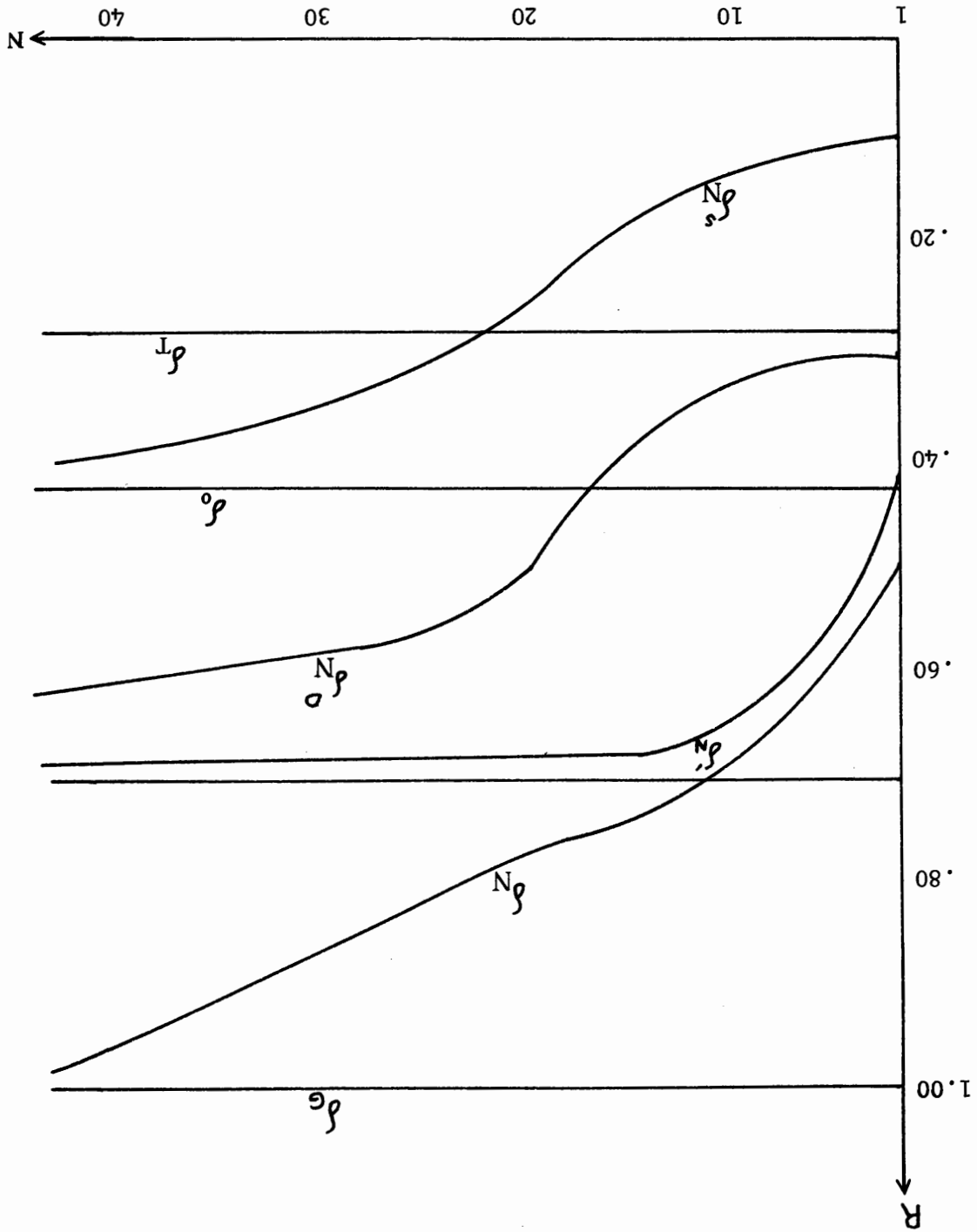
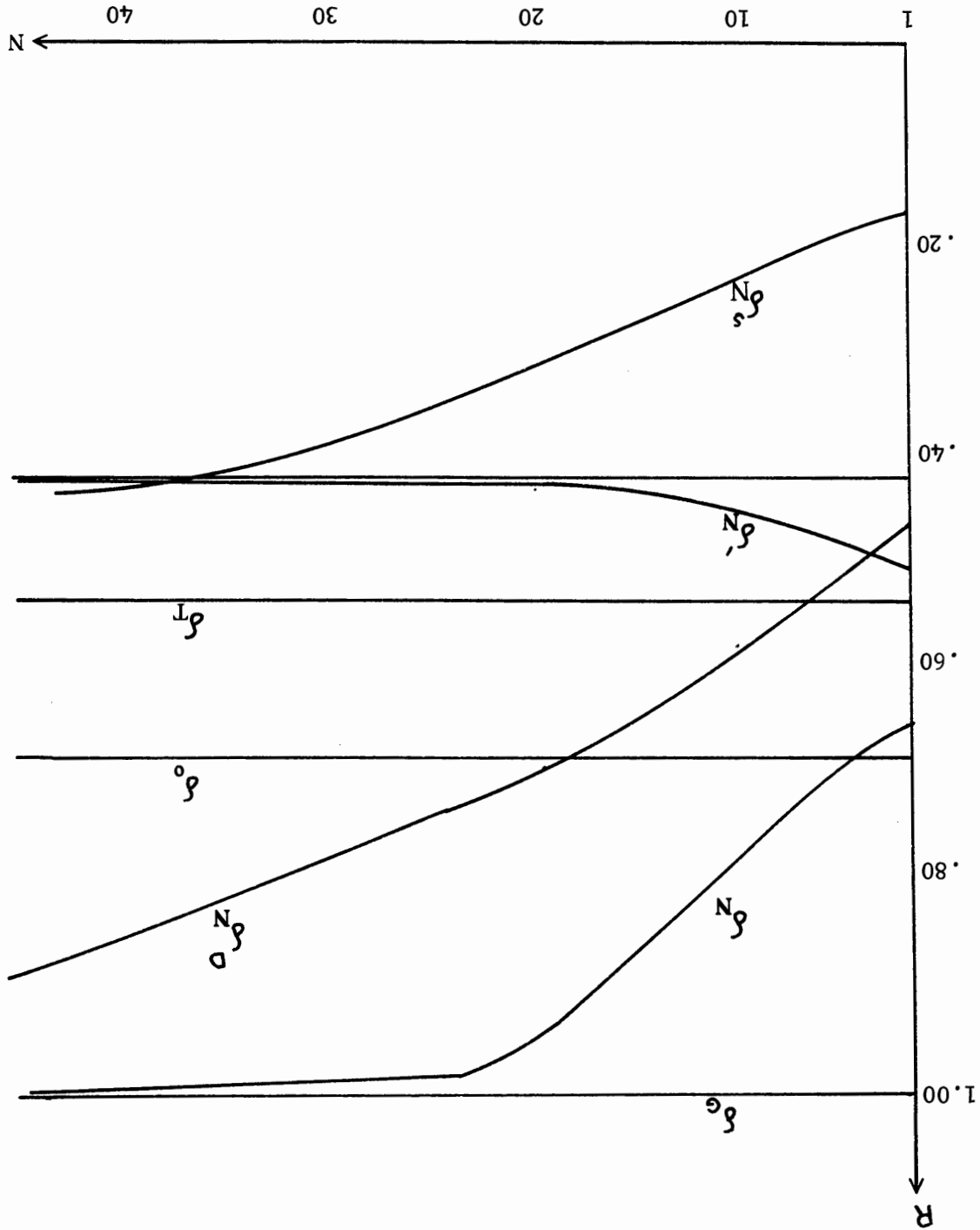
Figure 9. $R(\delta)$ as a function of $N(N'=12, R'=3, M=5)$.

Figure 10. $R(\delta)$ as a function of $N(N'=4, R'=1, M=5)$.

6.1 Introduction

As in Chapter IV we consider here the case that G is known only to be a member of some class $\mathcal{G}$ of distribution functions on Λ and no prior observations on $F_G(x)$ exist. The applicable theory is developed in section 2.4. The classes $\mathcal{G}$ to be considered are:

- (1) The class of all probability distributions functions on Λ .
- (11) $\mathcal{G}$: the conjugate class.
- (111) $\mathcal{G}$: the conjugate class, μ known.
- (1v) $\mathcal{G}$: μ known.
- (v) $\mathcal{G}$: μ and μ_2 known.

In each case θ -minimax estimators of λ are found and all are shown to be admissible. In addition, it is shown below that cases (i) and (ii) and cases (iii) and (iv) yield identical estimators so that only three distinct estimators are found in the five cases.

Asymptotic θ -minimax estimators are given in section 6.4.

The first case to be considered, $\mathcal{G}$ unrestricted, is given in this introductory section because of its simplicity. When $\mathcal{G}$ is the class of all probability distribution functions on Λ then obviously

$$\delta^0(t_{N+1}^0) = \frac{M + \sqrt{M}}{t_{N+1}^0 + \sqrt{M}/2} = \frac{\sqrt{M}(1 + \sqrt{M})}{t_{N+1}^0} + \frac{2(1 + \sqrt{M})}{1} \quad (6.2.1)$$

is just the ordinary minimax estimator

all degenerate distributions, the G -minimax estimator in this case

this is the degenerate distribution at λ^0 . Hence, since G contains

$R_0^0 = \lambda^0 N_0^0$ will have a mean of λ^0 and variance (by (5.1.10)) of

For any given $\lambda^0 \in [0, 1]$, a $\text{Be}(R_0^0, N_0^0 - R_0^0)$ distribution with

class of all beta distributions by G .

function in the beta class with parameters R' and $N' - R'$. Denote the

family of distributions on $\lambda = [0, 1]$. That is, G is a distribution

Suppose it is known only that G is a member of the conjugate

6.2 Estimation of λ_{N+1}^0 : G Conjugate

section.

This particular estimator is discussed in greater detail in the next

$$\frac{M + \sqrt{M}}{t_{N+1}^0 + \sqrt{M}/2} \quad (6.1.1)$$

mator is just the ordinary minimax estimator

G contains all degenerate distributions and hence the G -minimax esti-

as described earlier in (5.2.16) and (6.1.1).

Since the risk (see (5.2.17)) of δ_0 is constant for all $G \in \mathcal{G}$

the conclusion that δ_0 is G -minimax also follows from theorem 5.

And, since δ_0 is Bayes with respect to a $Be(\sqrt{M}/2, \sqrt{M}/2)$, denoted

by G_0 , and

$$(6.2.2) \quad R(\delta_0, G_0) = \frac{1}{4(\sqrt{M}+1)^2} = R(\delta_0, G) \quad \forall G \in \mathcal{G}$$

then it also follows from theorem 3 that δ_0 is G -minimax.

Further, since the estimator obtained when G is conjugate is

the same as that obtained with no restrictions on G , then this estima-

tor is obviously both G -admissible and admissible.

It should be noted that

$$(6.2.3) \quad \sup_{G \in \mathcal{G}} R(\delta, G) > \frac{1}{4(\sqrt{M}+1)^2}$$

since δ_0 is G -minimax and from (5.2.23)

$$(6.2.4) \quad R(\delta, G) < R(\delta_0, G) \quad \forall G \in \mathcal{G}.$$

That is, the Bayes risk of δ_0 , namely (6.2.2), is the maximum risk

for a Bayes d.f. with respect to any $G \in \mathcal{G}$. Recall (see 5.2.19) that

δ_0 is better than δ_I for prior distributions with large variances

and for samples with a small number (M) of observations. A rewriting

of (5.2.19) allows us to state that δ_0 is better than δ_I if

$$(6.2.5) \quad \frac{1}{1 + \frac{1}{\sqrt{M}}} < \frac{1}{4(\mu - \mu_I)^2}$$

For example, if $u - u_1' = \frac{4}{1}$ then δ_0 is better than δ_1 iff

$$(6.2.6) \quad \frac{1}{2} < \frac{\left(\frac{1 + \frac{1}{2}}{2} \right)^{\sqrt{M}}}{1}$$

which is true for all M . If $u - u_1' = \frac{8}{1}$ then δ_0 is better than δ_1 iff

$$(6.2.7) \quad \frac{1}{2} < \frac{\left(\frac{1 + \frac{1}{2}}{2} \right)^{\sqrt{M}}}{1} \text{ i.e., } M < 5.$$

In general, if

$$(6.2.8) \quad u - u_1' = \alpha \quad \text{where} \quad 0 < \alpha < \frac{4}{1}$$

then δ_0 is preferred to δ_1 iff:

$$(6.2.9) \quad \frac{1}{2} < \frac{\left(\frac{1 + \frac{1}{2}}{2} \right)^{\sqrt{M}}}{1}$$

or, equivalently, iff

$$(6.2.10) \quad M < \left(\frac{1 - 4\alpha}{8\alpha} \right)^2 + \frac{4\alpha}{1 - 4\alpha}.$$

In terms of α , δ_0 is preferred to δ_1 if and only if:

$$(6.2.11) \quad \alpha > \frac{4M + 8\sqrt{M} + 4}{M} = \frac{4 \left[1 + \frac{2}{\sqrt{M}} + \frac{1}{M} \right]}{M}$$

Figure 11 graphically illustrates the preference regions.

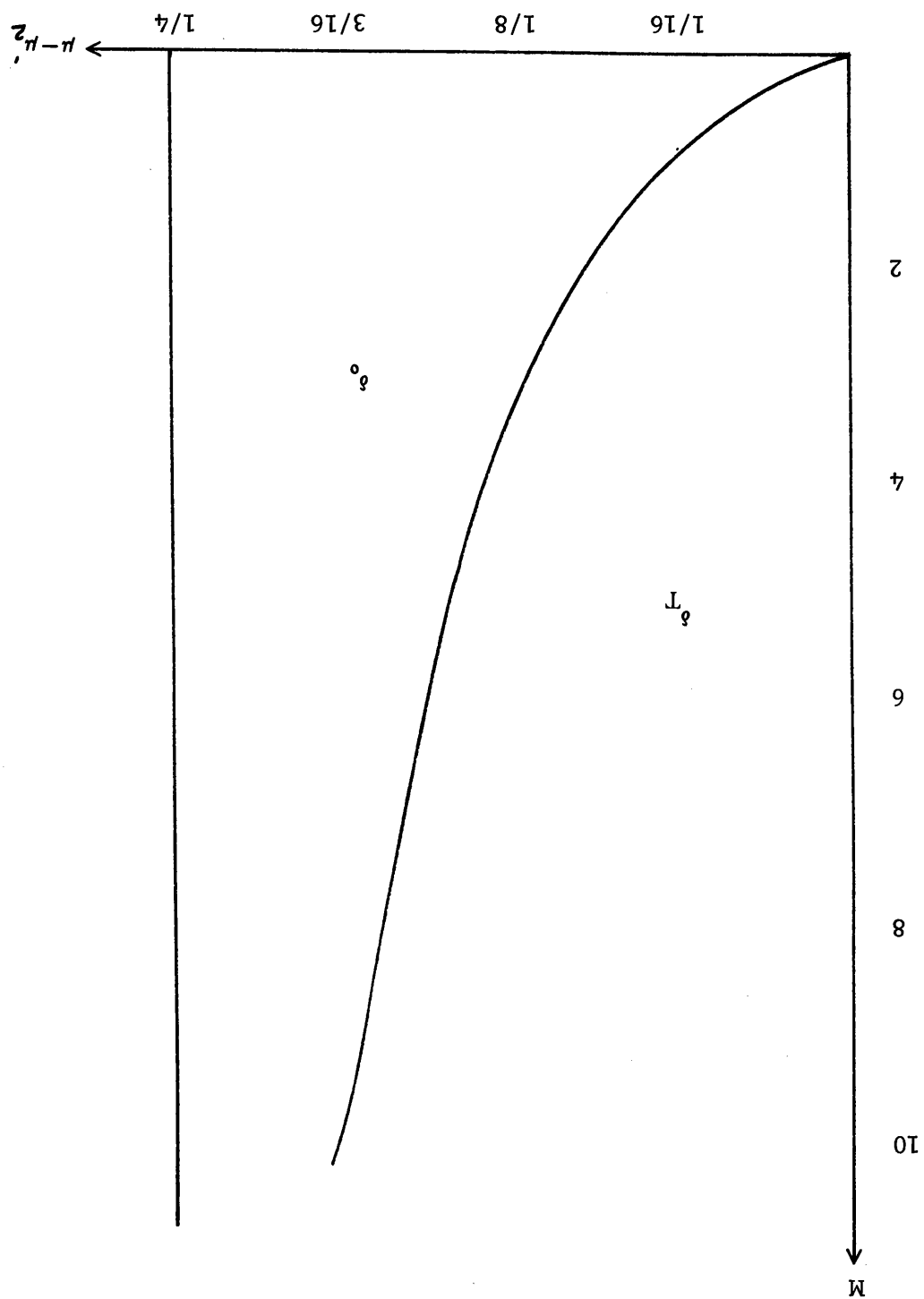


Figure 11. Preference regions for δ_0 and δ_T

Of course, M can take only integral values but for purposes of

illustration we have considered (6.2.11) for all real $M > 0$.

Next, suppose that the mean, say μ_0 , is also known and let

G' denote the class of all conjugate prior distributions with mean

μ_0 . That is, G' is the class of all $Be(\mu_0 N', (1-\mu_0) N')$ distributions.

Although G' neither contains all degenerate distributions nor

satisfies the sufficient condition for admissibility given after

(2.4.13), it is shown below that G' -admissibility implies admissi-

bility.

The Bayes estimator for any $G \in G'$ is of the form (see (5.2.7))

$$\delta_G^G(t_{N+1}) = \frac{M+\alpha}{t_{N+1} + \mu_0 \alpha} \quad (6.2.12)$$

where $\alpha = N'$. For a given α , and any $G' \in G'$, not necessarily G , we have

$$R(\delta_G, G') = E_\lambda \{ \text{Var}_G(\delta_G(t)) \} + E_\lambda \left\{ \frac{M+\alpha}{\lambda M + \mu_0 \alpha} - \lambda \right\} \quad (6.2.13)$$

$$= \mu_0 \alpha - \mu_0 \alpha^2 - M \mu_0^2 + \alpha^2 \mu_0^2$$

$$(M+\alpha)^2$$

and in order that the above Bayes risk be independent of the particular

G' in effect we must have

$$\alpha^2 - M = 0 \quad \text{i.e., } \alpha = \sqrt{M}. \quad (6.2.14)$$

The estimator

$$\delta_1^G(t_{N+1}) = \frac{t_{N+1} + \sqrt{M} \mu_0}{M + \sqrt{M}} \quad (6.2.15)$$

is Bayes with respect to $Be(\mu_0, \sqrt{M}(1-\mu_0))$, say G_1 , and has constant

Bayes risk

$$R(\delta_1, G) = \frac{M\mu_0 - M\mu_0^2}{(M + \sqrt{M})^2} = \frac{\mu_0 - \mu_0^2}{(\sqrt{M+1})^2} \quad (6.2.16)$$

over the class G' . Hence, by theorem 5, δ_1 is G' -minimax and G_1 is least favorable in G' . It is obvious that

$$\inf_{\delta} R(\delta, G) = R(\delta_1, G) = \sup_{G \in G'} \inf_{\delta} R(\delta, G) < \sup_{G \in G'} R(\delta, G) \quad (6.2.17)$$

i.e., that

$$\sup_{G \in G'} R(\delta, G) > \frac{\mu_0 - \mu_0^2}{(\sqrt{M+1})^2} \quad (6.2.18)$$

for every $\delta \neq \delta_1$.

As a comparison between δ_1 and δ_1^* , consider the ratio

$$\frac{R(\delta_1^*, G)}{R(\delta_1, G)} = \frac{\mu_0 - \mu_0^2}{(\sqrt{M+1})^2} \cdot \frac{M}{\mu_0 - \mu_0^2} \quad (6.2.19)$$

This ratio is ≤ 1 (i.e., δ_1 is better than δ_1^*) iff

Hence, even for reasonably large M the required N' (required to make δ_1

$$(6.2.23) \quad N' \geq \frac{(\sqrt{M} + 1)^2}{2\sqrt{M} + 1} - 1 = \frac{M}{2\sqrt{M} + 1}$$

that is,

$$(6.2.22) \quad \frac{1}{2\sqrt{M} + 1} < \frac{N' + 1}{(\sqrt{M} + 1)^2}$$

so that the inequality (6.2.20) becomes

$$(6.2.21) \quad \mu_2 = \frac{N' + 1}{\mu_0(1 - \mu_0)} \quad \text{for all } G \in G'$$

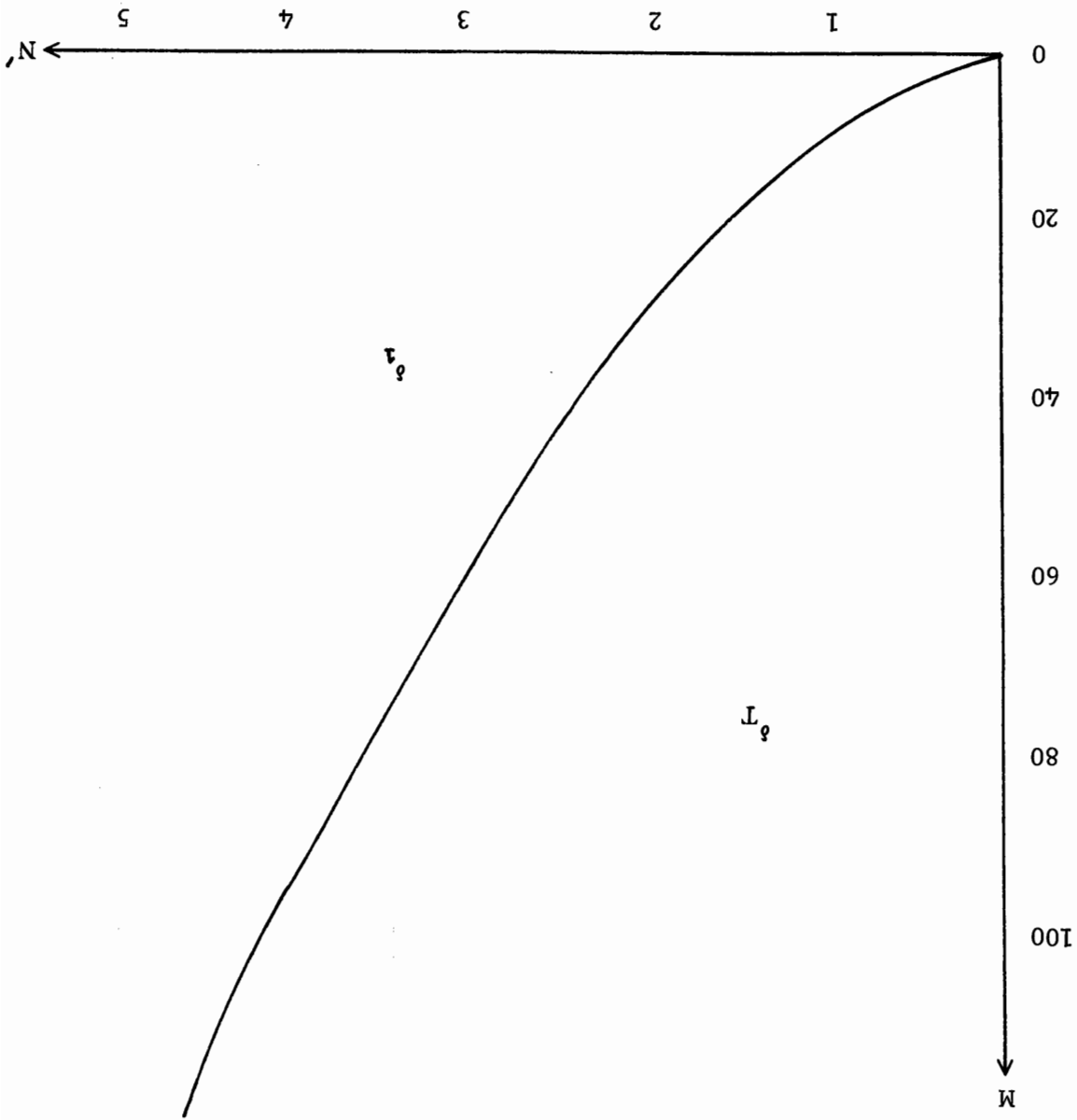
But from (5.1.20) we have

$$\Rightarrow \mu_2 \leq \left\{ \frac{2\sqrt{M} + 1}{(\sqrt{M} + 1)} \right\} \mu_0(1 - \mu_0)$$

$$\Rightarrow \mu_2 \leq \mu_0 - \mu_0^2 + \mu_0^2 \left[\frac{\sqrt{M} + 1}{\sqrt{M}} \right] - 1$$

$$\Rightarrow \mu_2 \leq \mu_0 - \mu_0^2 + \mu_0^2 \left[\frac{\sqrt{M} + 1}{\sqrt{M}} \right] - 1$$

$$(6.2.20) \quad \frac{\mu_0 - \mu_0^2}{M} \leq \frac{\mu_0 - \mu_0^2}{1}$$

Figure 12. Preference regions for δ_I and δ_J 

And, since δ_1 is Bayes with respect to a $\text{Be}(\sqrt{M_0}, \sqrt{M(1-\mu_0)})$ distribution, say G_0 , then by theorem 5, δ_1 is G_0 -minimax and G_0 is least favorable in G'' .

Just as in the normal theory case, the above result further justifies the use of the conjugate class of distributions as the class of prior distributions. Although $G' \subset G''$, the least favorable distribution in the class G'' is also in the class G' so the assumption $G \in G'$ does not enable us to do any better than the assumption $G \in G''$.

To see that G'' -admissibility implies admissibility let λ_0 be some point in Λ . Without loss of generality we can assume that $\lambda_0 < \mu_0$ and write

$$\mu_0 - \lambda_0 = \alpha > 0. \quad (6.3.2)$$

Also, define λ_1 by

$$\lambda_1 = \mu_0 + \frac{1-\epsilon}{\alpha\epsilon} \quad (6.3.3)$$

where $0 < \epsilon < 1$.

Then, for any probability distribution function G on Λ such that

$$G(\lambda_0) - G(\lambda_1) = \epsilon \quad (6.3.4)$$

$$G(\lambda_1) - G(\lambda_1^-) = 1-\epsilon$$

we have

$$\int \lambda dG(\lambda) = \lambda_0 \epsilon + \lambda_1 (1-\epsilon) = \mu_0. \quad (6.3.5)$$

in G'' .

G_2 . Hence, by theorem 5, δ_2 is G'' -minimax and G_2 is least favorable for all $G \in G''$ and is Bayes with respect to a prior $Be(R', N', R')$, say

$$(6.3.8) \quad R(\delta_2, G) = (n - n_1^2) \left\{ M + \frac{n_2}{n - n_1^2} \right\} \frac{\left[M + \frac{n_2}{n - n_1^2} \right]^2}{2}$$

has a constant Bayes risk of

$$(6.3.7) \quad \delta_2(t_{N+1}) = \frac{M + (n - n_1^2)}{n_2} \frac{n_2}{n(n - n_1^2)}$$

n_2 by G'' . The estimator

of all probability distribution functions with mean μ_0 and variance

Next, suppose that μ_0 and n_2 are both known and denote the class

Ferguson [13], p. 62, theorem 3, we know that δ_1 is admissible.

(2.4.13) we know that G'' -admissibility implies admissibility. By

above argument holds for an arbitrary $\lambda^0 \in \Lambda$, by the comments following

as required by (5.1.18) then we see that $G \in G''$. Finally, since the

$$(6.3.6) \quad \int_V (\lambda - \mu^0) dG(\lambda) = (\lambda^0 - \mu^0)^2 \epsilon + (\lambda^1 - \mu^0)^2 (1 - \epsilon) \\ = \frac{\alpha^2 \epsilon}{1 - \epsilon} < \overline{\mu - \mu^0}$$

And if we choose ϵ so that

of prior observations $t_1, \dots, t_N$ on $f_G(t)$ it may be desirable (at least

As pointed out at the end of section 2.4, even in the presence

6.4 Asymptotic G-Minimax Decision Functions

considerably more so than knowledge of the mean alone.

Hence knowledge of the mean and variance of the prior is quite valuable,

That is, δ_2 is better than the usual estimator δ_{π} for all M and all N .

$$(6.3.11) \quad R(\delta_2, G) = \frac{\frac{M + \pi - \pi^2}{N}}{\pi - \pi^2} > \frac{M}{\pi - \pi^2} = R(\delta_{\pi}, G).$$

Also note that

attain the minimum possible Bayes risk.

simple Bayesian estimation problem with G completely known and could

determined by its first two moments. That is, we would be in the

then G is completely specified since a beta distribution is completely

If it is also known that G is a member of the conjugate class

consequently, we also have the fact that δ_2 is admissible.

then we have the fact that G -admissibility implies admissibility.

$$(6.3.10) \quad \epsilon = \frac{\alpha_2 + \pi^2}{\pi^2}$$

that is,

$$(6.3.9) \quad \frac{\alpha_2 \epsilon}{1 - \epsilon} = \pi^2$$

and alter (6.3.6) to read

In addition, if we refer back to expressions (6.3.2) to (6.3.5)

for small N) to construct asymptotically G -minimax decision functions instead of finding asymptotically optimal EB decision functions.

Suppose that it is known that G is a class of priors and there exist N past observations $\{t_i\}$ on $f_G(t)$. Then, analogous to the δ_1

estimator defined by (6.2.15) we can write:

$$(6.4.1) \quad \delta_0(t_{N+1}) = \frac{M + \sqrt{M}}{t_{N+1} + \sqrt{M} \mu_0}$$

where μ_0 is defined by:

$$(6.4.2) \quad \mu_0 = \frac{1}{N+1} \sum_{i=1}^{N+1} t_i \cdot$$

It was shown earlier that

$$(6.4.3) \quad \mu_0 \xleftarrow{\text{a.s.}} \mu$$

and hence, by theorem 6,

$$(6.4.4) \quad \delta_0(t_{N+1}) \xleftarrow{\text{a.s.}} \delta_0(t_{N+1}) = \frac{M + \sqrt{M}}{t_{N+1} + \sqrt{M} \mu_2}$$

Now, $\delta_0(t_{N+1})$ can be rewritten

$$(6.4.5) \quad \delta_0(t_{N+1}) = t_{N+1} \left(1 + \frac{1}{t_{N+1} \sqrt{M}} \right) \frac{M + \sqrt{M}}{M + \sqrt{M}}$$

where

$$(6.4.6) \quad \mu_1 = \sum_{i=1}^N t_i / M(N+1)$$

and the global risk involved is found to be:

we can write (assuming $M > 1$)

Next, analogous to the G'' -minimax estimator δ_2 given by (6.3.7)

so that even in this case the difference is slight.

$$R(\delta_0, G) - R(\delta_1, G) = \frac{8}{\mu - \mu^2} < \frac{1}{32} \quad (6.4.11)$$

$M=N=1$ we find

decreases rapidly as either M or N increases. And, in the extreme case so that the maximum possible difference between the two risk functions

$$\frac{4(\sqrt{M+1})^2(N+1)}{1} < \frac{(M+\sqrt{M})^2(N+1)}{R(\delta_0, G) - R(\delta_1, G) = (M+N')\mu^2} \quad (6.4.10)$$

For a given prior $G \in G''$ and a fixed M and N we have

That is, δ_0 is asymptotically G'' -minimax.

$$\lim_{N \rightarrow \infty} R(\delta_0, G) = R(\delta_1, G) \quad \forall G \in G'' \quad (6.4.9)$$

Therefore,

$$E_{N+1}^{N+1}(\mu_0 - \mu_0)^2 = \text{Var}_T(\mu_0) = \frac{(M+N')R'(N'+1)M}{(N+1)(N')^2(N'+1)M} \quad (6.4.8)$$

where

$$R(\delta_0, G) = R(\delta_1, G) + \left\{ \frac{\sqrt{M}}{2} + \frac{\sqrt{M+1}}{2} \right\} E_{N+1}^{N+1}(\mu_0 - \mu_0)^2 \quad (6.4.7)$$

That is, δ_1^2 is an asymptotically G''-minimax decision function.

$$\lim_{N \rightarrow \infty} R(\delta_1^2, G) = R(\delta_2^2, G) \quad \text{AGEG''''} \quad (6.4.17)$$

it follows that

$$|\delta_1^2(t_{N+1})| \leq 1 \quad \forall N \text{ and all } t_{N+1} \quad (6.4.16)$$

then, since

$$\left\{ \delta_1^2(t_{N+1}) - \lambda \right\}_{a.s.} \left\{ \delta_2^2(t_{N+1}) - \lambda \right\} \quad (6.4.15)$$

Now, $\delta_1^2(t_{N+1})$ a.s. $\delta_2^2(t_{N+1})$ and therefore

$$\hat{\mu}_2 = \max \left\{ \hat{\mu}_2, \left(\frac{1}{N+1} \right) \sum_{i=1}^N t_i(t_i-1) \right\} \quad (6.4.14)$$

$$\hat{\mu} = \sum_{i=1}^{N+1} t_i \quad (6.4.14)$$

and

$$\left\{ \begin{array}{l} \hat{N}' = \frac{\hat{\mu}_2}{\hat{\mu} - \hat{\mu}_2} \\ \hat{R}' = \frac{\hat{\mu}_2}{\hat{\mu}(\hat{\mu} - \hat{\mu}_2)} \end{array} \right. \quad (6.4.13)$$

where

$$\frac{\delta_1^2(t_{N+1})}{t_{N+1} + \hat{R}'} = \frac{M + \hat{N}'}{M + \hat{N}'} \quad (6.4.12)$$

APPENDIX

COMPUTER PROGRAMS

All of the computer work in this paper was performed on an IBM 360 model 44 with the programs written in Fortran IV language. The programs used are listed below with brief explanations.

Program 1. RANDU subroutine

This well-known program is designed to generate a random observation from a $U(0,1)$ distribution. The method used is the power residue method described in IBM manual C20-8011, "Random number generation and testing".

Program 2. GAUSS subroutine

This program generates a normally distributed random variable with a given mean μ and standard deviation σ (i.e., a $N(\mu, \sigma^2)$) and requires the RANDU subroutine listed above. It should be noted that this program is not very efficient since it requires 12 uniform random numbers to compute 1 normal random number but it was found to be adequate for the purposes of this paper.

Program 3. The main program

The purposes of this program are:

- 1) to generate a random sample $\lambda_1, \dots, \lambda_{N+1}$ from $G(\lambda)$. In the beta

example given here this is accomplished by use of the well-known fact that the k th order statistic in a sample of N from a $U(0,1)$ distribution has a $Be(K, N-K+1)$ distribution.

(i) to generate observations $t_1, \dots, t_{N+1}$ from $f(t|\lambda_1), \dots, f(t|\lambda_{N+1})$ respectively.

(ii) to compute each decision function under consideration at t_{N+1} .

(iv) to estimate (or calculate directly) the global risk for each decision function.

The beta example given here has the following parameters:

Input

NP, RP - prior parameters; $Be(RP, NP-RP)$.

M - number of observations in each t_i (i.e., $t_i \sim B_i(M, \lambda_i)$).

N - the total number of past and current observations (the same as $N+1$ in the text of the paper).

$IREP$ - number of replications.

IX - entry value for RANDU subroutine.

Output

$R(I)$ - uniform random numbers $I=1, \dots, N$.

$T(I, j)$ - the j th observation in the I th sample. $I=1, \dots, N$; $j=1, \dots, M$ $T1=0$ or 1 .

$T(I)$ - the I th binomial observation. $I=1, \dots, N$.

$B(I)$ - beta random numbers. $B(I) \equiv \lambda_i$.

$\{EST1, \dots, EST8 \}$ the various estimators and their respective risks (either exact or estimated) in the following order:

1. δ_T 2. δ 3. δ_D 4. δ_N 5. δ_S 6. δ 7. δ 8. δ 9. δ_N

Program 4. Linkage program

This program provides the link between programs 3 and 5 by computing the entry values for program 5.

The K λ_j 's needed to compute G_K were found by evaluating δ_T at $t[1]$ and $t[N+1]$, the smallest and largest order statistics in the sample, and setting

$$\lambda_j = \delta_T(t[1]) + \frac{(j-1) | \delta_T(t[N+1]) - \delta_T(t[1]) |}{K-1} \quad (A.1)$$

$$j=1, \dots, K, K < N$$

In the binomial case, for example, $\delta_T(t)=t$ and hence

$$\lambda_j = \frac{(K-j)t[1] + (j-1)t[N+1]}{(K-1)}$$

$$\lambda_1 = t[1]$$

$$\lambda_2 = t[1] + \frac{t[N+1] - t[1]}{K-1}$$

$\vdots$

$$\lambda_{K-1} = t[1] + \frac{K-2}{K-1} (t[N+1] - t[1])$$

$$\lambda_K = t[N+1]$$

Parameters computed

XLAMDA - these are the λ_j described above.

A(I,J) - elements of the payoff matrix.

AS(I,J) - linear programming matrix computed directly from the

payoff matrix.

BS(I) - the $\bar{b}$ vector in the linear constrain equation $A\bar{x} = \bar{b}$

(A is the AS matrix).

C(I) - the $\bar{c}$ vector where $\bar{c}'\bar{x}$ is to be minimized.

Program 5. SIMPLE subroutine

In general, this program can supply a simplex solution to a linear

programming problem. In the present context, however, it is used to

compute the step function approximation G^K to the unknown G .

Before calling this subroutine, JH(J) and X(J) must be zeroed out.

Input parameters:

AS,BS,C - described in program 4 explanation.

INFILAG = 0 if no basis provided

1 if basis provided

MX - number of rows of linear programming matrix (4.N).

NN - number of columns of linear programming matrix (4N+K); the

total number of variables.

Output parameters

F - Player B's strategies. The first K of these values are the

desired weights $h_1, \dots, h_K$ at the points $\lambda_1, \dots, \lambda_K$.

PS - Player A's strategies.

```

Listings
Program 1 RANDU
SUBROUTINE RANDU (IX,IY,YFL)
  IX=IX*65539
  IF (IY) 5,6,6
  IY=IY+2147483647+1
  YFL=IY
  6 YFL=IY
  YFL=YFL*.4656613E-9
  RETURN
END
Program 2 GAUSS
SUBROUTINE GAUSS (IX,S,AM,V)
  A=0.0
  DΦ50I=1,12
  CALL RANDU (IX,IY,Y)
  IX=IY
  A=A+Y
  50 V=(A-6.0)*S+AM
  RETURN
END

```

C Program 3 The main program

```

DIMENSION R(100),T(100),TL(100)
DIMENSION EST1(150),EST2(150),EST3(150),EST4(150),EST5(150)
DIMENSION XLOSS4(150),XLOSS6(150),XLOSS7(150),XLOSS8(150)
DIMENSION R(100),T(100),TL(100,10)
NP = 12
RP = 3
M = 5
N = 2
N IS THE TOTAL NO. OF CURRENT AND PAST OBSERVATIONS
ZNP = NP
ZM = M
ZN = N
IRP = RP
NPL1 = NP - 1
V1 = RP/ZNP
V2 = RP*(RP + 1.0)/(ZNP*(ZNP+1.0))
VAR = V2 - V1**2
WRITE(6,19)
19 FORMAT(//7X,'BETA PRIOR DIST. = R(R,N-R) ')
WRITE(6,62) RP,ZNP
62 FORMAT(//7X,'R = ',F10.4,'10X,'N = ',F10.4)
WRITE(6,20) V1,VAR
20 FORMAT(//7X,'MEAN = ',F15.6,'10X,'VARIANCE = ',F15.6)
WRITE(6,900)
900 FORMAT(//3X,'TRUE PARAMETER:',4X,'USUAL:',8X,'BAYES:',9X,'SEB:',
1PLE,5X,'MEAN ONLY:',4X,'ICBNUGATE:',4X,'ICBNJ:',4X,'MEAN:',4X,'ICBNJ',
IRP = 100
IX = 423
DO 1000 JU = 1,IREP
DO 53 K = 1,N
DO 51 I = 1,NPL1
CALL RANDU(IX,IY,YFL)
51 R(I) = YFL
DO 52 I = 1,NPL1
IPL = I + 1
DO 52 J = 1,NPL1
IF(J.LT.IPL) GO TO 52
IF(R(I).LE.R(J)) GO TO 52
TEMP = R(I)
R(I) = R(J)
R(J) = TEMP
52 CONTINUE
53 B(K) = R(IRP)
DO 54 I=1,N

```

```

DE 54 J=1,M
CALL RANDU(IX,IY,YFL)
IX = IY
IF(YFL*LE*8(I)) T1(I,J) = 1.0
IF(YFL*GT*8(I)) T1(I,J) = 0.0
54 CONTINUE
DE 55 I=1,N
T(I) = 0.0
DE 55 J = 1,M
T(I) = T(I) + T1(I,J)
55 CONTINUE
T(I) = T(I) + T1(I,J)
DE 42 I = 1,N
EST1(I) = T(I)/ZM
EST2(I) = (RP + T(I))/(ZNP + ZM)
EST5(I) = (SORT(ZM)*V1 + T(I))/(ZM+SORT(ZM))
42 CONTINUE
C
COMPUTATION OF RBBINS SIMPLE EB ESTIMATOR
DE 400 I = 1,N
400 T1(I) = T(I) - T1(I,M)
SUM1 = 0.0
SUM2 = 0.0
TEMP = T1(I) + 1.0
DE 402 J=1,N
IF(T(J)*NE*TEMP) GO TO 401
SUM1 = SUM1 + 1.0
IF(T1(I)*NE*TL1(N)) GO TO 402
SUM2 = SUM2 + 1.0
402 CONTINUE
EST4(J) = ((T1(N)+1.0)/ZM)*(SUM1/SUM2)
IF(EST4(J)*GT*1.0) EST4(J) = 1.0
XLSSS4(J) = (EST4(J)-8(N))**2
TEMP1 = 0.0
TEMP2 = 0.0
DE 420 I = 1,N
TEMP1 = TEMP1 + T(I)
TEMP2 = TEMP2 + T(I)*(T(I)-1.0)
420 TEMP2 = TEMP2 + T(I)*(T(I)-1.0)
V11 = TEMP1/(ZM*ZN)
IF(ZM*EQ*1.0) GO TO 430
V22 = TEMP2/(ZN*ZM*(ZM-1.0))
V112 = V11**2
IF(V22*LE*V112) GO TO 430
ZNUM = V11*(V11-V22)+T(N)*(V22-V11**2)
DENOM = ZM*(V22-V11**2)+(V11-V22)
IF(DENOM*LT*1.0E-10) GO TO 430
EST6(J) = ZNUM/DENOM
XLSSS6(J) = (EST6(J)-8(N))**2
GO TO 431

```

```

430 EST6(JJ) = V11
XLSS6(JJ) = (V11 - B(N))**2
431 CONTINUE
V12 = V1**2
IF(ZM*EG.1.0) GO TO 432
IF(V22*LE.V12) GO TO 432
IF(V22*GT.V1) GO TO 433
ZNUM = V1*(V1-V22)+T(N)*(V22-V1**2)
DENOM = ZM*(V22-V1**2)+(V1-V22)
EST7(JJ) = ZNUM/DENOM
GO TO 434
432 EST7(JJ) = V1
GO TO 434
433 EST7(JJ) = T(N)/ZM
XLSS7(JJ) = (EST7(JJ) - B(N))**2
V1 = V11
UP = .5 + .5*SQRT(1.0-4.0*VAR)
DOWN = .5 - .5*SQRT(1.0-4.0*VAR)
IF(V11*GT.UP)
  VV = .5+.5*SQRT(1.0-4.0*VAR)
IF(V11*LT.DOWN)
  VV = .5-.5*SQRT(1.0-4.0*VAR)
ZNUM = VAR*(T(N)-VV)+(VV**2)*(1.0-VV)
DENOM = VAR*(ZM-1.0)+VV*(1.0-VV)
EST8(JJ) = ZNUM/DENOM
XLSS8(JJ) = (EST8(JJ) - B(N))**2
1000 CONTINUE

423 FORMAT(//1X,1AVG.RISK FOR ONLY MEAN KNOWN EST. = ,F15.6)
SUM4 = 0.0
SUM6 = 0.0
SUM7 = 0.0
SUM8 = 0.0
DO 408 I = 1,IREP
  SUM4 = SUM4 + XLSS4(I)
  SUM6 = SUM6 + XLSS6(I)
  SUM7 = SUM7 + XLSS7(I)
  SUM8 = SUM8 + XLSS8(I)
408 SUM8 = SUM8/ZZ
RISK4 = SUM4/ZZ
RISK6 = SUM6/ZZ
RISK7 = SUM7/ZZ
RISK8 = SUM8/ZZ
151

```

```

TEMP4 = 0.0
TEMP6 = 0.0
TEMP7 = 0.0
TEMP8 = 0.0
DO 416 I=1,IREP
  TEMP4 = (XLSS4(I) - RISK4)**2
  TEMP6 = (XLSS6(I) - RISK6)**2
  TEMP7 = (XLSS7(I) - RISK7)**2
  TEMP8 = TEMP8 + (XLSS8(I) - RISK8)**2
VAR4 = TEMP4/(ZZ-1.0)
VAR6 = TEMP6/(ZZ-1.0)
VAR7 = TEMP7/(ZZ-1.0)
VAR8 = TEMP8/(ZZ-1.0)
SE4 = SQRT(VAR4/ZZ)
SE6 = SQRT(VAR6/ZZ)
SE7 = SQRT(VAR7/ZZ)
SE8 = SQRT(VAR8/ZZ)
WRITE (6,409) IREP
409 FORMAT(//2X,'MONTE CARLO ESTIMATES OF BAYES RISKS',10X,'NO. OF
  REPLICATIONS = ',14,'//21X,'AVG. RISK',9X,'VAR',13X,'SE')
WRITE(6,417) RISK4,VAR4,SE4
417 FORMAT(//1X,'SIMPLE EB',6X,'F15.6)
WRITE(6,418) RISK6,VAR6,SE6
418 FORMAT(//1X,'CONJUGATE PRIOR',3F15.6)
WRITE(6,419) RISK7,VAR7,SE7
419 FORMAT(//1X,'CONJUGATE, MEAN',1X,'F15.6)
WRITE(6,421) RISK8,VAR8,SE8
421 FORMAT(//1X,'CONJUGATE, VAR',2X,'F15.6)
DO 200 J = 5,25,5
  ZJ = J
  DO 200 I = 1,20
    ZI = I
    ZNUM = RP*(ZNP-RP)*(ZNP+ZJ*(ZI**2)+ZI*ZJ+ZI*ZNP)
    DENOM = ZJ*((ZI+1.0)*ZNP)**2*(ZNP+1.0)
    RISK = ZNUM/DENOM
    WRITE(6,100) RISK
100 FORMAT(1X,'F16.7)
200 CONTINUE
STOP
END
C
C Program 4 Linkage program
TMIN = 1.0E75
TMAX = -1.0E75
DO 6 I = 1,N
  IF(EST1(I)-TMIN) 3,4,4
3 TMIN = EST1(I)
4 IF (EST1(I)-TMAX) 6,6,5
5 TMAX = EST1(I)

```

C
C Program 4 Linkage program

6 CONTINUE

DB 38 J=1,N1

XLAMDA(J) = TMIN + ((J-1)*ABS(TMAX-TMIN))/(N1-1)

38 CONTINUE

ORDERING OF THE N SAMPLES

DB 106 I = 1,N

106 T2(I) = T(I)

DB 103 I = 1,N

IPL = I + 1

DB 103 J = 1,N

IF(J.LT.IPL) GO TO 103

IF (T2(I).LE.T2(J)) GO TO 103

TEMP = T2(I)

T2(I) = T2(J)

T2(J) = TEMP

103 CONTINUE

DB 101 J = 1,N1
DB 101 I = 1,N2

Z1 = 1

IF(I.LE.N) A(I,J) = SFB(T2(I),XM,XLAMDA(J)) - (Z1-1.0)/ZN
IF(I.GT.N) A(I,J) = SFB(T2(I-N),XM,XLAMDA(J)) - (Z1-ZN)/ZN
A(I+N2,J) = -A(I,J)

101 CONTINUE

C THE FOLLOWING LOCATES LARGEST ELEMENT IN FIRST COL. OF THE A MATRI
C PLACES THE CORRESPONDING ROW IN THE LAST POSITION

DB 120 I = 1,N4

DB 120 J = 1,N4

IF(J.LE.I) GO TO 120

IF(A(I,1).LE.A(J,1)) GO TO 120

DB 121 K = 1,N1

ATEMP(K) = A(I,K)

A(I,K) = A(J,K)

121 A(J,K) = ATEMP(K)

120 CONTINUE

DB 140 I = 1,N4

IF(I.GT.1) GO TO 132

DB 133 J = 1,N5

IF(J.LE.N1) GO TO 134

AS(I,J) = 0.0

GO TO 133

134 AS(I,J) = 1.0

133 CONTINUE

GO TO 140

132 CONTINUE

IF (J.LE.N1) GO TO 136
 NIEMP = NI + 1
 IF (J.EQ.NIEMP) GO TO 137
 IF (J.EQ.N5) GO TO 138
 AS(I,J) = 0.0
 GO TO 140
 136 AS(I,J) = A(N4,J) - A(I-1,J)
 GO TO 140
 137 AS(I,J) = -1.0
 GO TO 140
 138 AS(I,J) = 1.0
 140 CONTINUE
 DO 130 I=1,N4
 IF (I=1) 127,127,126
 BS(I) = 0.0
 GO TO 130
 127 BS(I) = 1.0
 130 CONTINUE
 DO 150 I=1,N5
 IF (I=N1) 145,145,146
 C(I) = A(N4,I)
 GO TO 150
 146 IF (I.EQ.N5) GO TO 147
 C(I) = 0.0
 GO TO 150
 147 C(I) = 1.0
 150 CONTINUE
 DO 200 J=1,N5
 JH(J) = 0.0
 X(J) = 0.0
 Z(J) = 0.0
 200 CONTINUE
 CALL SIMPLE(0,N4,N5,AS,BS,C,K0,KB,PS,JH,X,Y,PE,E)
 DO 210 J = 1,N5
 IF (JH(J).EQ.0) GO TO 210
 I = JH(J)
 Z(I) = X(J)
 210 CONTINUE
 C THE FOLLOWING COMPUTES THE SMOOTH EB ESTIMATOR
 SUM1 = 0.0
 SUM2 = 0.0
 DO 300 J = 1,N1
 IF (XLAMDA(J).EQ.0.0.AND.1(N).EQ.0.0) GO TO 302

```

IF (XLAMDA(J).EQ.1.0.AND.(N).EQ.XM)      GO TO 303
ADD1      *      XLAMDA(J)*FB(T(N),XM,XLAMDA(J))*Z(J)
ADD2      *      FB(T(N),XM,XLAMDA(J))*Z(J)
GO TO 304
302 ADD1 = 0.0
GO TO 304
ADD2 = Z(J)
GO TO 304
303 ADD1 = Z(J)
GO TO 304
ADD2 = Z(J)
304 SUM1 = SUM1 + ADD1
300 SUM2 = SUM2 + ADD2
IF (SUM2.LT.1.0E-10) GO TO 309
EST3(J) = SUM1/SUM2
GO TO 306
309 EST3(J) = EST1(N)
306 MSEB(J) = (EST3(J)-B(N))*2

C
C Program 5 SIMPLE subroutine
C
SUBROUTINE SIMPLE(INFLAG,MX,NM,AS,SC,K0,K8,PS,JH,X,Y,PE,E)
AUTOMATIC SIMPLEX      REDUNDANT EQUATIONS CAUSE INFEASIBILITY
REAL BS(1),C(1),PS(1),X(1),Y(1),PE(1),E(1)
INTEGER INFLAG,MX,NM,K8(6),KB(1),JH(1)
EQUIVALENCE (XX,LL)
C THE FOLLOWING DIMENSION SHOULD BE THE SAME HERE AS IT IS IN CALLER.
REAL AS(80,90)
REAL AA,AJUT,BB,CBST,DT,RCBST,TEXP,TPIV,TY,XBLD,XX,XY,YI,YMAX
INTEGER I,IA,INVC,IR,ITER,J,JT,K,KAJ,LL,LL,M,M2,MM,N
INTEGER NCUT,NUMPV,NVER
LOGICAL FEAS,VER,NEG,TRIG,K0,ABSC
SET INITIAL VALUES, SET CONSTANT VALUES
ITER = 0
NUMVR = 0
NUMPV = 0
N = MX
M = NN
TEXP = .5**16
NCUT = 4*M + 10
NVER = M/2 + 5
M2 = M**2
FEAS = .FALSE.
IF (INFLAG.NE.0) GO TO 1400
DS 1402 L = 1,N
KB(J) = 0
K0 = .FALSE.
DS 1403 I = 1,M
IF (AS(I,J).EQ.0.0) GO TO 14C3
IF (K0.OR.AS(I,J).LT.0.0) GO TO 1402

```

```

1403 CONTINUE
      KQ = .TRUE.
1402 CONTINUE
      KB(J) = 1
1401 CONTINUE
      JH(I) = -1
1400 DB 1401 I = 1,M
      JH(I) = -1
1401 CONTINUE
      C* IVER1  CREATE INVERSE FROM (K9) AND (JH)
      (STEP 7)
      1320 VER = .TRUE.
      INVC = 0
      NUMVR = NUMVR + 1
      TRIG = .FALSE.
      DB 1101 I = 1,M2
      E(I) = 0.0
1101 CONTINUE
      MM=1
      DB 1113 I = 1,M
      E(MM) = 1.0
      PE(I) = 0.0
      X(I) = BS(I)
      IF (JH(I) .NE. 0) JH(I) = -1
      MM = MM + M + 1
1113 CONTINUE
      FORM INVERSE
      DB 1102 JT = 1,N
      IF (KB(JT) .EQ. 0) GB TB 1102
      CALL JMY
      C 600
      TY = 0.0
      KQ = .FALSE.
      DB 1104 I = 1,M
      IF (JH(I) .NE. (-1) .OR. ABS(Y(I)) .LE. TPIV) GB TB 1104
      IF (KQ) GB TB 1116
      IF (X(I) .EQ. 0) GB TB 1115
      IF (ABS(Y(I)/X(I)) .LE. TY) GB TB 1104
      TY = ABS(Y(I)/X(I))
      GB TB 1118
      KQ = .TRUE.
      GB TB 1117
      IF (X(I) .NE. 0 .OR. ABS(Y(I)) .LE. TY) GB TB 1104
      TY = ABS(Y(I))
      IR = 1
      CONTINUE
      KB(JT) = 0
      C
      TEST PIVOT
      GB TB 1102
      PIVOT
      GB TB 900

```

```

C 900 CALL PIV
1102 CONTINUE
C
DE 1109 I = 1,M
IF (JH(I).EQ.(-1)) JH(I) = 0
IF (JH(I).EQ.0) FEAS = .FALSE.
1109 CONTINUE
1200 VER = .FALSE.
C
***
PERFORM ONE ITERATION
C
C *XCK1 DETERMINE FEASIBILITY
(STEP 1)
NEG = .FALSE.
IF (FEAS) GO TO 500
FEAS = .TRUE.
DE 1201 I = 1,M
IF (X(I).LT.0.0) GO TO 1250
IF (JH(I).EQ.0) FEAS = .FALSE.
1201 CONTINUE
C *IGET1 GET APPLICABLE PRICES
GO TO 501
IF (.NOT.FEAS) GO TO 501
PS(I) = PE(I)
IF (X(I).LT.0.0) X(I) = 0.
503 CONTINUE
ABSC = .FALSE.
GO TO 599
1250 FEAS = .FALSE.
NEG = .TRUE.
DE 504 J = 1, M
PS(J) = 0.
504 CONTINUE
ABSC = .TRUE.
DE 505 I = 1,M
MM = 1
IF (X(I).GE.0.0) GO TO 507
ABSC = .FALSE.
DE 508 J = 1,M
PS(J) = PS(J) + E(MM)
MM = MM + M
508 CONTINUE
GO TO 505
IF (JH(I).NE.0) GO TO 505
IF (X(I).NE.0.0) ABSC = .FALSE.
DE 510 J = 1,M
PS(J) = PS(J) - E(MM)
MM = MM + M
510 CONTINUE
505 CONTINUE
C *IMINI FIND MINIMUM REDUCED COST
599 JT = 0
(STEP 3)

```

```

      BB = 0.0
      DO 701 J = 1, N
      IF (KB(J).NE.0) GO TO 701
      DT = 0.0
      DE 303 I = 1, M
      DT = DT + PS(I) * AS(I, J)
      CONTINUE
      IF (FEAS) DT = DT + C(J)
      IF (ABSC) DT = -ABS(DT)
      IF (DT.GE.BB) GO TO 701
      BB = DT
      JT = J
      701 CONTINUE
      C TEST FOR NB PIVOT COLUMN
      IF (JT.LE.0) GO TO 203
      C TEST FOR ITERATION LIMIT EXCEEDED
      IF (ITER.GE.NCUT) GO TO 160
      ITER = ITER + 1
      C* JUMP! MULTIPLY INVERSE TIMES A(:,JT)
      600 DO 610 I = 1, M
      Y(I) = 0.0
      610 CONTINUE
      LL = 0
      CBST = C(JT)
      DE 605 I = 1, M
      AJT = AS(I, JT)
      IF (AJT.EQ.0.0) GO TO 602
      CBST = CBST + AJT * PE(I)
      DE 606 J = 1, M
      LL = LL + 1
      Y(J) = Y(J) + AJT * E(LL)
      606 CONTINUE
      GO TO 605
      LL = LL + M
      605 CONTINUE
      C COMPUTE PIVOT TOLERANCE
      YMAX = 0.0
      DE 620 I = 1, M
      YMAX = AMAX1(ABS(Y(I)), YMAX)
      620 CONTINUE
      TPIV = YMAX * TEXP
      C RETURN TO INVERSION ROUTINE, IF INVERTING
      IF (VER) GO TO 1114
      CBST TOLERANCE CONTROL
      RCOST = YMAX/BB
      IF (TRIG.AND.BB.GE.(-TPIV)) GO TO 203
      TRIG = .FALSE.
      IF (BB.GE.(-TPIV)) TRIG = .TRUE.
      C* IRBM! SELECT PIVOT ROW
      (STEP 5)

```

```

IR = 0
AA = 0.0
KA = .FALSE.
DE 1050 I = 1,M
IF (X(I).NE.0.0.EP.Y(I).LE.TP.IV) 58 10 1053

```

1045 IF (Y(I).LE.AA) GO TO 1050
IF (KJ) GO TO 1050
IF (JH(I).EQ.O) GO TO 1044

1044 IF (KQ) GG TB 1045

KG - TRUE.

1047 AA = Y(I)

IR - I

IF (IR·NE·0) 66 TB 1099

AA = 1.0E+20

C FIND MIN. PIVOT AMONG POSITIVE EQUATIONS

DB 1010 I 1010

(I) A (I) X - A A

MR. J. H. HARRIS

1010 CONTINUE

IF (.NOT.NEG) GO TO 1099

C FIND PIVOT AMONG NEGATIVE EQUATIONS, IN WHICH X/Y IS LESS THAN THE C MINIMUM X/Y IN THE POSITIVE EQUATIONS, THAT HAS THE LARGEST ABSF(Y)

BB - TP IV

DB 1030 I - 1, M

(1) 人 88

IF (X(I)•GE•0••BR•Y(I)•GE•BB•BR•Y(I)•AA•GT•X(I)) GG TB 1030

1030 CONTINUE

C TEST FOR NO PIVOT ROW

1099 IF (IR.LE.0) GO TO 207

C* PIV: PIVOT ON (IR, JT)

IA - JH (IR)

$$IF(IA+GT,0) \quad KB(IA) = 0$$

900 NUMPY + 1

JH(IR) • JT

KB(JT) - IR

YI = Y (IR)

Y(IIR) = 100

0 77

REF ID: A66804

$$81 + 77 = 158$$

IF (E(L)•NE•0•0)

W + 77 = 77

159

STOP

160

```

905  G8 T8 904
      XY = E(L) / YI
      PE(J) = PE(J) + CBST * XY
      E(L) = 0.0
      D8 906 I = 1, M
      XOLD = X(I)
      X(I) = XOLD + XY * Y(I)
      IF (.NOT. VER.AND.X(I).LT.0.0.AND.XOLD.GE.0.0) X(I) = 0.0
906  CONTINUE
904  CONTINUE
      TRANSFORM X
      XY = X(IR) / YI
      D8 908 I = 1, M
      XOLD = X(I)
      X(I) = XOLD + XY * Y(I)
      IF (.NOT. VER.AND.X(I).LT.0.0.AND.XOLD.GE.0.0) X(I) = 0.0
908  CONTINUE
      Y(IR) = -YI
      X(IR) = -XY
      IF (VER) G8 T8 1102
      IF (NUMPV.LE.M) G8 T8 1200
      C TEST FOR INVERSION ON THIS ITERATION
      INVC = INVC + 1
      IF (INVC.EQ.NVER) G8 T8 1320
      G8 T8 1200
      C* END OF ALGORITHM, SET EXIT VALUES
      207 IF (.NOT.FEAS.OR.RCBST.LE.(-1000.)) G8 T8 203
      C INFINITE SOLUTION
      K = 2
      G8 T8 250
      C PROBLEM IS CYCLING
      160 K = 4
      G8 T8 250
      C FEASIBLE OR INFEASIBLE SOLUTION
      203 K = 0
      250 IF (.NOT.FEAS) K = K + 1
      D8 1399 J = 1, N
      XX = 0.0
      KBJ = KBJ
      IF (KBJ.NE.0) XX = X(KBJ)
      KBJ(J) = LL
      1399 CONTINUE
      KBJ(1) = K
      KBJ(2) = ITER
      KBJ(3) = INVC
      KBJ(4) = NUMVR
      KBJ(5) = NUMPV
      KBJ(6) = JT
      RETURN
      END

```

- [1] Blischke, W.R. "Estimating the Parameters of Mixtures of Binomial Distributions." JASA, 59 (1964), 510-528.
- [2] Blischke, W.R. "Moment Estimators for the Parameters of a Mixture of Two Binomial Distributions." AMS, 33(1962), 444-454.
- [3] Blum, J.R. and Judah Rosenblatt. "On Partial a Prior Information in Statistical Inference." AMS, 38 (1967), 1671-1678.
- [4] Boes, D.C. "Minimum Unbiased Estimator of Mixing Distribution for Finite Mixtures." Sankhya, A 29 (1967), 417-420.
- [5] Boes, D.C. "On the Estimation of Mixing Distributions." AMS, 37 (1966), 177-188.
- [6] Clemmer, B.A. and R.G. Krutchkoff. "The Use of Empirical Bayes Estimators in a Linear Regression Model." Biometrika, 55 (1968), 525-535.
- [7] Cote, L. and M. Skibinsky. "On the Inadmissibility of Some Standard Estimates in the Presence of Prior Information." AMS, 34 (1963), 539-548.
- [8] Cozui, C.H. and G.R. Cooper. "The Use of Prior Estimates in Successive Point Estimations of a Random Signal." J. SIAM, 14 (1966), 112-130.
- [9] Cramer, H. Mathematical Methods of Statistics. Princeton: Princeton University Press, 1946.
- [10] Deely, John J. "Non-Parametric Empirical Bayes Procedures for Selecting the Best of K Populations." AMS, 37 (1966), 540-541.
- [11] Deely, J.J. and R.L. Kruse. "Construction of Sequences Estimating the Mixing Distribution." AMS, 39 (1968), 286-288.

LIST OF REFERENCES

- [12] Deely, J.J. and W.J. Zimmer. "Shorter Confidence Intervals Using Prior Observations." JASA, 64 (1969), 378-386.
- [13] Ferguson, T.S. Mathematical Statistics: A Decision Theoretic Approach. New York: Academic Press, 1967.
- [14] Hald, A. "The Mixed Binomial Distribution and the Posterior Distribution of p for a Continuous Prior Distribution." JRSS (B), 30 (1968), 359-367.
- [15] Hoagley, B. "The Compound Multinomial Distribution and Bayesian Analysis of Categorical Data From Finite Populations." JASA, 64 (1969), 216-229.
- [16] Hodges, J.L. and E.L. Lehmann. "The Use of Previous Experience in Reaching Statistical Decisions." AMS, 23 (1952), 396-407.
- [17] Johns, M.V. "An Empirical Bayes Approach to Non-Parametric Two-Way Classification" in Studies in Stem Analysis and Prediction (ed. by H. Solomon). Stanford U. Press, 1961, 221-232.
- [18] Johns, M.V. "Non-Parametric Empirical Bayes Procedures." AMS, 28 (1957), 649-669.
- [19] Kagan, A.M. "An Empirical Bayes Approach to the Estimation Problem." Doklady of the Academy of Sciences of the USSR, 147: 5 (1962), 1020-1021.
- [20] Krutchkoff, R.G. "A Supplementary Sample Non-Parametric Empirical Bayes Approach to Some Statistical Decision Problems." Biometrika, 54 (1967), 451-458.
- [21] Lindley, D.V. Introduction to Probability and Statistics from a Bayesian Viewpoint: Volumes I and II. Cambridge: University Press, 1965.
- [22] Maritz, J.S. "On the Smooth Empirical Bayes Approach to Testing of Hypotheses and the Compound Decision Problem." Biometrika, 55 (1968), 83-100.

- [23] Maritz, J.S. "Smooth Empirical Bayes Estimation for Continuous Distributions." Biometrika, 54 (1967), 435-450.
- [24] Maritz, J.S. "Smooth Empirical Bayes Estimation for One-Parameter Discrete Distributions." Biometrika, 53 (1966), 417-429.
- [25] Miyasawa, K. "An Empirical Bayes Estimator of the Mean of a Normal Population." Bull. Inst. Int. Stat., 38 (1961), 181-188.
- [26] Neyman, J. "Two Breakthroughs in the Theory of Statistical Decision Making." Rev. Int. Statist. Inst., 30 (1962), 11-27.
- [27] Parzen, E. Modern Probability Theory and Its Applications. New York: John Wiley and Sons, 1960.
- [28] Raiffa, Howard and R. Schlaifer. Applied Statistical Decision Theory. Boston: Harvard University Press, 1961.
- [29] Rider, Paul R. "Estimating the Parameters of Mixed Poisson, Binomial and Weibull Distributions by the Method of Moments." BISI, 39 (1962), 225-232.
- [30] Rider, Paul R. "The Method of Moments Applied to a Mixture of 2 Exponential Distributions." AMS, 32 (1961), 143-147.
- [31] Robbins, H. "Asymptotically Subminimax Solutions of Compound Statistical Decision Problems" in Proc. of 2nd Berkeley Symp. on Math. Stat. and Prob., Berkeley, 1951, 131-148.
- [32] Robbins, H. "The Empirical Bayes Approach to Statistical Decision Problems." AMS, 35 (1964) 1-20.
- [33] Robbins, H. "An Empirical Bayes Approach to Statistics" in Proc. 3rd Berkeley Symp. on Math. Stat. and Prob., 1956, 157-164.

- [34] Robbins, H. "The Empirical Bayes Approach to Testing Statistical Hypotheses." Rev. Int. Stat. Inst., 31 (1963), 195-208.
- [35] Rolph, J.E. "Bayesian Estimation of Mixing Distributions." AMS, 39 (1968), 1289-1302.
- [36] Rutherford, J.R. "Monte Carlo Analysis of the Rate of Convergence of Some Empirical Bayes Point Estimators." AMS, 37 (1966), 766-767.
- [37] Rutherford, J.R. and R.G. Krutchkoff. "Asymptotic Optimality of Empirical Bayes Estimators." Biometrika, 56 (1969), 220-223.
- [38] Rutherford, J.R. and R.G. Krutchkoff. "Some Empirical Bayes Techniques in Point Estimation." Biometrika, 56 (1969), 133-137.
- [39] Samuel, E. "An Empirical Bayes Approach to the Testing of Certain Parametric Hypotheses." AMS, 34 (1963), 1370-1385.
- [40] Skibinsky, M. "Minimax Estimation of a Random Probability." SIAM Journal Appl. Math., 16 (1968), 134-145.
- [41] Suzuki, G. "On Testing Certain Hypotheses." Ann. I. Stat., 20 (1968), 71-80.
- [42] Teicher, H. "Identifiability of Finite Mixtures." AMS, 34 (1963), 1265-1269.
- [43] Teicher, H. "Identifiability of Mixtures." AMS, 32 (1961), 244-248.
- [44] Teicher, H. "On the Mixtures of Distribution." AMS, 31 (1960), 55-73.
- [45] Tucker, H.G. "An Estimate of the Compounding Distribution of a Compound Poisson Process." Theor. Prob. Appl., 8 (1963), 195-200.

- [46] Tucker, H.G. A Graduate Course in Probability. New York: Academic Press, 1967.
- [47] Ula, N. and M. Kim. "An Empirical Bayes Approach to Adaptive Control." Journal of the Franklin Institute, 280 (1965), 189-204.
- [48] Weiler, H. "The Use of Incomplete Beta Functions for Prior Distributions in Binomial Sampling." Technometrics, 7 (1965), 335-347.
- [49] Winkler, R.L. "The Assessment of Prior Distributions in Bayesian Analysis." JASA, 62 (1967), 776-800.

DOCUMENT CONTROL DATA - R & D

(Security classification of title, body of abstract and indexing annotation must be entered when the overall report is classified)

1. ORIGINATING ACTIVITY (Corporate author)

SOUTHERN METHODIST UNIVERSITY

2b. GROUP

UNCLASSIFIED

UNCLASSIFIED

2a. REPORT SECURITY CLASSIFICATION

3. REPORT TITLE

Partial Prior Information: Some Empirical Bayes and G-Minimax Decision Functions

4. DESCRIPTIVE NOTES (Type of report and inclusive dates)

Technical Report

5. AUTHOR(S) (First name, middle initial, last name)

Stephen L. George

6. REPORT DATE

October 30, 1969

8a. CONTRACT OR GRANT NO.

N00014-68-A-0515

b. PROJECT NO.

NR 042-260

d.

10. DISTRIBUTION STATEMENT

This document has been approved for public release and sale; its distribution is unlimited.

11. SUPPLEMENTARY NOTES

12. SPONSORING MILITARY ACTIVITY

Office of Naval Research

13. ABSTRACT

If the prior probability distribution function G of the parameters in a statistical decision theory model is not completely specified then a Bayes decision function cannot be obtained. However, in many cases there may be some partial (i.e., incomplete) prior information concerning G .

The types of partial prior information considered in this paper is of two kinds: (i) N past observations on the compound distribution $f_G(t)$ or (ii) knowledge of a restrictive class G to which the unknown G is assumed to belong.

When past observations are available, empirical Bayes decision functions are found which are: (i) asymptotically optimal (ii) easy to apply and (iii) better than some optimal non-Bayes decision function for reasonably small N . When only knowledge of the class G is assumed, G-minimax decision functions are found.

In all cases, estimators for the parameters of the normal and Bernoulli processes are found using a quadratic loss function. These estimators are then compared to each other and to other well-known estimators by means of their Bayes risks and, in the empirical Bayes case, by means of their global risks for small N .