

AD-A100 947

PURDUE UNIV LAFAYETTE IN DEPT OF STATISTICS
ON THE PROBLEM OF SELECTING GOOD POPULATIONS.(U)
AUG 80 W KIN
MMS-80-20

F/8 12/1

N00014-75-C-0455

NL

UNCLASSIFIED

1 of 1
AD-A
INFO



END
DATE
FILMED
7-81
DTIC

AD A100942

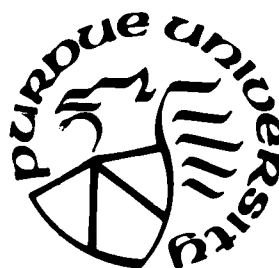
LEVEL II

①

PURDUE UNIVERSITY

DTIC
ELECTE
JUL 06 1981
S D
E

To DTIC



DEPARTMENT OF STATISTICS

60

DIVISION OF MATHEMATICAL SCIENCES

DTIC FILE COPY

81 6 29 242

X

ON THE PROBLEM OF SELECTING GOOD POPULATIONS*,

by

Shanti S. Gupta
Purdue University, W. Lafayette, Indiana

and

Woo-Chul Kim
Seoul National University, Seoul, Korea

Department of Statistics
Division of Mathematical Sciences
(9) Mimeograph Series #80-20

August 1980

APPROVED FOR PUBLIC RELEASE
DISTRIBUTION UNLIMITED

*This research was supported by the Office of Naval Research contract
N00014-75-C-0455 at Purdue University. Reproduction in whole or in
part is permitted for any purpose of the United States Government.

ON THE PROBLEM OF SELECTING GOOD POPULATIONS

Shanti S. Gupta

Purdue University
West Lafayette, Indiana

Woo-Chul Kim

Seoul National University
Seoul, Korea

ABSTRACT

The problem of selecting good populations out of k normal populations is considered in a Bayesian framework under exchangeable normal priors and additive loss functions. Some basic approximations to the Bayes rules are discussed. These approximations suggest that some well-known classical rules are "approximate" Bayes rules. Especially, it is shown that Gupta-type rules are extended Bayes with respect to a family of the exchangeable normal priors for any bounded and additive loss function. Furthermore, for a simple loss function, the results of a Monte Carlo comparison of Gupta-type rules and Seal-type rules are presented. They indicate that, in general, Gupta-type rules perform better than Seal-type rules.

Accession For	
NTIS GRA&I	☒
DTIC TAB	☐
Unannounced	☐
Justification	
By	
Distribution/	
Availability Codes	
Dist	Avail and/or Special
A	

1. INTRODUCTION

In many practical situations, the experimenter often faces the problem of comparing several competing populations, treatments or processes. The statistical methodology of ranking and selection procedures provides useful techniques for solving such problems. One of the basic formulations of the selection problem is the subset selection formulation of Gupta (1956), under which we select a random sized non-empty (small) subset of the populations while controlling the probability of including the 'best' population in the selected subset. One modification of the basic goal is concerned with the selection of "good" populations which are defined below. Contributions in this direction have been made by Fabian (1962), Desu (1970), Carroll, Gupta and Huang (1975), and Panchapakesan and Santner (1977), among others. Moreover, Berger (1979) and Bjørnstad (1980) have studied the minimaxity of some well-known classical rules under various control conditions.

Let $\pi_1, \dots, \pi_k$ be k independent normal populations with unknown means $\theta_1, \dots, \theta_k$, respectively, and a common known variance. Let the ordered θ_i be denoted by $\theta_{[1]} \leq \dots \leq \theta_{[k]}$. A population π_i is said to be

a good population if $\theta_i \geq \theta_{[k]} - \Delta$,

a bad population if $\theta_i < \theta_{[k]} - \Delta$,

where Δ is a given positive constant. This definition implies that the experimenter is willing to accept all the populations which are sufficiently close to the 'best' population (i.e. the populations associated with $\theta_{[k]}$) while screening out the bad populations. Thus it seems reasonable that any suitable loss function should contain two components: one depending on the bad populations in the selected subset and another depending only on the good populations excluded. Then the loss function of the following type seems to be appropriate for our purpose: for $\underline{\theta} = (\theta_1, \dots, \theta_k)$ and $a \subset \{1, 2, \dots, k\}$ ($a \neq \phi$), let

$$(1.1) \quad L(\underline{\theta}, a) = \sum_{i \in a} L_B(\theta_i - \theta_{[k]}^+ \Delta) + \sum_{i \notin a} L_G(\theta_i - \theta_{[k]}^+ \Delta),$$

where L_B is non-increasing, L_G is non-decreasing, $L_B(y) = 0$ for $y \geq 0$ and $L_G(y) = 0$ for $y < 0$. Here the action $a \in G$ means that we select the set $\{\pi_i, i \in a\}$ as the set of good populations where G is the action space consisting of all non-empty subsets of $\{1, 2, \dots, k\}$. Any loss function of the type given in (1.1) will be called additive.

Miescke (1979) has shown that the selection rule proposed by Gupta (1956) is asymptotically Bayes as the sample size increases for an additive loss function when the unknown means are assumed to have an exchangeable normal prior. Also, he has studied the approximations to the Bayes rules for an additive and linear loss function i.e. the one corresponding to $L_B(y) = y^-$ and $L_G(y) = y^+$, with y^+ (y^-) denoting the usual positive (negative) part of y .

In this paper, we assume that the unknown means $\theta_1, \dots, \theta_k$ have an exchangeable normal prior. Under this assumption some basic approximations to the Bayes rules for any additive loss function are discussed. Gupta-type rules are shown to be extended Bayes with respect to the family of exchangeable normal priors for a class of additive loss functions. Also, a Monte Carlo comparison of two well-known classical rules with the Bayes rule is carried out for a simple loss function. This empirical study as well as our theoretical results support the earlier results of the Monte Carlo studies by Chernoff and Yahav (1977), and Gupta and Hsu (1978), which indicate that Gupta-type rules would perform as well as the Bayes rule for a wide class of loss functions.

2. BAYES RULE AND APPROXIMATIONS TO BAYES RULE

Suppose that we have k independent samples of size n from each population. By sufficiency we can reduce the problem to that based on the sample means $X_1, \dots, X_k$, whose common variance may be assumed to be 1 without loss of generality.

It is assumed that the loss function is given by (1.1), and

we further assume that the unknown means $\underline{\theta} = (\theta_1, \dots, \theta_k)$ have an exchangeable normal distribution such that, for $-(k-1)^{-1} < \rho < 1$,

$$(2.1) \quad E(\theta_i) = m, \text{Var}(\theta_i) = \sigma_0^2 \quad (i = 1, \dots, k), \text{ and}$$

$$\text{Cov}(\theta_i, \theta_j) = \rho \sigma_0^2 \quad (1 \leq i < j \leq k).$$

Gupta and Hsu (1978) and Miescke (1979) have used a representation of θ_i 's similar to one given in (2.2) to reduce the prior to an iid normal prior;

$$(2.2) \quad (\theta_i - m)/\sigma_0 = \sqrt{1-\rho} Z_i - (\sqrt{\rho^-} + \sqrt{\rho^+}) Z_0,$$

where $Z_0, Z_1, \dots, Z_k$ are standard normal random variables with $Z_1, \dots, Z_k$ being independent and $\text{Cov}(Z_0, Z_i) = \sqrt{\rho^-}/\sqrt{1-\rho}$ ($i=1, \dots, k$). Note that we can restrict ourselves to the translation invariant rules δ i.e. $\delta(\underline{x}) = \delta(\underline{x} + \underline{w})$ for any $\underline{w} = (w, \dots, w) \in R^k$ in this framework. Such a reduction is well presented in the next result.

Lemma 1. Let δ be a translation invariant rule. Then the overall risk wrt the prior in (2.1) can be written as follows; for $\sigma^2 = (1-\rho)\sigma_0^2$,

$$(2.3) \quad r(\sigma^2, \delta) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} L(\underline{\theta}, \delta(\underline{x})) dN(\underline{\theta} | b\underline{x}, bI) dN(\underline{x} | \underline{0}, (1-b)^{-1}I),$$

where $b = \sigma^2/(1+\sigma^2)$, I is the $k \times k$ identity matrix and $N(\cdot | \underline{\mu}, \Sigma)$ denotes the distribution function of a multivariate normal distribution with mean $\underline{\mu}$ and covariance matrix Σ .

Let $x_{[1]} \leq \dots \leq x_{[k]}$ denote the ordered observations of $X_1, \dots, X_k$ and $\pi_{(i)}$ and $\theta_{(i)}$ denote the π and θ associated with $x_{[i]}$ for $i = 1, \dots, k$. Then it follows from Lemma 1 that $\theta_{(i)}$, a posteriori, has normal distribution with mean $b x_{[i]}$ and variance $b = \sigma^2/(1+\sigma^2)$. Now the Bayes rule given in the next result can be easily derived by using the methods similar to those in Goel and Rubin (1977).

Theorem 1. Assume that the loss function is given by (1.1). Then the Bayes rule δ^* wrt the prior in (2.1) always selects $\pi_{(k)}$, and

moreover it selects $\pi_{(i)}$ if and only if $E[\ell(\theta_{(i)} - \theta_{[k]} + \Delta) | \underline{x}] < 0$, for $i = 1, \dots, k-1$ with $\ell(\cdot) = L_B(\cdot) - L_G(\cdot)$.

Remark 1. Properties of the Bayes rule such as the monotonicity and orderedness can be derived from Theorem 1 (see, for example, Miescke (1979) and Kim (1979)).

Eventhough Theorem 1 gives a general description of the Bayes rule for any additive loss function, a complete specification of the Bayes rule requires the explicit form of $\ell(\cdot)$. Also, it usually involves difficult computations to implement the Bayes rule. So it is of practical significance to examine some basic approximations to the Bayes rule, which are applicable to every additive loss function provided they are easily computable. For this purpose, we introduce the following notations: for $i = 1, \dots, k-1$,

$$(2.4) \quad \underline{x}_i^* = (x_{[1]} - x_{[i]}, \dots, x_{[i-1]} - x_{[i]}, x_{[i+1]} - x_{[i]}, \dots, x_{[k]} - x_{[i]}),$$

$$\underline{x}_i^a = \left(\frac{1}{k-1} \sum_{j \neq i} x_{[j]} - x_{[i]}, \dots, \frac{1}{k-1} \sum_{j \neq i} x_{[j]} - x_{[i]} \right),$$

$$\underline{x}_i^u = (0, 0, \dots, 0, \sum_{j=i+1}^k (x_{[j]} - x_{[i]})),$$

$$\underline{x}_i^m = (x_{[k]} - x_{[i]}, \dots, x_{[k]} - x_{[i]}),$$

$$D(\underline{x}_i^*) = E[\ell(\theta_{(i)} - \theta_{[k]} + \Delta) | \underline{x}]$$

$$= E[\Delta - \sqrt{b} \{ \max_{j \neq i} (Z_j - Z_i + \sqrt{b} x_{[j]} - \sqrt{b} x_{[i]}) \}^+],$$

where $Z_1, \dots, Z_k$ are iid normal random variables. It follows from a result in Marshall and Olkin (1974) that $D(\underline{x}_i^*)$ is Schur-convex in $\underline{x}_i^*$. Also, note that $D(\underline{x}_i^*)$ is non-decreasing in each argument of $\underline{x}_i^*$. Furthermore, it can be easily shown that we have the following inequalities:

$$(2.5) \quad \underline{x}_i^a \preceq^m \underline{x}_i^* \preceq^m \left(\sum_{j < i} (x_{[j]} - x_{[i]}), 0, \dots, 0, \sum_{j > i} (x_{[j]} - x_{[i]}) \right),$$

where $\alpha \preceq^m \beta$ means that α is majorized by β . Hence, it follows

from the Schur convexity of $D(\underline{x}_i^*)$, (2.5) and the monotonicity of $D(\underline{x}_i^*)$ that, for $i = 1, 2, \dots, k-1$,

$$(2.6) \quad D(\underline{x}_i^a) \leq D(\underline{x}_i^*) \leq D(\underline{x}_i^u),$$

$$E[\Delta - (\theta_{(k)} - \theta_{(i)})^+ | \underline{x}] \leq D(\underline{x}_i^*) \leq D(\underline{x}_i^m).$$

These bounds on $D(\underline{x}_i^*)$ and Theorem 1 suggest the "approximate" Bayes rules as follows; for $b = \sigma^2/(1+\sigma^2)$,

(2.7) δ_1 : Select $\pi_{(k)}$ and, for $i = 1, \dots, k-1$, select $\pi_{(i)}$ iff

$$x_{[i]} \geq x_{[k]} - d_1(b)/\sqrt{b},$$

δ_2 : Select $\pi_{(k)}$ and, for $i = 1, \dots, k-1$, select $\pi_{(i)}$ iff

$$x_{[i]} \geq x_{[k]} - d_2(b)/\sqrt{b},$$

δ_3 : Select $\pi_{(k)}$ and, for $i = 1, \dots, k-1$, select $\pi_{(i)}$ iff

$$x_{[i]} \geq \left(\sum_{j=i+1}^k x_{[j]} - d_3(b)/\sqrt{b} \right) / (k-i),$$

δ_4 : Select $\pi_{(k)}$ and, for $i = 1, \dots, k-1$, select $\pi_{(i)}$ iff

$$x_{[i]} \geq \frac{1}{k-1} \sum_{j \neq i} x_{[j]} - d_1(b)/\sqrt{b},$$

where $d_j(b)$ ($j = 1, 2, 3$) are determined so that

$$(2.8) \quad H_1(b, d) = E[\Delta - \sqrt{b} (\max_{2 \leq j \leq k} Z_j - Z_1 + d)^+] \leq 0 \quad \text{iff } d \leq d_1(b),$$

$$H_2(b, d) = E[\Delta - \sqrt{b} (Z_2 - Z_1 + d)^+] \leq 0 \quad \text{iff } d \leq d_2(b),$$

$$E[\Delta - \sqrt{b} \{\max_{2 \leq j \leq k-1} (Z_j - Z_1, Z_k - Z_1 + d)\}^+] \leq 0 \quad \text{iff } d \leq d_3(b);$$

with $Z_1, \dots, Z_k$ being iid standard normal random variables. The next result follows from (2.6) and Theorem 1.

Corollary 1. For any additive loss function given in (1.1), the Bayes rule δ^* wrt an exchangeable normal prior satisfies the following relations with probability one; for $\alpha = 1, 3$ and $\beta = 2, 4$,

$$(2.9) \quad \delta_\alpha(\underline{x}) \subseteq \delta^*(\underline{x}) \subseteq \delta_\beta(\underline{x}).$$

It should be pointed out that more approximations to the Bayes rule can be obtained for a particular additive loss function (see, for instance, Miescke (1979) and Kim (1979)). Also note that, for $k = 2$, the "approximate" Bayes rules coincide so that the Bayes rule can be explicitly specified.

It seems interesting to note that the "approximate" Bayes rules except δ_3 are members of the class of the following well known classical selection rules;

$$(2.10) \quad \delta_d^a: \text{Select } \pi_i \text{ iff } x_i = x_{[k]} \text{ and/or } x_i \geq \frac{1}{k-1} \sum_{j \neq i} x_j - d,$$

$$\delta_d^m: \text{Select } \pi_i \text{ iff } x_i \geq x_{[k]} - d.$$

Rules δ^m were proposed by Gupta (1956) for the goal of selecting a subset containing the best population, and later studied by Desu (1970) for selecting a subset consisting of only good populations. Rules δ^a are modified versions of the selection rules proposed by Seal (1955).

Note that one Gupta-type rule δ_1 always selects a smaller subset than the Bayes rule while another rule δ_2 of the same type selects a larger subset. Thus, one might expect that the Gupta-type rules δ^m would be close to the Bayes rule in some sense. This is proved in Theorem 2 given below. First, we recall the next definition (see, for example, Ferguson (1967)).

Definition 1. A decision rule δ_0 is called an extended Bayes rule wrt a family $\mathfrak{F}$ of prior distributions if, for every $\epsilon > 0$, there exists a prior $\tau \in \mathfrak{F}$ such that $r(\tau, \delta_0) < \inf_{\delta} r(\tau, \delta) + \epsilon$.

Also, we need the following conditions.

Conditions A.

(A-1) The loss function is additive with $\mathfrak{L}(\cdot) = L_B(\cdot) - L_G(\cdot)$ being bounded.

(A-2) Let d_j^* denote the number $d_j(b)$ determined by (2.8) for $b = 1$, and assume that they are finite ($j=1,2$).

Theorem 2. Suppose that Conditions A hold. Then, for $d_1^* < d \leq d_2^*$, the selection rules δ_d^m are extended Bayes wrt the family of exchangeable normal priors; in fact,

$$(2.11) \quad \lim_{\sigma \rightarrow \infty} [r(\sigma, \delta_d^m) - r(\sigma, \delta^*)] = 0,$$

where $r(\sigma, \delta)$ denotes the overall risk of a rule given by (2.3).

Proof. It follows from Theorem 1, Corollary 1 and (2.6) that, for

$$b = \frac{\sigma^2}{1+\sigma^2},$$

$$(i) \quad \{\pi_{(i)} \notin \delta^*(\underline{x}), \pi_{(i)} \in \delta_d^m(\underline{x})\} \subset \{0 \leq D(\underline{x}_i^*) \leq H_1(b, \sqrt{b} d)\},$$

$$(ii) \quad \{\pi_{(i)} \in \delta^*(\underline{x}), \pi_{(i)} \notin \delta_d^m(\underline{x})\} \subset \{H_2(b, \sqrt{b} d) \leq D(\underline{x}_i^*) < 0\},$$

where H_j ($j = 1, 2$) are defined in (2.8). Let A_i and B_i denote the events $\{\pi_{(i)} \notin \delta^*(\underline{x}), \pi_{(i)} \in \delta_d^m(\underline{x})\}$ and $\{\pi_{(i)} \in \delta^*(\underline{x}), \pi_{(i)} \notin \delta_d^m(\underline{x})\}$, respectively. Then, (i) and (ii) imply that

$$(2.12) \quad \begin{aligned} r(\sigma, \delta_d^m) - r(\sigma, \delta^*) &= E \left[\sum_{i=1}^{k-1} D(\underline{x}_i^*) (I_{A_i} - I_{B_i}) \right] \\ &\leq E \left[\sum_{i=1}^{k-1} \{H_1(b, \sqrt{b} d) I_{A_i} - H_2(b, \sqrt{b} d) I_{B_i}\} \right]. \end{aligned}$$

Since the condition (A-1) implies the continuity of $H_j(b, d)$ ($j = 1, 2$), it can be easily shown that the monotonicity of $H_j(b, d)$ in b and d , the definition of $d_j(b)$ in (2.8) and the condition (A-2) imply the following facts;

(iii) $d_j(b)$ is non-increasing in b ($j = 1, 2$),

(iv) $\lim_{\substack{b \rightarrow 1 \\ b < 1}} d_j(b) = d_j^* \quad (j = 1, 2),$

(v) For $d \in (d_1^*, d_2^*]$, $d_1(b)/\sqrt{b} < d \leq d_2(b)/\sqrt{b}$ for b sufficiently close to 1 and $b < 1$.

It follows from (v), (2.9) and (2.12) that, for sufficiently large σ , i.e., for b sufficiently close to 1,

$$\begin{aligned} & r(\sigma, \delta_d^m) - r(\sigma, \delta^*) \\ & \leq \{H_1(b, \sqrt{bd}) - H_2(b, \sqrt{bd})\} \sum_{i=1}^{k-1} P[d_1(b)/\sqrt{b} < X_{[k]} - X_{[i]} \leq d_2(b)/\sqrt{b}] \\ & = \{H_1(b, \sqrt{bd}) - H_2(b, \sqrt{bd})\} \sum_{i=1}^{k-1} P[\sqrt{\frac{1-b}{b}} d_1(b) < Z_{[k]} - Z_{[i]} \leq \\ & \qquad \qquad \qquad \sqrt{\frac{1-b}{b}} d_2(b)], \end{aligned}$$

where $Z_{[1]} \leq \dots \leq Z_{[k]}$ are the ordered iid standard normal random variables. Hence, it follows from Conditions A that

$$\lim_{\sigma \rightarrow \infty} [r(\sigma, \delta_d^m) - r(\sigma, \delta^*)] = 0,$$

which completes the proof.

3. MONTE CARLO RESULTS FOR A SIMPLE LOSS FUNCTION

The well-known selection rules δ^m and δ^a in Section 2 are, in some sense, natural approximations to the Bayes rule, and it is shown that the Gupta-type rules δ^m would perform as well as the Bayes rule when the prior variance is large. In this respect, the comparative performance of the Gupta-type rules δ^m and the Seal-type rules δ^a was studied using Monte Carlo technique for a simple loss function given as follows:

$$(3.1) \quad L(\theta, a) = c_1 \sum_{i \in a} I_{(-\infty, 0)}(\theta_i - \theta_{[k]} + \Delta) + c_2 \sum_{i \notin a} I_{[0, \infty)}(\theta_i - \theta_{[k]} + \Delta),$$

where $c_1 > 0$, $c_2 > 0$ and $c_1 + c_2 = 1$.

For this simple loss function, it follows from (2.3) that for any translation invariant rule, the posterior risk is given by

$$(3.2) \quad c_1 \sum_{i \in \delta(x)} P[\theta(i) < \theta_{[k]} - \Delta | x] + c_2 \sum_{i \notin \delta(x)} P[\theta(i) \geq \theta_{[k]} - \Delta | x],$$

where $X_1, \dots, X_k$ are iid normal random variables with mean 0 and variance $1 + \sigma^2$. Also, by Theorem 1, the Bayes rule is determined by

$$\begin{aligned}
 (3.3) \quad D(\underline{x}_1^*) &= c_1 - P[\theta_{(i)} \geq \theta_{[k]} - \Delta | \underline{x}] \\
 &= c_1 - \int_{-\infty}^{\infty} \prod_{j \neq i} \phi(z + \sqrt{b}(x_{[i]} - x_{[j]} + \Delta/\sqrt{b})) d\phi(z),
 \end{aligned}$$

where $b = \sigma^2/(1+\sigma^2)$.

We carried out the Monte Carlo comparisons for $k = 3$ and $k = 9$. The relevant parameters in this comparison study are σ^2 , Δ and c_1 ($=1-c_2$), while $c = c_2/c_1$ was used instead of c_1 since c , being the ratio of two different sources of losses, seems more appealing than c_1 . The range of the parameter values in this report is as follows:

$$\begin{cases} \Delta = 0.5, 1.0, \\ \sigma = (1.5)^i \quad (i = -2(1)6), \\ c = 1, 2, 4, 8 \end{cases}$$

For each of parameter sets (Δ, σ, c) , 400 simulations for $k = 3$ and 100 simulations for $k = 9$ were carried out. In each simulation, the generation of k iid normal random variables $X_1, \dots, X_k$ with mean 0 and variance $1+\sigma^2$ was involved. The Bayes rule and its posterior risk are obtained from (3.2) and (3.3) by numerically computing $D(\underline{x}_1^*)$'s, where some of the computations can be omitted by using (2.9). Then, the estimated Bayes risk can be obtained by taking the average of the posterior risks. Two sufficiently fine grids of the constants d are used to obtain optimal values of d for δ^m and δ^a which minimize the average regrets, where the range of these trial values is determined by (2.9).

The estimated Bayes risk, the estimated regrets incurred by the optimal δ^m and the optimal δ^a are given in Table I along with the sample standard deviations of these estimates. The cells left blank in the table correspond to cases in which the Bayes rule selects only one population or all the populations. It can be observed from Table I that the performance of the Gupta-type rule δ^m is almost as good as that of the Bayes rules, and that it performs remarkably better when the prior variance becomes larger.

This agrees with Theorem 2. Also, we observe that, for $k = 3$, the Seal-type rule δ^a performs reasonably well though its performance is not as good as the Gupta-type rule δ^m . However, for $k = 9$, the rule δ^a performs very badly and it was observed that, for most values of σ , the optimal δ^a tends to select much larger subsets than the Bayes rule as c becomes larger. In Figures 1a-1f, the estimated risks of the Bayes, the optimal δ^m and the optimal δ^a are shown graphically for $\Delta = 1.0$ and some selected values of c . Table II gives the average number of the bad populations selected and the average number of the good populations excluded by each rule. Also, the proportions of times that the optimal δ^m and the optimal δ^a coincide with the Bayes rule are presented graphically in Figures 2a-2h.

Chernoff and Yahav (1977), and Gupta and Hsu (1978) have observed the performance of the rule δ^m similar to the one in this study under certain loss functions for the goal of selecting a subset containing the 'best' population. However, the performance of δ^a in the present study is worse than that observed by Gupta and Hsu (1978), and it seems that the Seal-type rule has little to recommend for the goal of selecting good populations. The results of the present study and the previous ones mentioned above indicate that the Gupta-type rule performs fairly well in various formulations at least in the Bayesian framework considered in this paper. However, it should be pointed out that the proper constant d should be chosen according to the operating characteristics of the particular problem considered before one chooses the constant d based on any intuitive control condition. For this reason the estimated optimal d -values of the rules δ^m are provided in Table III, which can be used if one accepts the framework in this section.

ACKNOWLEDGMENT

This research was supported by the Office of Naval Research Contract N00014-75-C-0455 at Purdue University.

BIBLIOGRAPHY

- Berger, R. L. (1979). Minimax subset selection for loss measured by subset size. *Ann. Statist.* 7, 1333-1338.
- Björnstad, J. F. (1980). A class of Schur-procedures and minimax theory for subset selection. Mimeo. Series No. 80-17, Department of Statistics, Purdue University, W. Lafayette, IN.
- Carroll, R. J., Gupta, S. S. and Huang, D. Y. (1975). On selection procedures for the t best populations and some related problems. *Commun. Statist.* 4, 987-1008.
- Chernoff, H. and Yahav, J. A. (1977). A subset selection problem employing a new criterion. *Statistical Decision Theory and Related Topics II*, (ed. S. S. Gupta and D. S. Moore), 93-119. Academic Press.
- Desu, M. M. (1970). A selection problem. *Ann. Math. Statist.* 41, 1596-1603.
- Fabian, V. (1962). On multiple decision methods for ranking population means. *Ann. Math. Statist.* 33, 248-254.
- Ferguson, T. S. (1967). *Mathematical Statistics: A Decision-Theoretic Approach*. Academic Press, New York.
- Goel, P. K. and Rubin, H. (1977). On selecting a subset containing the best population - a Bayesian approach. *Ann. Statist.* 5, 969-983.
- Gupta, S. S. (1956). On a decision rule for a problem in ranking means. Mimeo. Series No. 150, Inst. of Statistics, University of North Carolina, Chapel Hill, North Carolina.
- Gupta, S. S. and Hsu, J. C. (1978). On the performance of some selection procedures. *Commun. Statist.-Simul. Computa.*, B7(6), 561-591.
- Kim, W. C. (1979). On Bayes and gamma-minimax subset selection rules. Mimeo. Series No. 79-7, Department of Statistics, Purdue University, W. Lafayette, IN.
- Marshall, A. W. and Olkin, I. (1974). Majorization in multivariate distribution. *Ann. Statist.* 2, 1189-1200.
- Miescke, K. J. (1979). Bayesian subset selection for additive and linear loss functions. *Commun. Statist.-Theor. Meth.* A8(12), 1205-1226.

Panchapakesan, S. and Santner, T. J. (1977). Subset selection procedures for Δ_p -superior populations. *Commun. Statist.-Theor. Meth.* A6, 1081-1090.

Seal, K. C. (1955). On a class of decision procedures for ranking means of normal populations. *Ann. Math. Statist.* 26, 387-398.

TABLE I

The entry on top of each cell is the estimated Bayes risk and the numbers in the second and the third row are the regrets incurred by the optimal δ^m and by the optimal δ^a , respectively. The estimated standard deviations of the estimates are given in the parentheses.

$$k = 3, \Delta = 0.5$$

$\sigma \backslash c$	1	2	4	8
.44	.4609 (.0016) .0024 (.0005) .0334 (.0029)	.3297 (.0012) .0006 (.0002) .0033 (.0008)	.1995 (.0008) .0000 (.0000) .0001 (.0001)	
.67	.4686 (.0053) .0038 (.0006) .0487 (.0028)	.3897 (.0026) .0029 (.0005) .0185 (.0017)	.2558 (.0011) .0010 (.0003) .0024 (.0006)	.1456 (.0001) .0001 (.0001) .0000 (.0000)
1.00	.3872 (.0076) .0036 (.0006) .0175 (.0018)	.3424 (.0056) .0027 (.0005) .0160 (.0016)	.2511 (.0028) .0022 (.0004) .0053 (.0008)	.1532 (.0014) .0011 (.0003) .0016 (.0004)
1.50	.2813 (.0088) .0027 (.0006) .0115 (.0015)	.2710 (.0068) .0023 (.0005) .0180 (.0024)	.2005 (.0045) .0020 (.0004) .0078 (.0013)	.1361 (.0024) .0009 (.0002) .0033 (.0006)
2.25	.2033 (.0083) .0010 (.0005) .0103 (.0014)	.1833 (.0072) .0010 (.0004) .0318 (.0034)	.1561 (.0049) .0011 (.0003) .0148 (.0020)	.0987 (.0030) .0009 (.0003) .0082 (.0011)
3.38	.1470 (.0083) .0006 (.0002) .0084 (.0013)	.1274 (.0065) .0004 (.0002) .0343 (.0038)	.1081 (.0050) .0007 (.0003) .0318 (.0030)	.0662 (.0028) .0005 (.0002) .0159 (.0016)
5.06	.0984 (.0072) .0000 (.0000) .0045 (.0009)	.1207 (.0060) .0008 (.0004) .0241 (.0032)	.0726 (.0044) .0000 (.0000) .0379 (.0039)	.0457 (.0025) .0002 (.0001) .0203 (.0020)
7.59	.0663 (.0061) .0003 (.0001) .0045 (.0010)	.0564 (.0050) .0002 (.0001) .0184 (.0030)	.0462 (.0038) .0002 (.0001) .0299 (.0041)	.0298 (.0021) .0000 (.0000) .0277 (.0032)
11.39	.0335 (.0043) .0000 (.0000) .0026 (.0013)	.0392 (.0043) .0000 (.0000) .0109 (.0022)	.0287 (.0030) .0000 (.0000) .0248 (.0041)	.0200 (.0019) .0001 (.0001) .0300 (.0038)

TABLE I

The entry on top of each cell is the estimated Bayes risk and the numbers in the second and the third row are the regrets incurred by the optimal δ^m and by the optimal δ^a , respectively. The estimated standard deviations of the estimates are given in the parentheses.

$$k = 9, \Delta = 0.5$$

$\sigma \backslash c$	1	2	4	8
.44	1.5947 (.0148) .0198 (.0033) .0603 (.0052)	1.5322 (.0118) .0296 (.0046) .2614 (.0170)	1.0429 (.0037) .0070 (.0018) .0306 (.0049)	.5953 (.0016) .0011 (.0001) .0017 (.0006)
.67	1.1473 (.0193) .0083 (.0021) .0083 (.0021)	1.2235 (.0173) .0223 (.0041) .2021 (.0099)	.9849 (.0118) .0222 (.0031) .3029 (.0146)	.6407 (.0051) .0112 (.0023) .0735 (.0025)
1.00	.7783 (.0234) .0072 (.0028) .0108 (.0029)	.8272 (.0240) .0238 (.0039) .1071 (.0092)	.7342 (.0179) .0191 (.0035) .3424 (.0164)	.5097 (.0106) .0138 (.0026) .2757 (.0126)
1.50	.5637 (.0237) .0064 (.0021) .0096 (.0028)	.6126 (.0261) .0186 (.0041) .0654 (.0078)	.5207 (.0194) .0107 (.0024) .2469 (.0148)	.3625 (.0118) .0101 (.0023) .4252 (.0230)
2.25	.4021 (.0249) .0043 (.0018) .0054 (.0016)	.3444 (.0202) .0051 (.0024) .0443 (.0068)	.3389 (.0181) .0062 (.0019) .1228 (.0141)	.2324 (.0124) .0067 (.0016) .1368 (.0172)
3.38	.2420 (.0209) .0015 (.0010) .0053 (.0018)	.2662 (.0211) .0032 (.0010) .0424 (.0081)	.1925 (.0154) .0032 (.0013) .0815 (.0130)	.1559 (.0120) .0019 (.0007) .1092 (.0136)
5.06	.1868 (.0191) .0004 (.0003) .0128 (.0029)	.1284 (.0150) .0005 (.0004) .0304 (.0071)	.1590 (.0152) .0032 (.0012) .0683 (.0109)	.0924 (.0093) .0017 (.0008) .0880 (.0134)
7.59	.1114 (.0153) .0001 (.0001) .0041 (.0021)	.0972 (.0134) .0012 (.0009) .0357 (.0078)	.1101 (.0123) .0009 (.0005) .0852 (.0143)	.0583 (.0071) .0003 (.0003) .0783 (.0121)
11.39	.0622 (.0121) .0001 (.0001) .0059 (.0026)	.0754 (.0115) .0000 (.0000) .0236 (.0069)	.0698 (.0113) .0014 (.0007) .0441 (.0099)	.0421 (.0064) .0000 (.0000) .0629 (.0107)

TABLE I

The entry on top of each cell is the estimated Bayes risk and the numbers in the second and the third row are the regrets incurred by the optimal δ^m and by the optimal δ^a , respectively. The estimated standard deviations of the estimates are given in the parentheses.

$$k = 3, \Delta = 1.0$$

$\sigma \backslash c$	1	2	4	8
.44	.1462 (.0018) .0000 (.0000) .0001 (.0001)			
.67	.3213 (.0027) .0016 (.0003) .0315 (.0034)	.2276 (.0023) .0011 (.0003) .0035 (.0009)	.1435 (.0017) .0000 (.0000) .0001 (.0001)	
1.00	.3515 (.0046) .0018 (.0004) .0774 (.0044)	.2969 (.0032) .0023 (.0005) .0293 (.0028)	.1989 (.0019) .0011 (.0003) .0087 (.0012)	.1158 (.0012) .0002 (.0001) .0013 (.0005)
1.50	.2948 (.0066) .0022 (.0005) .0385 (.0034)	.2613 (.0052) .0030 (.0006) .0228 (.0024)	.1922 (.0032) .0022 (.0005) .0135 (.0016)	.1227 (.0018) .0008 (.0002) .0056 (.0008)
2.25	.2117 (.0075) .0009 (.0003) .0389 (.0041)	.1922 (.0061) .0009 (.0003) .0272 (.0032)	.1488 (.0042) .0011 (.0003) .0150 (.0019)	.1015 (.0025) .0009 (.0002) .0067 (.0010)
3.38	.1438 (.0069) .0005 (.0002) .0296 (.0033)	.1344 (.0060) .0008 (.0003) .0435 (.0046)	.1049 (.0040) .0009 (.0003) .0198 (.0024)	.0661 (.0026) .0004 (.0002) .0141 (.0015)
5.06	.0997 (.0067) .0005 (.0003) .0231 (.0030)	.0888 (.0055) .0005 (.0003) .0478 (.0052)	.0682 (.0039) .0004 (.0002) .0385 (.0036)	.0482 (.0025) .0003 (.0001) .0222 (.0020)
7.59	.0525 (.0052) .0001 (.0001) .0182 (.0028)	.0608 (.0046) .0002 (.0001) .0384 (.0047)	.0476 (.0033) .0000 (.0000) .0366 (.0043)	.0335 (.0022) .0001 (.0001) .0251 (.0022)
11.39	.0376 (.0041) .0000 (.0000) .0086 (.0019)	.0416 (.0042) .0001 (.0001) .0329 (.0048)	.0323 (.0029) .0001 (.0001) .0438 (.0059)	.0207 (.0019) .0001 (.0001) .0358 (.0025)

TABLE I

The entry on top of each cell is the estimated Bayes risk and the numbers in the second and the third row are the regrets incurred by the optimal δ^m and by the optimal δ^a , respectively. The estimated standard deviations of the estimates are given in the parentheses.

$$k = 9, \Delta = 1.0$$

$\sigma \backslash c$	1	2	4	8
.44	1.0626 (.0097) .0073 (.0022) .0175 (.0048)	.7235 (.0076) .0000 (.0000) .0005 (.0005)		
.67	1.4650 (.0147) .0256 (.0052) .6476 (.0328)	1.2372 (.0097) .0284 (.0062) .1879 (.0180)	.8253 (.0065) .0053 (.0013) .0242 (.0050)	.4712 (.0045) .0020 (.0009) .0041 (.0015)
1.00	1.0771 (.0300) .0212 (.0038) .1833 (.0116)	1.0578 (.0200) .0283 (.0050) .5966 (.0264)	.8267 (.0119) .0245 (.0045) .3253 (.0190)	.5314 (.0063) .0122 (.0025) .0991 (.0085)
1.50	.7150 (.0280) .0178 (.0042) .1088 (.0095)	.7205 (.0264) .0154 (.0035) .3235 (.0219)	.6264 (.0203) .0119 (.0025) .6319 (.0311)	.3998 (.0117) .0103 (.0020) .3399 (.0165)
2.25	.4724 (.0264) .0071 (.0019) .0685 (.0088)	.4437 (.0269) .0061 (.0019) .1925 (.0177)	.4031 (.0199) .0055 (.0014) .4051 (.0293)	.2852 (.0132) .0082 (.0022) .5105 (.0170)
3.38	.2580 (.0227) .0021 (.0013) .0299 (.0062)	.2926 (.0237) .0050 (.0018) .1266 (.0172)	.2417 (.0154) .0043 (.0013) .2275 (.0239)	.1732 (.0112) .0022 (.0009) .3120 (.0349)
5.06	.1680 (.0167) .0006 (.0005) .0257 (.0061)	.1920 (.0165) .0006 (.0003) .1062 (.0159)	.1211 (.0124) .0024 (.0014) .1277 (.0193)	.1027 (.0078) .0010 (.0006) .1334 (.0158)
7.59	.1149 (.0158) .0011 (.0006) .0279 (.0066)	.1331 (.0158) .0015 (.0007) .0642 (.0120)	.1126 (.0115) .0019 (.0009) .1009 (.0170)	.0566 (.0070) .0005 (.0003) .1124 (.0198)
11.39	.0617 (.0113) .0000 (.0000) .0111 (.0041)	.0668 (.0122) .0001 (.0001) .0262 (.0073)	.0694 (.0088) .0000 (.0000) .0931 (.0185)	.0364 (.0049) .0000 (.0000) .0501 (.0118)

TABLE II

The rows in each cell correspond to the Bayes, the optimal δ^m and the optimal δ^a from top to bottom. The entries in the first column and the second column are the average number of bad populations selected and that of good populations excluded, respectively.

$$k = 3, \Delta = 0.5$$

$\sigma \backslash c$	1		2		4		8	
.44	.7662	.1556	.9525	.0184	.9940	.0009		
	.7661	.1604	.9631	.0140	.9940	.0009		
	.9684	.0202	.9974	.0008	.9950	.0000		
.67	.5349	.4023	.9472	.1109	1.2105	.0172	1.2916	.0024
	.5113	.4337	.9617	.1080	1.2335	.0126	1.3029	.0011
	.5249	.5098	1.1306	.0469	1.2409	.0125	1.2961	.0018
1.00	.3694	.4051	.6767	.1752	.9962	.0649	1.2349	.0179
	.3417	.4399	.6810	.1770	1.0264	.0601	1.2605	.0160
	.3007	.5089	.7794	.1479	1.0775	.0512	1.2254	.0209
1.50	.2529	.3098	.4931	.1599	.7029	.0749	.9663	.0323
	.2518	.3162	.4971	.1614	.7208	.0728	1.0072	.0282
	.1932	.3926	.5378	.1646	.7927	.0622	1.0207	.0292
2.25	.1848	.2219	.3204	.1147	.5417	.0597	.7023	.0233
	.1908	.2178	.3331	.1099	.5727	.0533	.7143	.0227
	.1238	.3034	.3289	.1582	.6485	.0516	.7927	.0211
3.38	.1498	.1443	.2131	.0845	.3733	.0418	.4674	.0168
	.1542	.1412	.2186	.0824	.3740	.0425	.4708	.0162
	.1019	.2089	.1675	.1588	.5551	.0361	.6128	.0158
5.06	.0983	.0984	.1693	.0527	.2370	.0315	.2846	.0159
	.0971	.0996	.1684	.0543	.2351	.0321	.2671	.0183
	.0603	.1454	.0834	.1318	.3569	.0489	.5049	.0111
7.59	.0639	.0686	.0948	.0372	.1453	.0212	.1862	.0102
	.0617	.0714	.0950	.0374	.1460	.0214	.1862	.0102
	.0409	.1006	.0393	.0923	.0997	.0701	.3873	.0163
11.39	.0306	.0364	.0724	.0226	.0919	.0129	.1302	.0062
	.0306	.0364	.0724	.0226	.0919	.0129	.1259	.0069
	.0196	.0504	.0266	.0618	.0387	.0527	.2758	.0218

TABLE II

The rows in each cell correspond to the Bayes, the optimal δ^m and the optimal δ^d from top to bottom. The entries in the first column and the second column are the average number of bad populations selected and that of good populations excluded, respectively.

$$k = 9, \Delta = 0.5$$

$\sigma \backslash c$	1		2		4		8	
.44	.9132	2.2761	3.0660	.7653	4.8169	.0994	5.2816	.0095
	1.0130	2.2159	3.0756	.8049	4.8959	.0883	5.3538	.0016
	.3486	2.9614	5.3807	.0000	5.3676	.0000	5.3721	.0000
.67	.4778	1.8168	1.6250	1.0223	3.3360	.3971	4.9058	.1076
	.3761	1.9351	1.7540	.9918	3.5822	.3633	4.7214	.1432
	.3761	1.9351	.3805	1.9482	6.4389	.0000	6.4282	.0000
1.00	.4314	1.1253	.9308	.1755	2.0683	.4006	3.2729	.1643
	.4686	1.1025	.9346	.8092	2.0634	.4257	3.1622	.1936
	.3322	1.2461	.3246	1.2393	.3307	1.2630	7.0685	.0000
1.50	.3508	.7766	.7108	.5636	1.3636	.3100	2.2131	.1312
	.3322	.8080	.6561	.6188	1.4223	.3087	2.2498	.1379
	.2804	.8662	.2761	.8789	.2985	.8849	.5227	.8208
2.25	.2795	.5248	.4351	.2990	.9762	.1795	1.3742	.0896
	.2937	.5191	.4602	.2941	.9984	.1817	1.3010	.1063
	.2149	.6003	.1661	.4992	.7470	.3903	1.4132	.2386
3.38	.1924	.2916	.3815	.2086	.6387	.0810	.8702	.0666
	.1990	.2881	.3913	.2085	.6179	.0902	.8632	.0696
	.1477	.3469	.2572	.3343	.5281	.2105	1.1838	.1503
5.06	.2040	.1697	.1862	.0995	.4834	.0779	.5406	.0364
	.2043	.1701	.1867	.1000	.5266	.0711	.5690	.0348
	.1218	.2775	.0833	.1966	.3837	.1881	.5487	.1345
7.59	.0900	.1328	.1805	.0556	.3384	.0530	.3745	.0188
	.0951	.1279	.1950	.0501	.3313	.0559	.3659	.0202
	.0641	.1668	.0695	.1646	.2955	.1702	.7284	.0827
11.39	.0639	.0605	.1262	.0501	.2148	.0335	.2707	.0135
	.0589	.0655	.1262	.0501	.1842	.0429	.2618	.0147
	.0398	.0964	.0497	.1236	.0590	.1277	.4669	.0597

TABLE II

The rows in each cell correspond to the Bayes, the optimal δ^m and the optimal δ^a from top to bottom. The entries in the first column and the second column are the average number of bad populations selected and that of good populations excluded, respectively.

$$k = 3, \Delta = 1.0$$

$\sigma \backslash c$	1		2		4		8	
.44	.2912	.0012						
	.2912	.0002						
	.2925	.0000						
.67	.5272	.1155	.6525	.0152	.7129	.0011		
	.5475	.0984	.6636	.0113	.7129	.0011		
	.6861	.0195	.6861	.0037	.7169	.0002		
1.00	.4200	.2830	.7025	.0941	.8704	.0311	1.0064	.0045
	.4355	.2710	.7314	.0831	.8795	.0302	1.0111	.0041
	.6386	.2191	.8651	.0568	.9811	.0143	1.0188	.0044
1.50	.3104	.2793	.5297	.1270	.7252	.0590	.9434	.0202
	.3064	.2877	.5494	.1218	.7453	.0566	.9642	.0184
	.3776	.2890	.6358	.1082	.8686	.0399	1.0624	.0099
2.25	.2184	.2049	.3531	.1118	.5182	.0564	.7213	.0240
	.2093	.2158	.3389	.1201	.5113	.0595	.7200	.0252
	.1772	.3238	.4469	.1056	.6153	.0510	.8235	.0187
3.38	.1356	.1519	.2606	.0713	.3476	.0442	.4153	.0225
	.1324	.1562	.2480	.0787	.3564	.0431	.4268	.0215
	.0640	.2828	.3974	.0681	.4713	.0380	.5916	.0163
5.06	.0964	.1030	.1598	.0533	.2341	.0268	.3199	.0142
	.0970	.1034	.1669	.0504	.2305	.0282	.3180	.0148
	.0358	.2098	.0543	.1778	.3926	.0353	.5687	.0080
7.59	.0674	.0577	.1089	.0368	.1583	.0200	.2426	.0073
	.0638	.0615	.1141	.0345	.1603	.0195	.2449	.0071
	.0281	.1333	.0306	.1335	.2729	.0371	.4766	.0064
11.39	.0403	.0349	.0837	.0206	.1065	.0138	.1402	.0058
	.0403	.0349	.0855	.0198	.1107	.0129	.1425	.0056
	.0140	.0785	.0199	.1018	.0363	.0861	.4804	.0035

TABLE II

The rows in each cell correspond to the Bayes, the optimal δ^m and the optimal δ^a from top to bottom. The entries in the first column and the second column are the average number of bad populations selected and that of good populations excluded, respectively.

$$k = 9, \Delta = 1.0$$

$\sigma \backslash c$	1		2		4		8	
.44	1.9926	.1325	2.1648	.0029				
	2.0050	.0849	2.1648	.0029				
	2.1601	.0000	2.1719	.0000				
.67	1.6325	1.2974	2.8473	.4321	3.8555	.0678	4.1572	.0104
	1.6381	1.3430	2.8424	.4772	3.9488	.0511	4.1682	.0113
	4.2251	.0000	4.2752	.0000	4.2477	.0000	4.2768	.0000
1.00	.8150	1.3392	1.7229	.7252	2.9470	.2967	4.0555	.0909
	.8561	1.3404	1.8779	.6902	3.0915	.2912	3.9255	.1209
	.1683	2.3525	.1729	2.3952	5.7603	.0000	5.6146	.0000
1.50	.5275	.9024	1.0839	.5388	2.0030	.2823	2.6489	.1187
	.4802	.9852	1.1259	.5409	1.9110	.3202	2.8192	.1090
	.1563	1.4912	.1540	1.4890	.1469	1.5362	6.6510	.0008
2.25	.3851	.5597	.6039	.3637	1.1882	.2068	1.8056	.0951
	.3572	.6019	.5833	.3831	1.2177	.2063	1.8583	.0978
	.1286	.9532	.1097	.8995	.1212	.9799	7.1606	.0000
3.38	.1921	.3239	.4279	.2250	.7161	.1232	1.1018	.0571
	.1992	.3210	.4329	.2300	.7603	.1174	1.0418	.0671
	.0720	.5038	.1345	.5616	.3436	.5007	1.1293	.4047
5.06	.1376	.1984	.3050	.1355	.3050	.0752	.6648	.0324
	.1483	.1891	.2722	.1527	.3234	.0735	.6924	.0301
	.0583	.3291	.0912	.4017	.2887	.2388	1.5093	.0769
7.59	.1225	.1073	.2137	.0928	.3130	.0625	.3697	.0174
	.1086	.1234	.2286	.0876	.3149	.0644	.3791	.0169
	.0455	.2403	.0713	.2604	.3099	.1894	.7931	.0909
11.39	.0569	.0665	.1308	.0349	.2497	.0243	.2103	.0147
	.0569	.0665	.1376	.0316	.2497	.0243	.2192	.0136
	.0230	.1225	.0236	.1227	.0788	.1834	.2337	.0682

TABLE III

Estimated d values for the optimal δ^m

$\sigma \backslash c$	1	2	4	8	1	2	4	8
	k = 3, $\Delta = 0.5$				k = 9, $\Delta = 0.5$			
.44	2.1572	3.7576	4.6502		.7906	2.0710	3.4991	5.0995
.67	1.0391	2.1839	3.3917	4.2796	.0000	1.2464	2.2379	2.8148
1.00	.6040	1.4808	2.3435	3.1461	.2222	.8727	1.6505	2.2304
1.50	.4939	1.2570	1.9361	2.6181	.1333	.6862	1.4193	2.0923
2.25	.5304	1.1760	1.8682	2.2840	.2595	.9927	1.4305	1.9776
3.38	.4552	1.1332	1.6755	2.3222	.4240	.9194	1.5243	1.8893
5.06	.4940	1.0801	1.6866	2.1657	.4329	.8559	1.3859	1.9873
7.59	.3775	1.0029	1.6789	2.2435	.4683	1.0987	1.3360	1.9157
11.39	.4938	1.1011	1.6307	2.0799	.2981	1.1815	1.3571	2.0548
	k = 3, $\Delta = 1.0$				k = 9, $\Delta = 1.0$			
.44	5.3177				4.2282	5.2869		
.67	2.9075	3.9441	5.0438		1.8799	2.9796	4.1874	5.1159
1.00	1.7525	2.6293	3.3965	4.1991	1.0949	2.0000	2.7778	3.1879
1.50	1.2341	2.0514	2.6403	3.3854	.7834	1.5887	2.2136	2.8867
2.25	1.0389	1.6517	2.3356	2.9320	.7598	1.3617	2.0292	2.6420
3.38	.9678	1.5832	2.2664	2.8035	.7958	1.3381	2.0995	2.4958
5.06	.9447	1.6608	2.1144	2.6623	1.0135	1.4671	2.0125	2.7515
7.59	.9013	1.5872	2.1949	2.7446	.7652	1.4258	1.9957	2.6211
11.39	.9977	1.5598	2.2023	2.7268	1.0730	1.6401	2.2073	2.7092

FIG. 1a. Estimated risks of Bayes, optimal Gupta-type and optimal Seal-type rules.

$$k = 3, \Delta = 1.0, c = 1.0$$

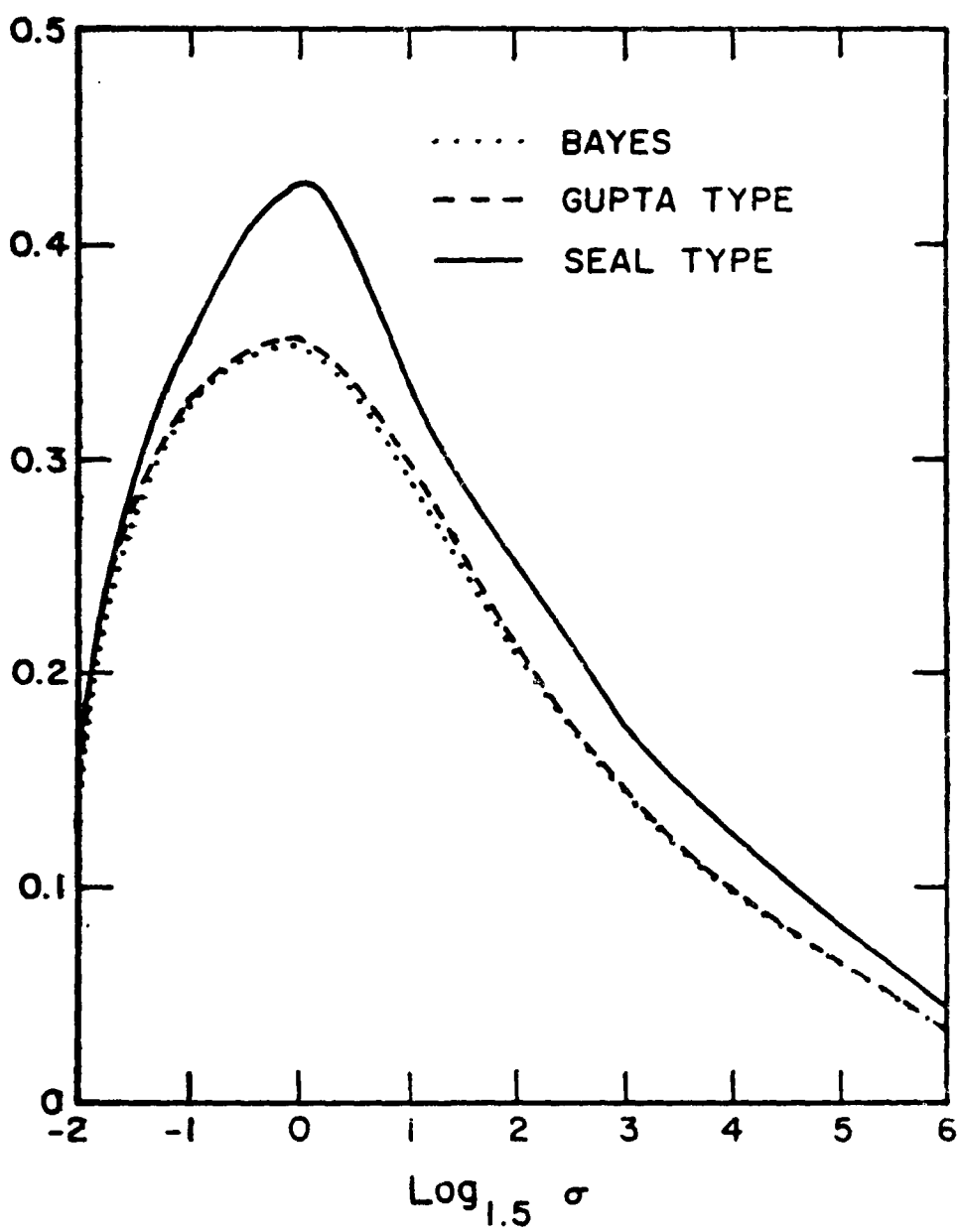


FIG. 1b. Estimated risks of Bayes, optimal Gupta-type and optimal Seal-type rules.

$$k = 3, \Delta = 1.0, c = 4.0$$

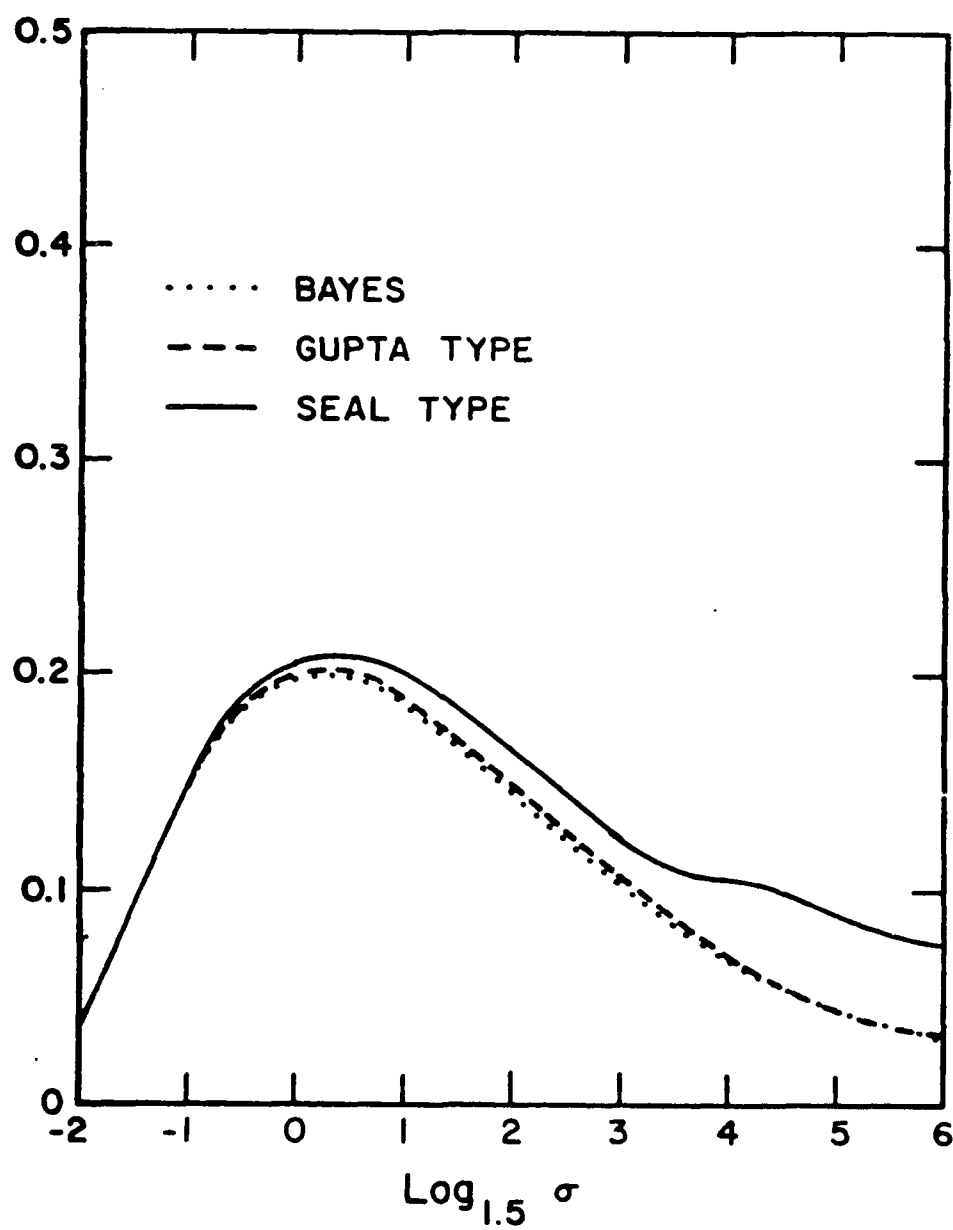


FIG. 1c. Estimated risks of Bayes, optimal Gupta-type and optimal Seal-type rules.

$$k = 3, \Delta = 1.0, c = 8.0$$

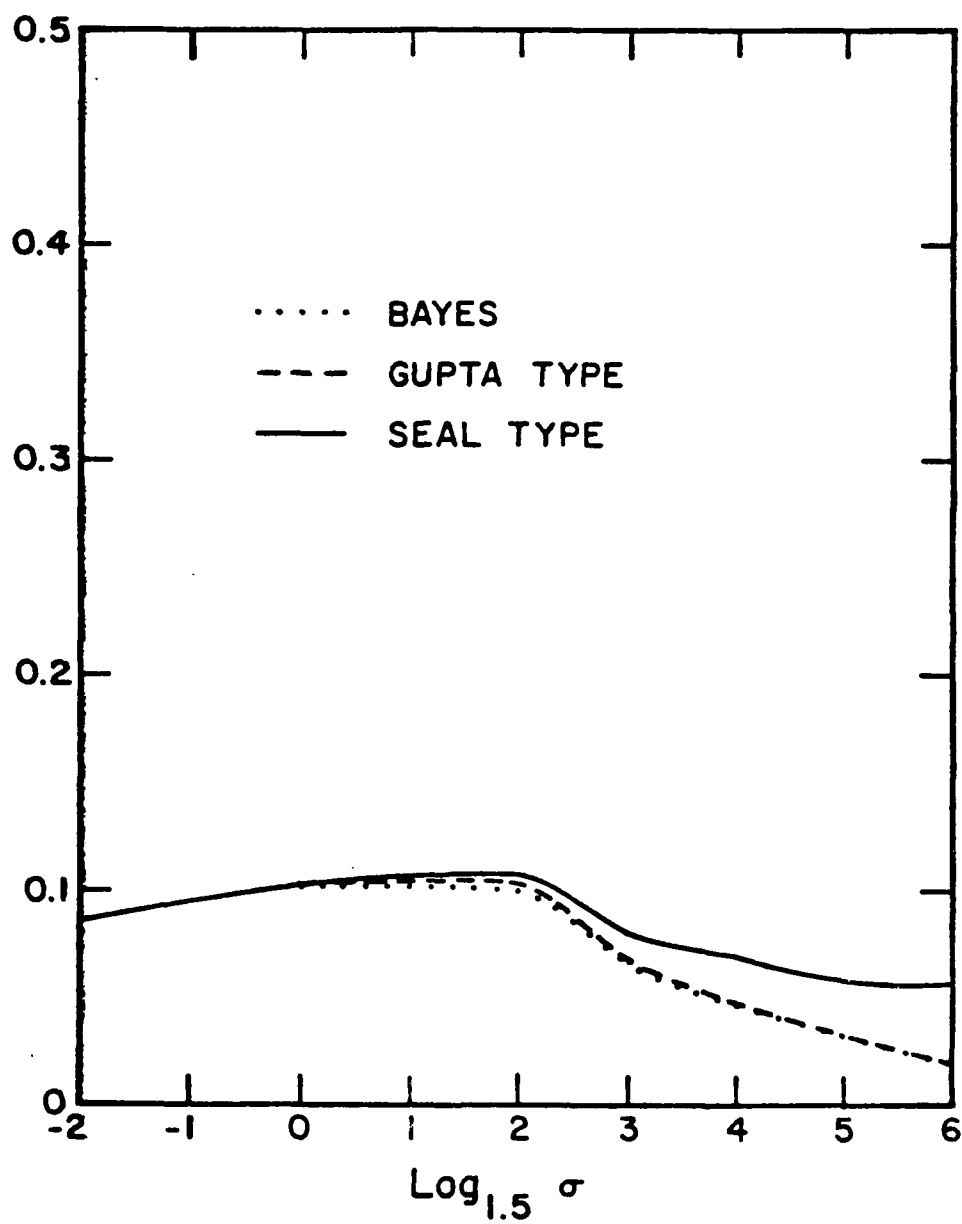


FIG. 1d. Estimated risks of Bayes, optimal Gupta-type and optimal Seal-type rules.

$$k = 9, \Delta = 1.0, c = 1.0$$

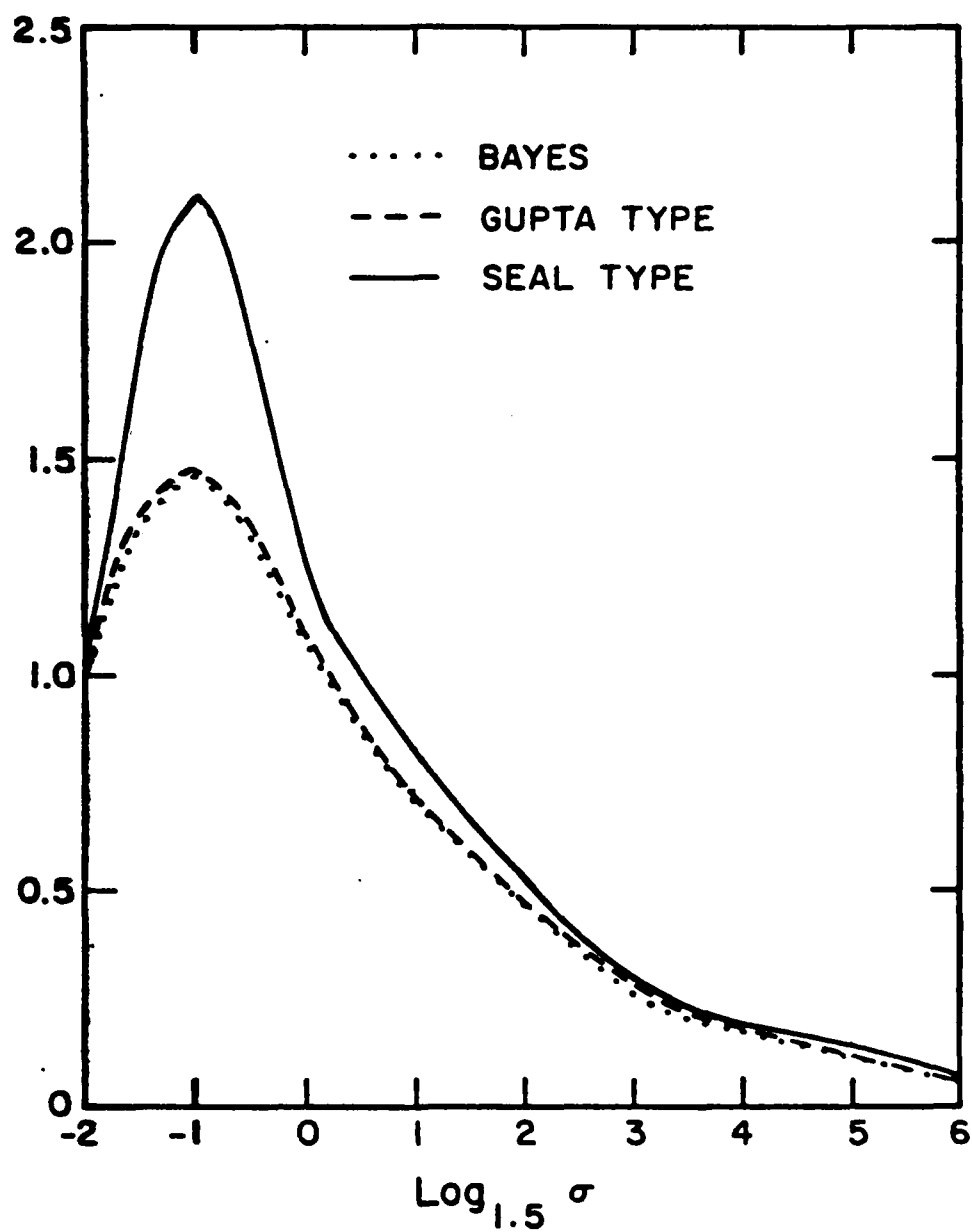


FIG. 1e. Estimated risks of Bayes, optimal Gupta-type and optimal Seal-type rules.

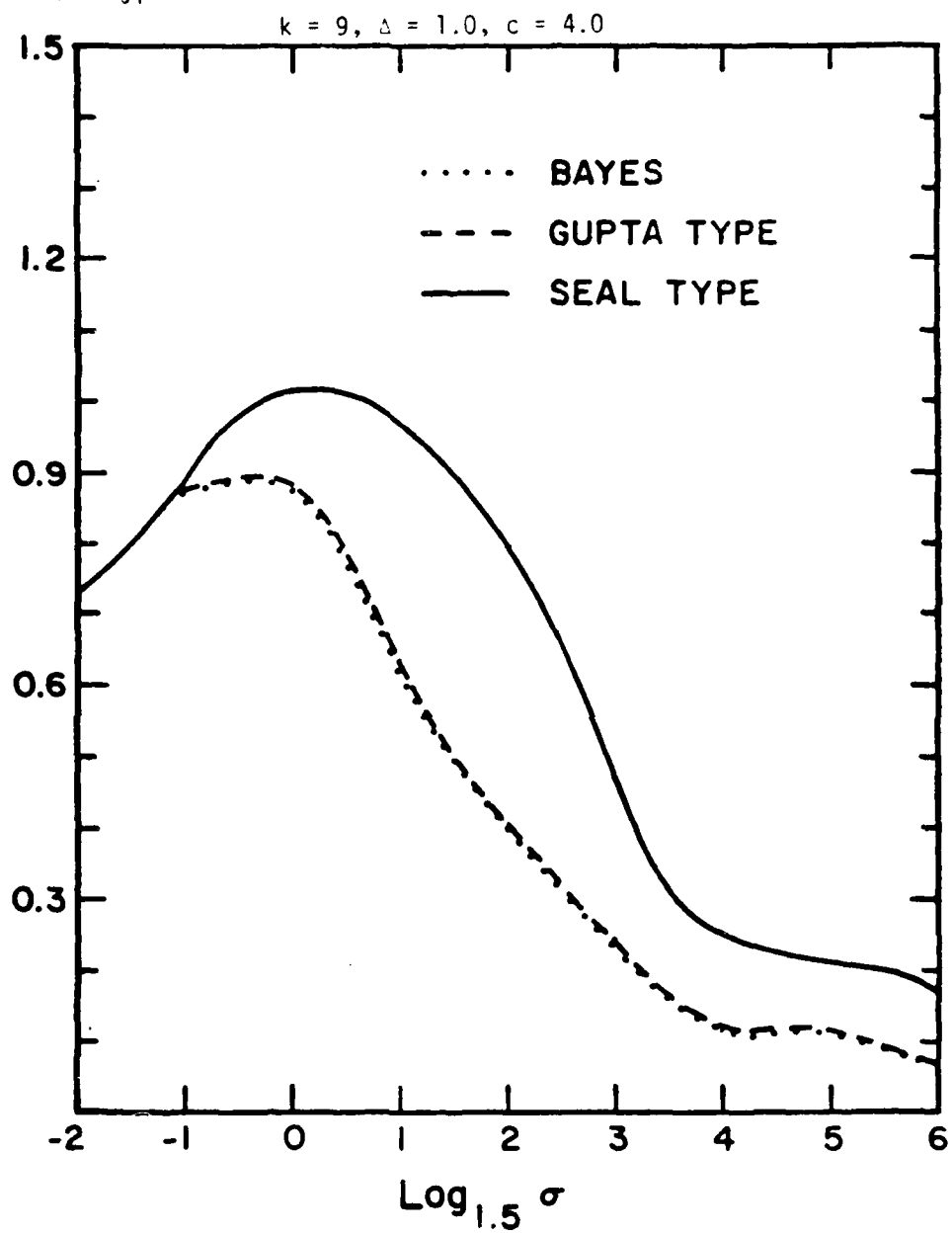


FIG. 1f. Estimated risks of Bayes, optimal Gupta-type and optimal Seal-type rules.

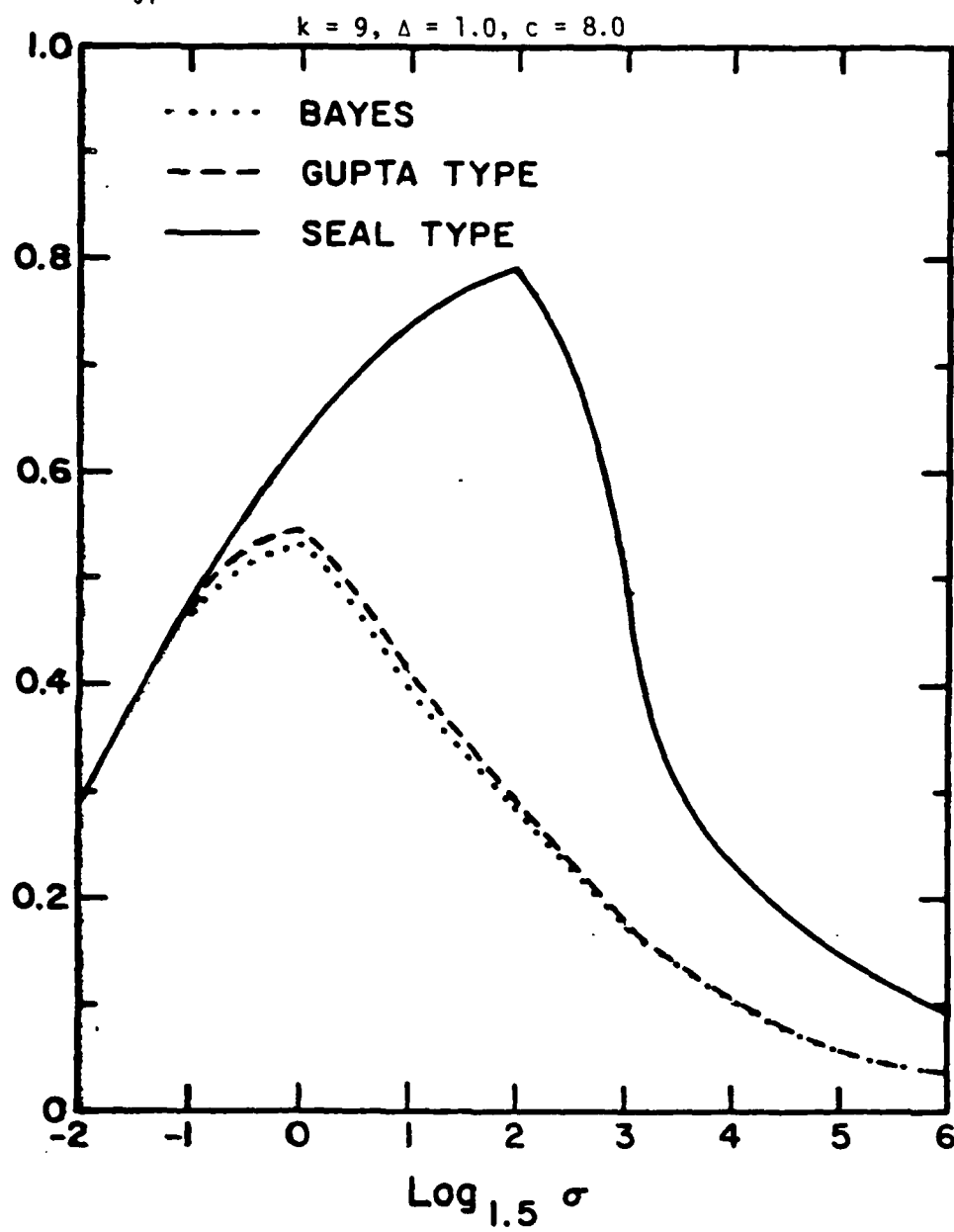


FIG. 2a. Proportion of times the optimal Gupta-type rule coincides with the Bayes rule.

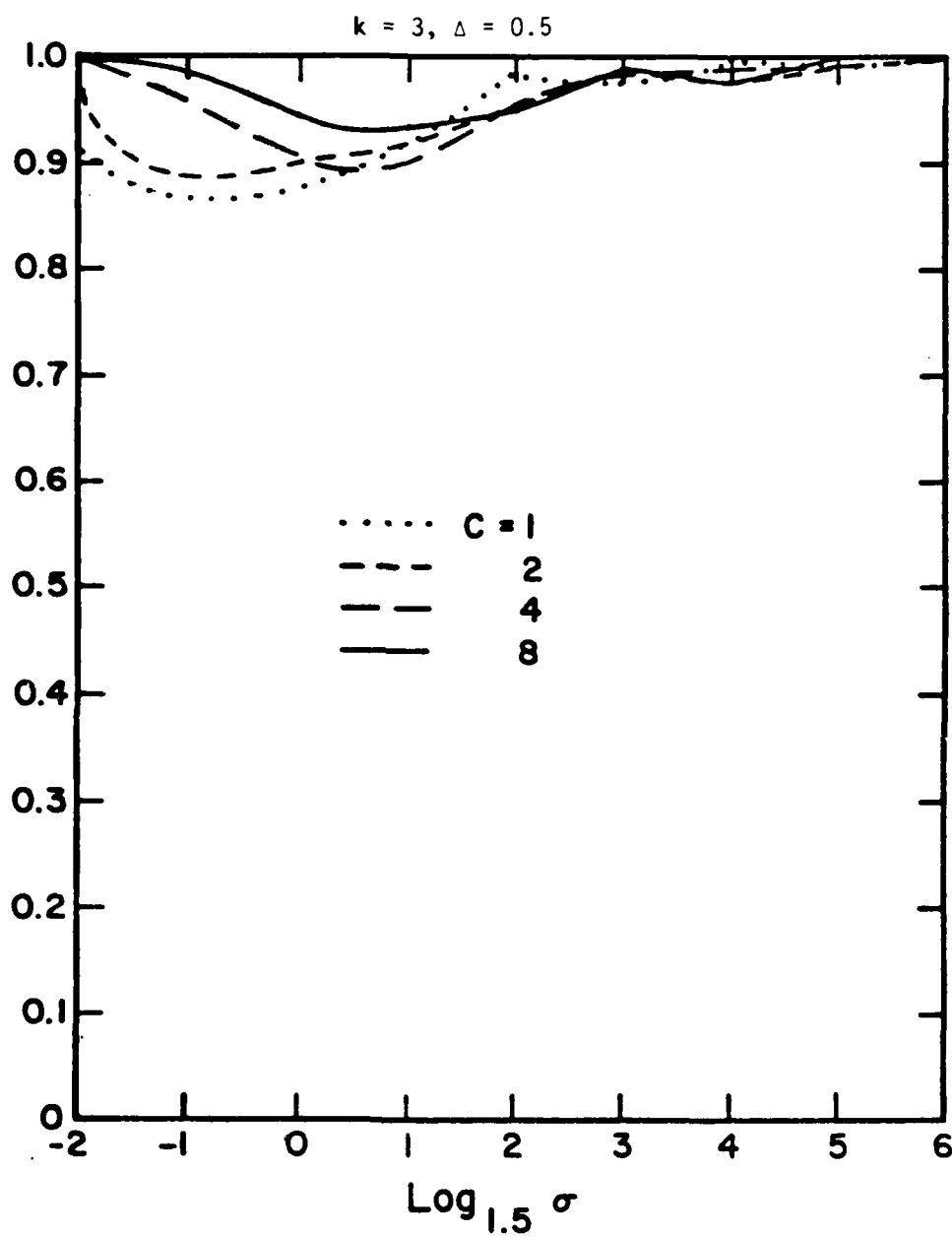


FIG. 2b. Proportion of times the optimal Seal-type rule coincides with the Bayes rule.

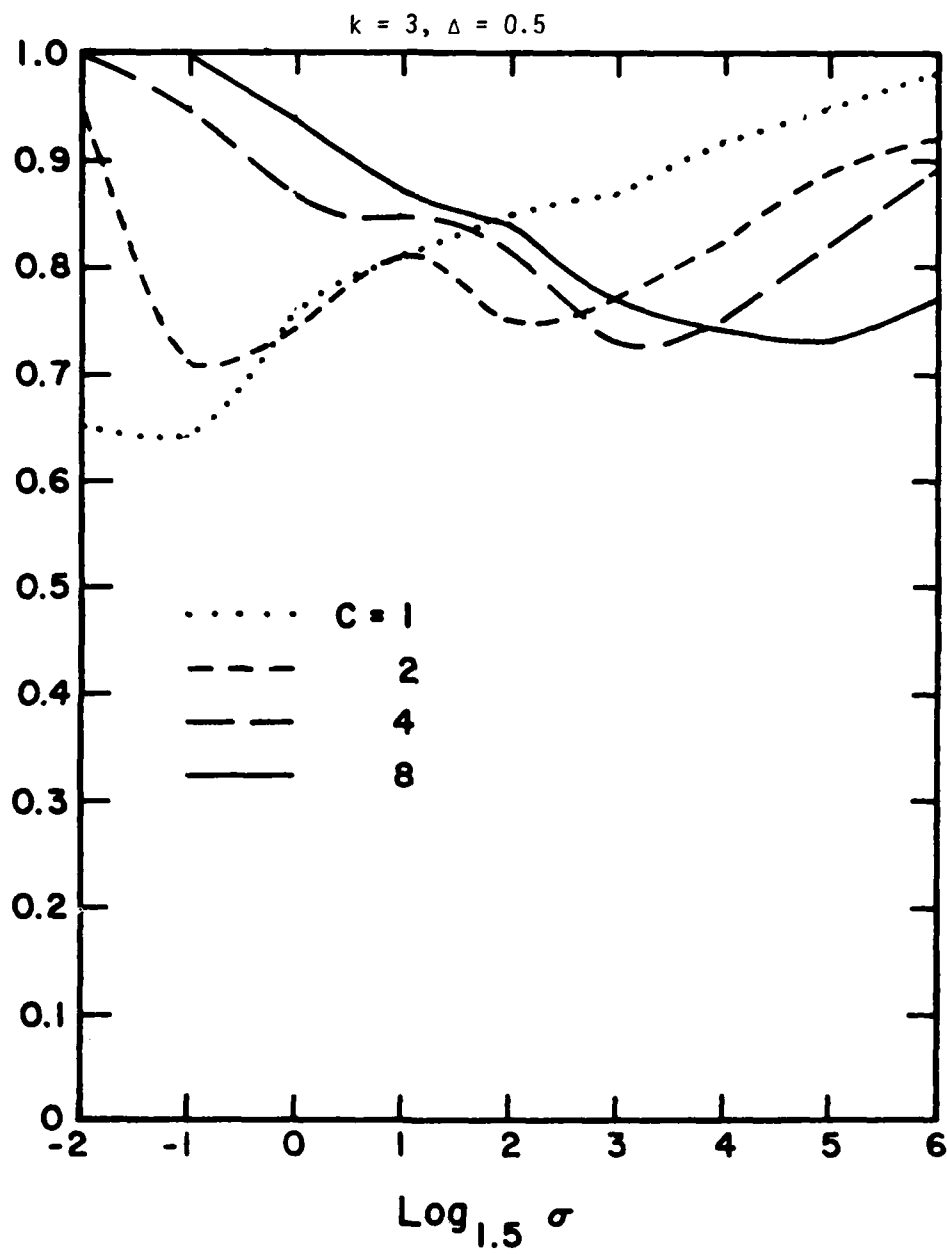


FIG. 2c. Proportion of times the optimal Gupta-type rule coincides with the Bayes rule.

$$k = 9, \Delta = 0.5$$

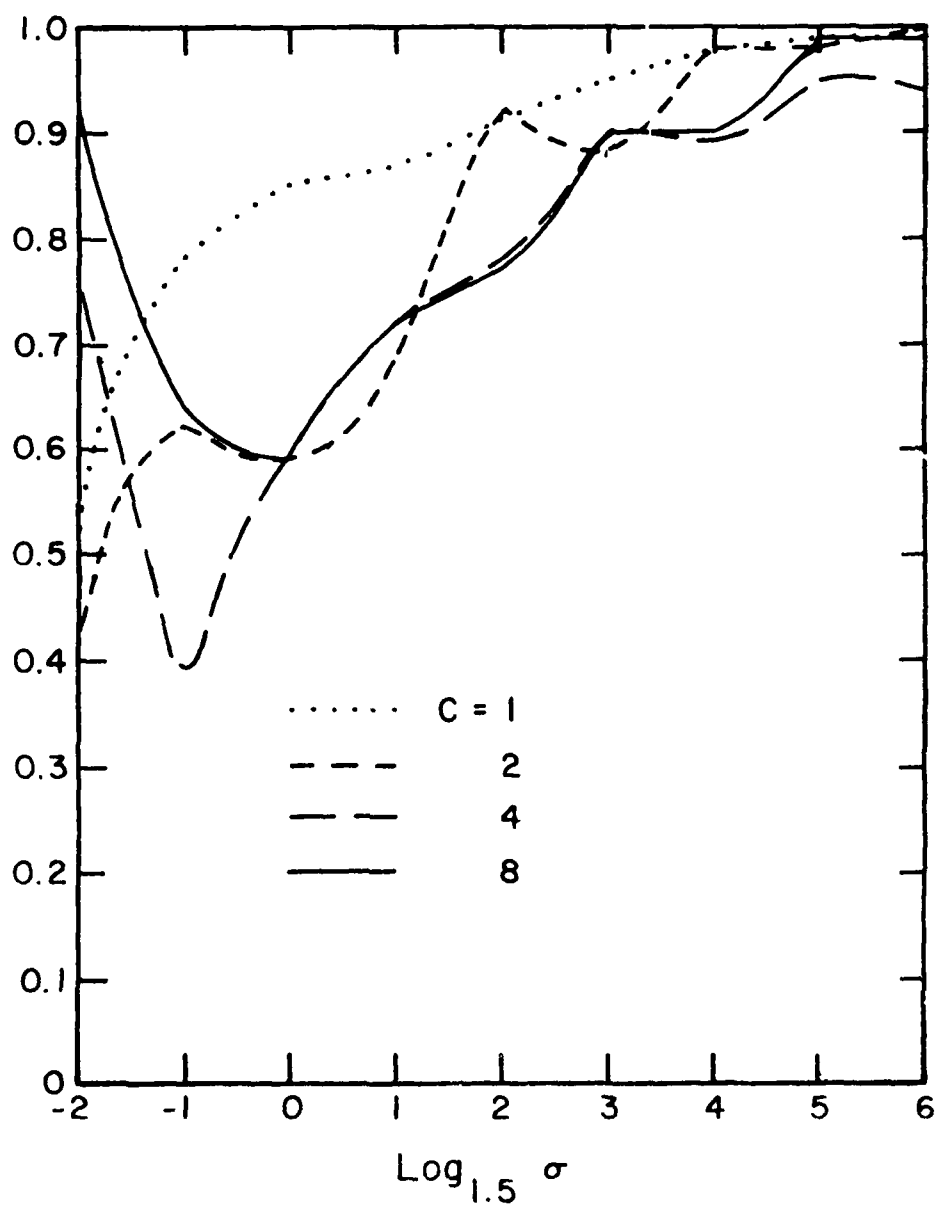


FIG. 2d. Proportion of times the optimal Seal-type rule coincides with the Bayes rule.

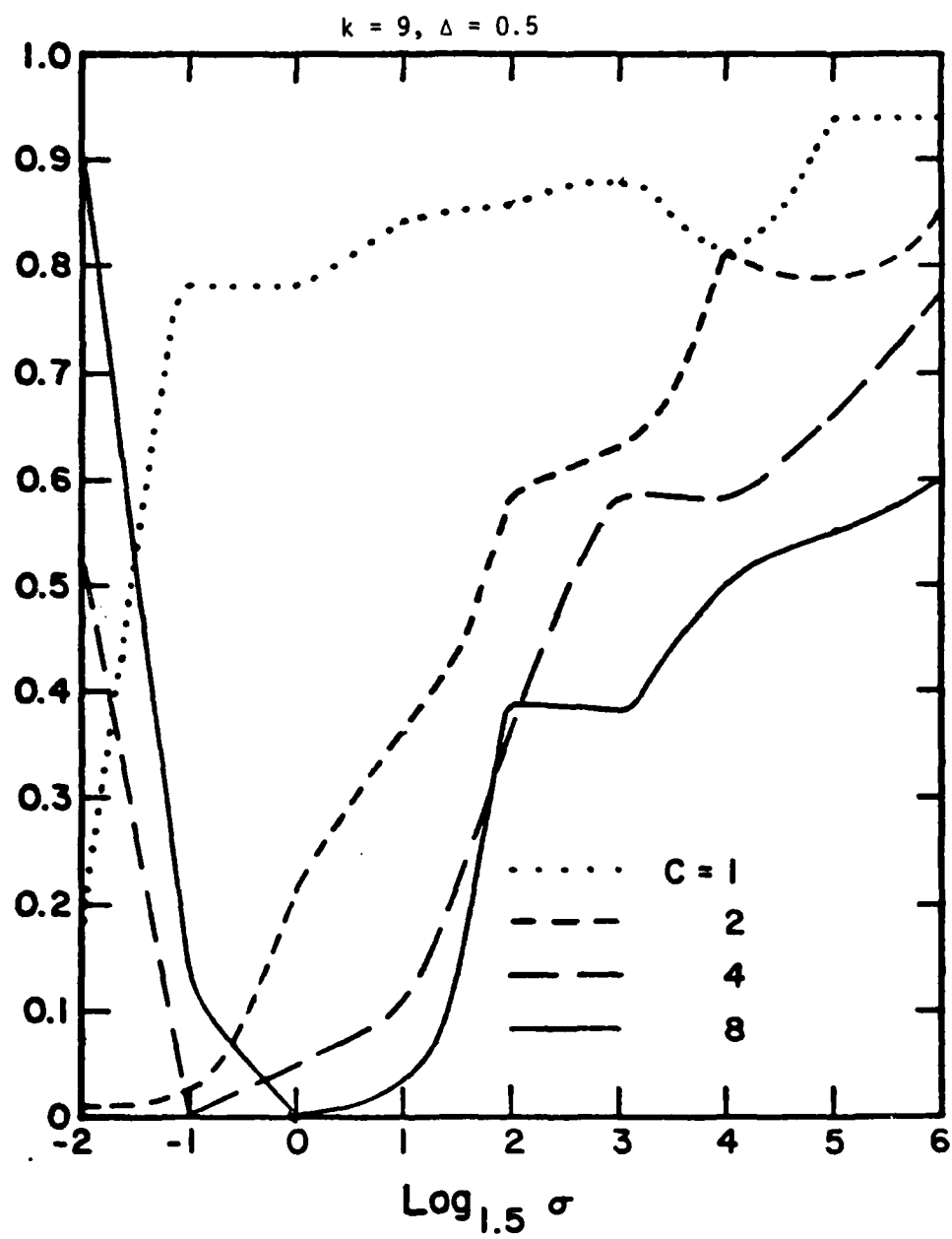


FIG. 2e. Proportion of times the optimal Gupta-type rule coincides with the Bayes rule.

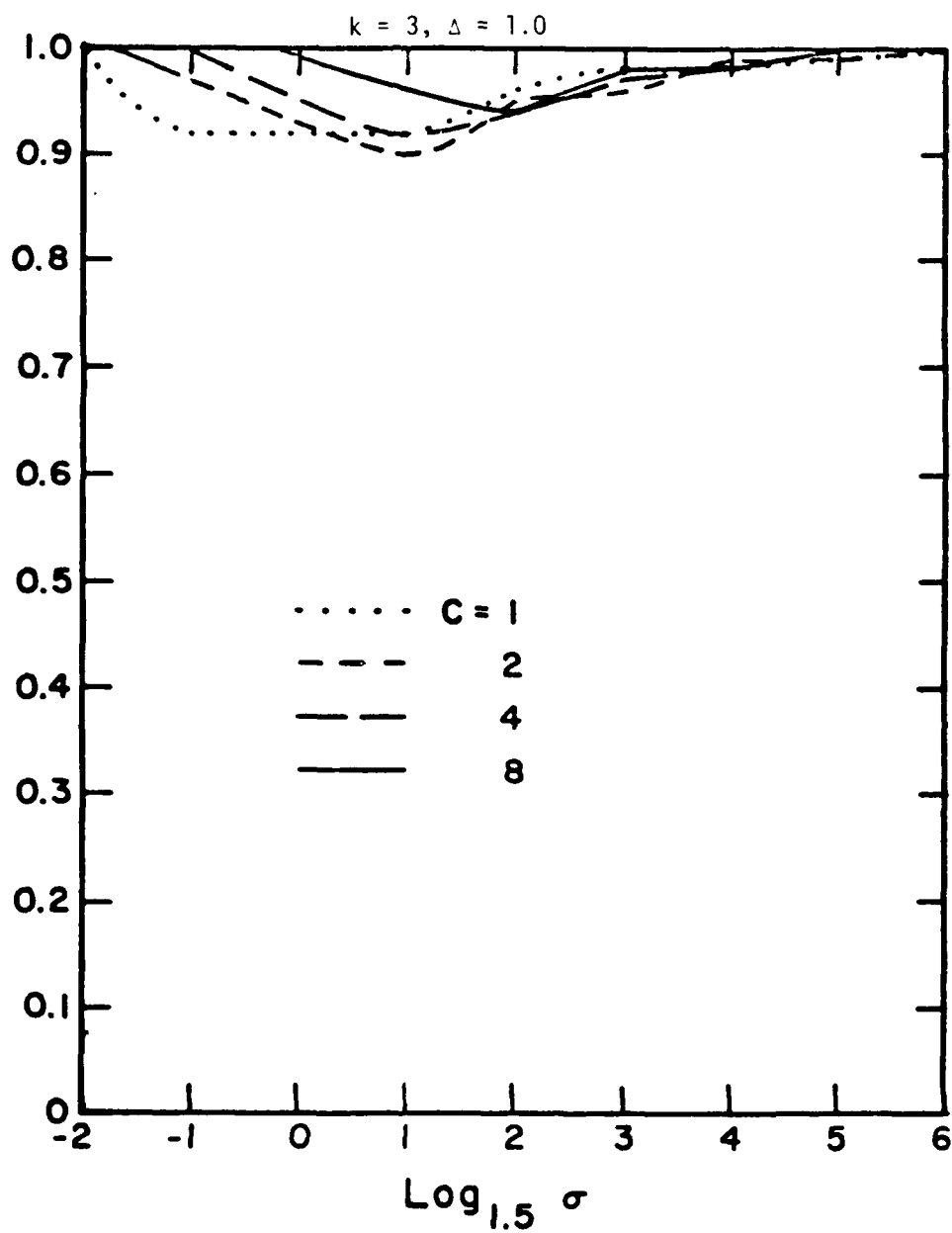


FIG. 2f. Proportion of times the optimal Seal-type rule coincides with the Bayes rule.

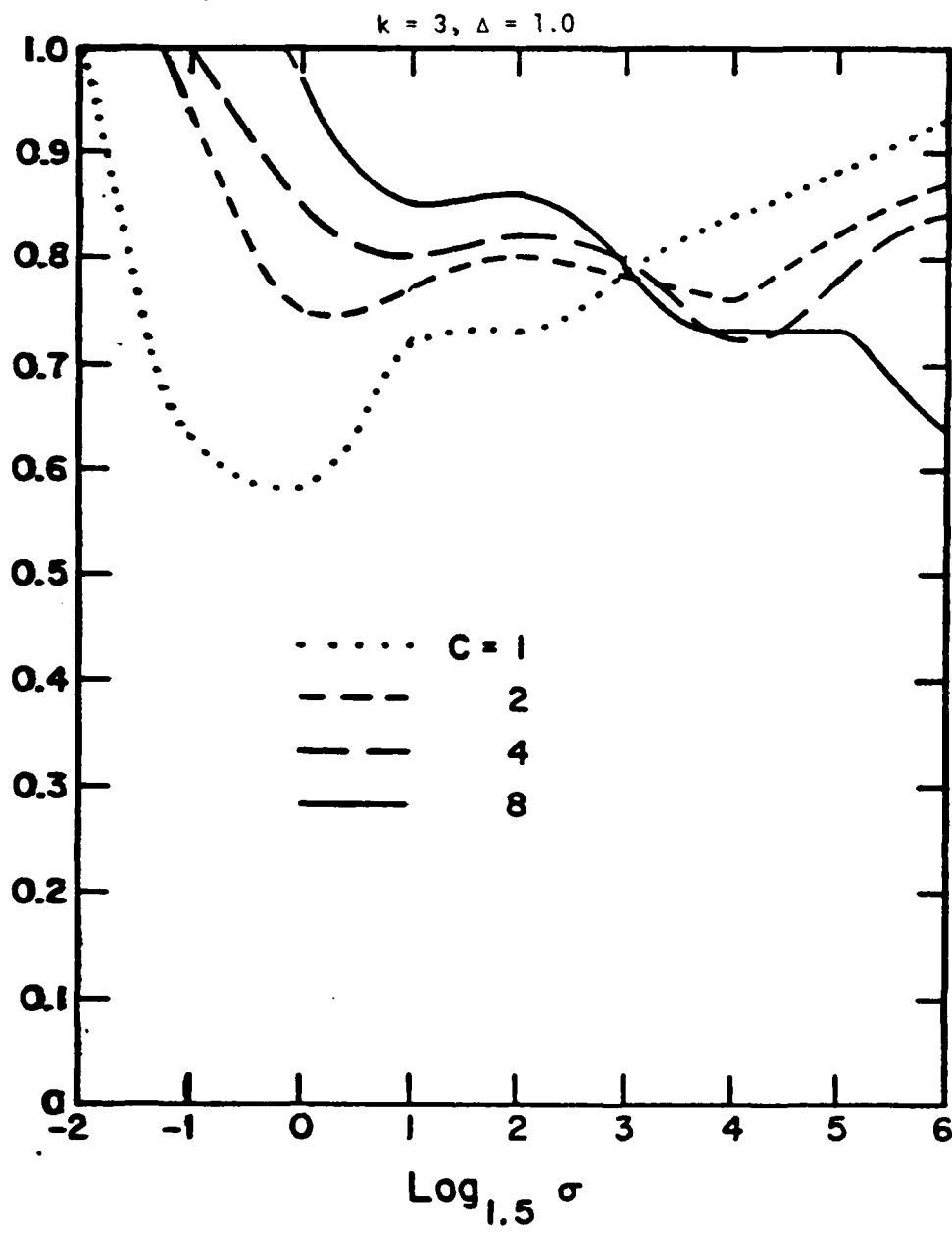


FIG. 2g. Proportion of times the optimal Gupta-type rule coincides with the Bayes rule.

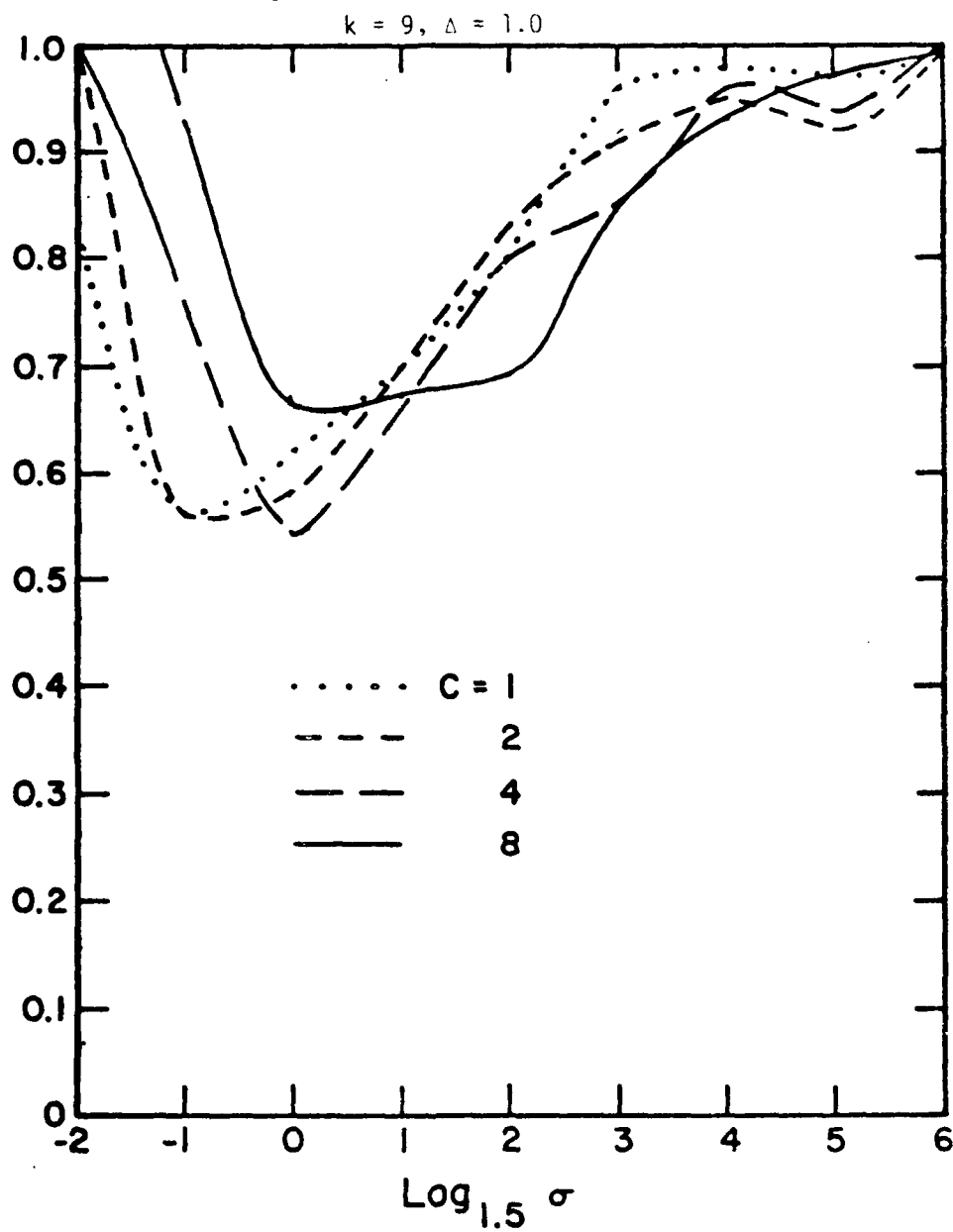
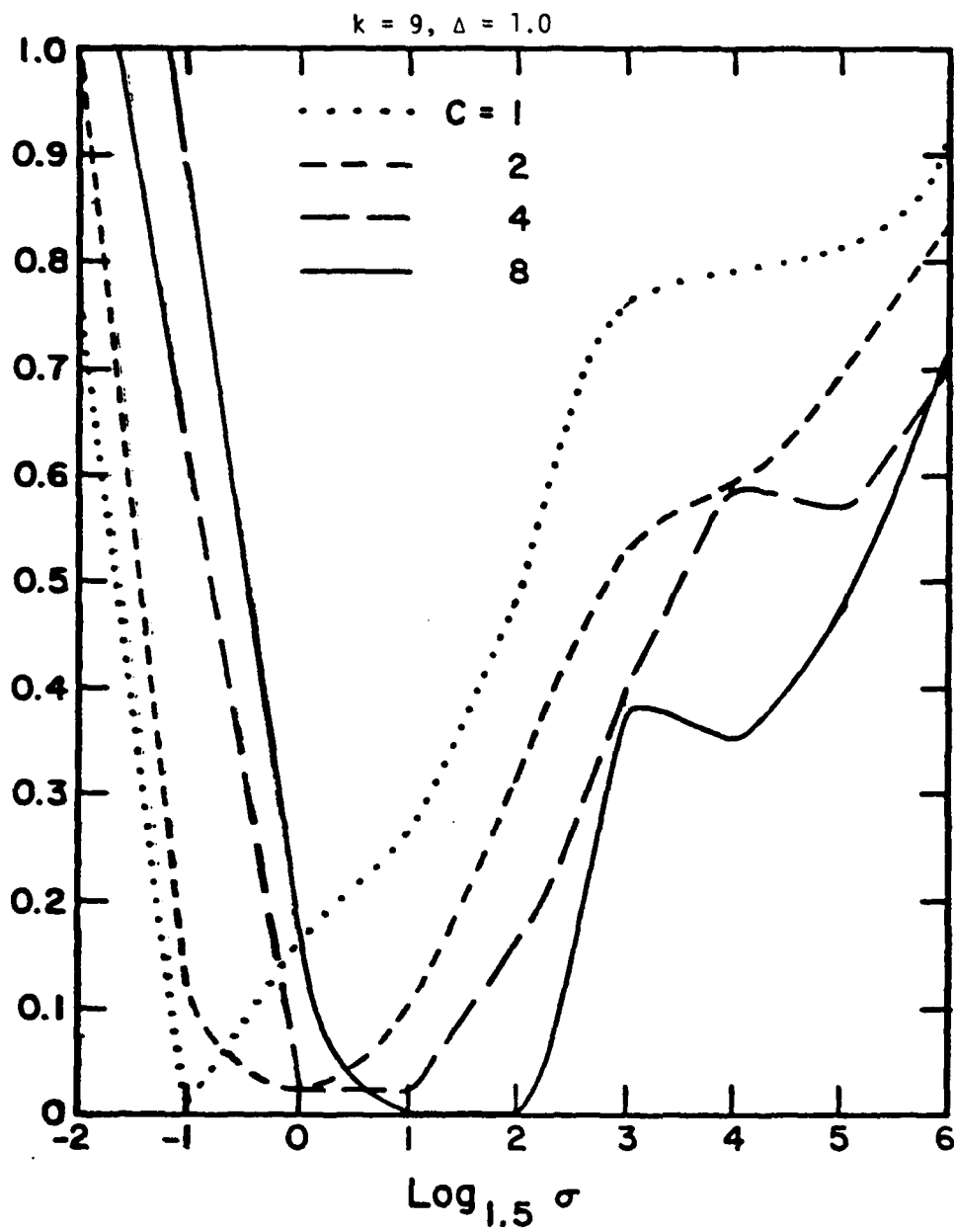


FIG. 2h. Proportion of times the optimal Seal-type rule coincides with Bayes rule.



UNCLASSIFIED

SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

REPORT DOCUMENTATION PAGE		READ INSTRUCTIONS SERIALS ACQUISITION FORM
1. REPORT NUMBER Mimeograph-Series #80-20	2. GOVT ACCESSION NO. AD-A100947	3. REPORTS CATEGORY NUMBER
4. TITLE (and Subtitle) On the problem of selecting good populations	5. TYPE OF REPORT (A, D, etc.) Technical	6. PERFORMING ORG. REPORT NUMBER Mimeo. Series # 80-20
7. AUTHOR(s) Shanti S. Gupta and Won-Chul Kim	8. CONTRACT OR GRANT NUMBER ONR N00014-78-1-0049E	
9. PERFORMING ORGANIZATION NAME AND ADDRESS Purdue University Department of Statistics, West Lafayette, IN 47907	10. PROGRAM ELEMENT PROJECT TASK AREA & WORK UNIT NUMBERS see last page	
11. CONTROLLING OFFICE NAME AND ADDRESS Office of Naval Research Washington, DC	12. REPORT DATE August 1980	13. NUMBER OF PAGES 36
14. MONITORING AGENCY NAME & ADDRESS (if different from Controlling Office)	15. SECURITY CLASS. of this report Unclassified	16. DECLASSIFICATION/DOWNGRADING SCHEDULE
16. DISTRIBUTION STATEMENT (of this Report) Approved for public release, distribution unlimited.		
17. DISTRIBUTION STATEMENT (of the abstract entered in Block 20, if different from Report)		
18. SUPPLEMENTARY NOTES		
19. KEY WORDS (Continue on reverse side if necessary and identify by block number) Subset selection of good populations; additive loss; Bayes rules; extended Bayes rules; Gupta-type rules; Seal-type rules; Monte Carlo comparison.		
20. ABSTRACT (Continue on reverse side if necessary and identify by block number) The problem of selecting good populations out of k normal populations is considered in a Bayesian framework under exchangeable normal priors and additive loss functions. Some basic approximations to the Bayes rules are discussed. These approximations suggest that some well-known classical rules are "approximate" Bayes rules. Especially, it is shown that Gupta-type rules are extended Bayes with respect to a family of the exchangeable normal priors for any bounded and additive loss function. Furthermore, for a simple loss function, the		

DD FORM 1473

1 JAN 73

EDITION OF 1 NOV 69 IS OBSOLETE
GPO 1102-014-6601

DISTRIBUTION STATEMENT

SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

UNCLASSIFIED

SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

results of a Monte Carlo comparison of Gupta-type rules and Seal-type rules are presented. They indicate that, in general, Gupta-type rules perform better than Seal-type rules.

UNCLASSIFIED

SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

```

-- 1 OF 1
-- 11 - TITLE (U) MULTIPLE DECISION THEORY, ORDER STATISTICS AND RELATED
-- PROBLEMS
-- 1 - AGENCY ACCESSION NO. DN523558
-- 6 - SECURITY OF WORK UNCLASSIFIED
-- 12 - 5 + 7 AREAS.
-- 009700 MATHEMATICS AND STATISTICS
-- 011700 OPERATIONS RESEARCH
-- 21E - MILITARY/CIVILIAN APPLICATIONS CIVILIAN
-- 2 - DATE OF SUMMARY 06 JAN 81
-- 39 - PROCESSING DATE (RANGE) 14 JAN 81
-- 10A1 - PRIMARY PROGRAM ELEMENT 61153H
-- 10A2 - PRIMARY PROJECT NUMBER RR01405
-- 10A24 - PRIMARY PROJECT AGENCY AND PROGRAM RR01405
-- 10A3 - PRIMARY TASK AREA RR0140501
-- 10A4 - WORK UNIT NUMBER NP-042-243
-- 17A1 - CONTRACT/GRANT EFFECTIVE DATE JUL 75
-- 17A2 - CONTRACT/GRANT EXPIRATION DATE JAN 84
-- 17B - CONTRACT/GRANT NUMBER N00014-75-C-0455
-- 17C - CONTRACT TYPE COST TYPE
-- 17D2 - CONTRACT/GRANT AMOUNT $ 53,166
-- 17E - KIND OF AWARD EXT
-- 17F - CONTRACT/GRANT CUMULATIVE DOLLAR TOTAL $ 357,193

-- 19A - DOD ORGANIZATION OFFICE OF NAVAL RESEARCH (456)
-- 19B - DOD ORG ADDRESS ARLINGTON, VA 22217
-- 19C - RESPONSIBLE INDIVIDUAL WEGMAN, E J
-- 19D - RESPONSIBLE INDIVIDUAL PHONE 202-696-4315
-- 19U - DOD ORGANIZATION LOCATION CODE 5110
-- 19S - DOD ORGANIZATION SORT CODE 35832
-- 19T - DOD ORGANIZATION CODE 265250
-- 20A - PERFORMING ORGANIZATION PURDUE UNIVERSITY DEPT OF STATISTICS
-- 20B - PERFORMING ORG ADDRESS LAFAYETTE, IN 47907
-- 20C - PRINCIPAL INVESTIGATOR GUPTA, S S
-- 20D - PRINCIPAL INVESTIGATOR PHONE 317-494-8622
-- 20U - PERFORMING ORGANIZATION LOCATION CODE 1802
-- 20A - PERFORMING ORGANIZATION TYPE CODE 0
-- 20S - PERFORMING ORG SORT CODE 39418
-- 20T - PERFORMING ORGANIZATION CODE 291736
-- 22 - KEYWORDS (U) TWO-STAGE PROCEDURES (U) MULTIPLE DECISION (U)
-- RANKING AND SELECTION (U) ORDER STATISTICS (U) RELIABILITY
-- 37 - DESCRIPTORS (U) DECISION MAKING (U) *DECISION THEORY (U)
-- DISTRIBUTION THEORY (U) GOVERNMENT PROCUREMENT (U) *LOGISTICS
-- (U) *OPERATIONS RESEARCH (U) PROBABILITY (U) QUALITY CONTROL
-- (U) RELIABILITY (U) SAMPLING (U) STANDARDS (U) STATISTICAL
-- ANALYSIS
--*****

```

APPROVED FOR RELEASE
DISSEMINATION RESTRICTED

