

**NAVAL AEROSPACE MEDICAL RESEARCH LABORATORY
51 HOVEY ROAD, PENSACOLA, FL 32508-1046**

NAMRL Special Report 00-3

**AN UPDATE TO THE LANDING CRAFT AIR CUSHION (LCAC)
NAVIGATOR SELECTION SYSTEM PREDICTION ALGORITHM**

D. J. Blower

ABSTRACT

This report updates the status of the Landing Craft Air Cushion Vehicle (LCAC) Navigator Selection System prediction algorithm. The last revision took place at the end of 1997. With the receipt of new training data, the prediction algorithm was changed in June 2000 to take advantage of these new data. Some rough estimates of the new attrition rate and the rejection rate due to the selection system may be made on the basis of the new data. The attrition rate is estimated as 17.24% and the rejection rate as 36.96%. The failure rate during training appears to be about 39.13%. If the rejection rate of about 37% is acceptable, then the LCAC Navigator Selection System can reduce the attrition rate from around 40% to around 17%. These estimates are based on rather small numbers and therefore are subject to substantial revision as more data accrue. The prediction of a success or failure for any given candidate by the LCAC Navigator Selection System is built on the foundation of statistical decision theory (SDT). The selection system is trying to make a decision about a candidate whose composite score on the test battery is known but whose training outcome is unknown. There is then, by definition, some uncertainty about the training outcome for this candidate. Uncertainty and the problem of making decisions in an uncertain world is the province of probability theory and SDT. The only mathematically self-consistent way that has been found to treat uncertainty is through probability theory. From the basic axioms of probability theory, one is able to construct a generally accepted framework dealing with uncertainty called the Bayesian approach. One of the core concepts within the Bayesian approach is the predictive probability density function. This paper presents some numerical examples of how a Bayesian predictive density can be calculated for the LCAC Navigator Selection System.

Acknowledgments

I would like to acknowledge the congenial work environment provided by LT Henry P. Williams, Head of the Selection Division at NAMRL at the time when this work was completed. In addition, I would like to extend a warm note of gratitude to my colleague Ms. Amanda Albert for her help with the LCAC Navigator data base. As always, the support provided by Dr. Ken Ford of the Institute for Human and Machine Cognition at the University of West Florida is gratefully acknowledged.

INTRODUCTION

This report updates the status of the Landing Craft Air Cushion Vehicle (LCAC) Navigator Selection System prediction algorithm. The last revision (Biggerstaff *et. al* [1]) took place at the end of 1997. With the receipt of new training data, the prediction algorithm was changed in June 2000 to take advantage of these new data.

Selection systems may be roughly divided into two parts. The first part concerns the choice of tests to be used in the test battery. The tests chosen are the ones that tap into some of the fundamental cognitive and psychomotor skills necessary for success in training. These tests are validated by administering the test battery to an initial set of subjects. All subjects tested during the validation stage enter training and their training outcome is recorded. On the basis of the validation stage, a prediction algorithm is developed that will be used in the operational stage. Here candidates are actually selected to enter training or are rejected based on the scores they achieve on the test battery.

The second part of a selection system comes into play with the dynamic updating of the original prediction algorithm as the training results come in for those candidates who were selected to enter training. The candidates who were rejected by the selection system obviously cannot provide any training data, but by a set of reasonable assumptions (see Blower [2]) they can be allocated to a guessed pass or fail during training. A policy decision was made by the LCAC training community to admit into training all candidates no matter what their score on the test battery. This policy worked to our benefit as it was now not necessary to perform as many allocations of rejected candidates into either actual training attritions or graduations. With these new data, the prediction algorithm can be revised under the assumption that it will do a better job than the previous algorithm.

The first part of this paper concerns itself with the data used to update the prediction algorithm. Composite scores are listed for all 46 LCAC Navigator candidates who have taken the test battery and completed training. Next, we show how these composite scores are calculated. Numerical examples are given for a typical candidate who is a predicted pass and for a typical candidate who is a predicted fail as based on their test battery scores. The new weighting coefficients for each of the eight tasks making up the test battery are documented. At the end of the first part of the paper, 2×2 classification matrices are used to show the performance of the algorithm in terms of the number of correct and incorrect predictions. We illustrate how changing the cut-off score in order to meet different policy needs impacts the number of correct and incorrect predictions.

The second part of the paper is more theoretical in nature and lays out the statistical justification for the LCAC Navigator Selection System. It is important to emphasize that the system makes *optimal* decisions given a fixed set of test battery scores. The only way that it (or any other competing system) could make better decisions is if it had better information from different tests. We can make this claim because the LCAC Navigator Selection System is based on the theoretical foundations of Statistical Decision Theory.

It turns out that the prediction algorithm that is ultimately derived from Statistical Decision Theory depends on probability distributions. Such probability distributions capture the uncertainty surrounding data and parameters in a statistical model. A large and ever-growing majority of experts in the statistical community has determined that the best way to deal with this kind of uncertainty is through probability theory. The corpus of knowledge the statistical community has accumulated over the years to process probabilities is called the Bayesian approach. We employ a rigorous Bayesian approach in developing the prediction algorithm for the LCAC Navigator Selection System.

One of the core components of the prediction algorithm is the ratio of the likelihood of obtaining a test battery score (the composite score, for example) given that the candidate really comes from the PASS group to the likelihood of obtaining the same composite score given that the candidate really comes from the FAIL group. From the Bayesian (*i.e.*, correct) perspective, this likelihood ratio is a ratio of *predictive probability densities*. We present in some detail how the Bayesian predictive density is derived for the LCAC Navigator Selection System.

A computer program was written to numerically evaluate the integrals involved in the two predictive densities. Numerical examples are given showing how this program works to produce these densities needed to form the likelihood ratio. Finally, we employ the program to compute that particular composite score that serves as a cut-off

score for the newly revised prediction algorithm. This cut-off score determines whether a candidate is a predicted fail or a predicted pass.

DATA FOR UPDATING THE PREDICTION ALGORITHM

Listing of Composite Scores

The composite score for each LCAC Navigator candidate is formed by a linear combination of the eight tests comprising the selection test battery. The weights are derived from a linear discriminant analysis of two groups. The two groups are the known successes in the initial LCAC Navigator training curriculum (labeled as the PASS group), and the known attritions during training (labeled as the FAIL group). At the time of this report, there were $N_{Pass} = 28$ in the PASS group and $N_{Fail} = 18$ in the FAIL group. Table 1 shows the composite score for each member of the PASS and FAIL group. The scores are listed in ascending order. At the bottom of the table is the sample mean and the sample standard deviation of the composite scores for each of the two groups. The sample mean of the PASS group is +0.599 with a standard deviation of .912, while the sample mean of the FAIL group is -0.932 with a standard deviation of 1.125. The discriminant analysis creates these composite scores such that the means are as far apart as possible, both groups have theoretical standard deviations of 1.00, and the composite scores for each group are distributed according to a Normal probability distribution.

Computation of the Composite Scores

The composite scores, as mentioned in the previous section, are computed as a linear sum of eight test scores multiplied by the corresponding eight weighting coefficients. These weighting coefficients are the unstandardized canonical discriminant function coefficients as reported by the SPSS Discriminant Analysis module. A constant is then added to this sum for the final composite score. The composite score is therefore computed according to Equation (1)

$$\text{Composite score} = \sum_{i=1}^8 c_i S_i + \text{Constant.} \quad (1)$$

The coefficients c_1 to c_8 , together with the tests they weight, are listed in Table 2. The first column gives the subscript attached to the coefficient, c_i , and the score, S_i , on the i th test. The better the performance on the Absolute Difference, Dichotic Listening, Stick with DLT, and Time Estimation tasks, the better the composite score. The better the performance on the Manikin test, the *worse* the composite score. The percentage correct on the CVT test does not appear at this point to impact the composite score too greatly. Quicker RTs when *answering* the CVT questions results in a higher composite score, while slower RTs to *reading* the CVT questions also results in a higher composite score. Thus, it appears that success in navigator training is correlated with taking time to understand a problem and then responding quickly once the problem is understood. The Absolute Difference task and the Time Estimation task remain the two most influential indicators of success in training.

Some speculation on why the Manikin test and the percentage correct on the CVT do not show better weighting coefficients is that these tests are more prone to improvement as information is passed from candidate to candidate about the content of the test. The more you know about the nature of these tests, the more you can prepare yourself beforehand. There is little one can do to prepare for the other tests even if one knows the general nature of the test. We will, therefore, consider the effects of removing the Manikin and percentage correct CVT scores from the algorithm in the future.

Tables 3 and 4 present the calculation of the composite score for two candidates who have not yet entered training. The first candidate, shown in Table 3, is a predicted fail because his composite score is -1.187, which is below the cut-off score of -.218. The second candidate, shown in Table 4, is a predicted pass because his composite score is +0.840 which is above the cut-off score. The determination of the cut-off score is deferred to a later section of the paper.

Table 1: A listing of the composite score for each member of the PASS and FAIL groups sorted in ascending order.

Number	PASS	FAIL
1	-0.956	-3.800
2	-0.844	-2.419
3	-0.651	-2.054
4	-0.387	-1.792
5	-0.206	-1.425
6	-0.139	-1.377
7	-0.062	-1.329
8	-0.031	-1.059
9	-0.016	-0.873
10	+0.041	-0.838
11	+0.279	-0.500
12	+0.498	-0.338
13	+0.555	-0.219
14	+0.761	-0.189
15	+0.828	+0.064
16	+0.856	+0.262
17	+0.861	+0.532
18	+0.898	+0.585
19	+0.938	
20	+1.001	
21	+1.016	
22	+1.042	
23	+1.062	
24	+1.305	
25	+1.308	
26	+1.590	
27	+1.785	
28	+3.437	
Sample mean	+0.599	-0.932
Sample standard deviation	+0.912	+1.125

Table 2: The coefficients, c_i , of the composite score for the newly revised prediction algorithm as determined by discriminant analysis.

Subscript	Test	Coefficient
1	Absolute Difference	+ .514
2	Percentage Correct on CVT	− .069
3	RT to CVT Questions	− .181
4	RT to CVT Answers	+ .322
5	Dichotic Listening Task	+ .187
6	Manikin	− .267
7	Multi-Task Stick with DLT	+ .152
8	Time Estimation Task	+ .561
9	Constant	− .401

Table 3: The calculation of the composite score for an LCAC navigator candidate who is a predicted FAIL.

Subscript	Test	Score	Coefficient	Multiplication
1	Absolute Difference	−0.2846	+0.514	−0.146
2	Percentage Correct on CVT	0.4250	−0.069	−0.029
3	RT to CVT Questions	1.0006	−0.181	−0.181
4	RT to CVT Answers	0.5729	+0.322	+0.184
5	Dichotic Listening Task	0.5712	+0.187	+0.107
6	Manikin	2.4700	−0.267	−0.659
7	Multi-Task Stick with DLT	0.1992	+0.152	+0.030
8	Time Estimation Task	−0.1638	+0.561	−0.092
9	Constant	1.0000	−0.401	−0.401
$\sum c_i S_i + \text{Constant} =$				−1.187

Table 4: The calculation of the composite score for an LCAC navigator candidate who is a predicted PASS.

Subscript	Test	Score	Coefficient	Multiplication
1	Absolute Difference	0.6999	+0.514	+0.360
2	Percentage Correct on CVT	1.0500	−0.069	−0.072
3	RT to CVT Questions	0.4927	−0.181	−0.089
4	RT to CVT Answers	−0.0741	+0.322	−0.024
5	Dichotic Listening Task	0.5712	+0.187	+0.107
6	Manikin	−0.7100	−0.267	+0.190
7	Multi-Task Stick with DLT	0.3087	+0.152	+0.047
8	Time Estimation Task	1.2896	+0.561	+0.723
9	Constant	1.0000	−0.401	−0.401
$\sum c_i S_i + \text{Constant} =$				+0.840

Classification of Correct and Incorrect Predictions

Table 5 presents a 2×2 table showing the breakdown for the two correct and the two incorrect predictions for a given cut-off score. A technically rigorous derivation of the cut-off score will be presented in a later section, but

Table 5: *The 2×2 classification table of the correct and incorrect predictions for 46 candidates using the revised algorithm with a cut-off score of $-.218$.*

		Prediction		
		PASS	FAIL	
Actual	PASS	24	4	28
	FAIL	5	13	18
		29	17	46

for now we can simply note from Table 1 that setting the cut-off score at $-.218$ results in a good trade-off between the two kinds of errors. Anyone obtaining a cut-off score below $-.218$ is a predicted failure and conversely, anyone scoring above $-.218$ is a predicted pass. This particular cut-off score yields nine errors. Four errors occur when we predict fail and the candidate actually passes, and the other five errors occur when we predict a pass and the candidate actually fails training. These numbers can be verified by looking back at Table 1.

Some rough estimates of the new attrition rate and the rejection rate due to the selection system may be made on the basis of the data in Table 5. The attrition rate is estimated as $5/29 = 17.24\%$, and the rejection rate as $17/46 = 36.96\%$. The failure rate during training appears to be about $18/46 = 39.13\%$. If the rejection rate of about 37% is acceptable, then the LCAC Navigator Selection System can reduce the attrition rate from around 40% to around 17%. Because these estimates are based on rather small numbers, they are subject to substantial revision as more data accrue.

Furthermore, raising or lowering the cut-off score will implement different policy decisions. To take two extreme examples, suppose that it is desired to either maximize the number of correct predicted passes, or to maximize the number of correct predicted failures. For the first case of maximizing the number of correct predicted passes, a cut-off score placed at a composite score of -1.00 will result in the classification table shown in Table 6. The trade-off for getting all the predicted passes right is that although the attrition rate in training deteriorates from 17.24% to $10/38 = 26.32\%$, the rejection rate decreases from 36.96% to $8/46 = 17.39\%$.

Table 6: *Maximizing the number of correct predicted passes by setting the cut-off score at -1.00 .*

		Prediction		
		PASS	FAIL	
Actual	PASS	28	0	28
	FAIL	10	8	18
		38	8	46

For the second case of maximizing the number of correct predicted failures, a cut-off score placed at the composite score of $+0.586$ will result in the classification table shown in Table 7. The trade-off for getting all the predicted failures right and bringing the attrition rate down to 0% is that the rejection rate deteriorates from $17/46 = 36.96\%$ to $31/46 = 67.39\%$. These two examples illustrate why, in most cases, a cut-off score in the middle of these two extreme cut-off scores, as illustrated in Table 5, engages a reasonable trade-off between lowering the attrition rate and raising the rejection rate.

Table 7: Maximizing the number of correct predicted failures by setting the cut-off score at +0.586.

		Prediction		
		PASS	FAIL	
Actual	PASS	15	13	28
	FAIL	0	18	18
		15	31	46

STATISTICAL DECISION THEORY

The prediction of a success or failure for any given candidate by the LCAC Navigator Selection System is built on the foundation of statistical decision theory (SDT). SDT is thoroughly explained in the context of personnel selection in Blower [3]. After the theory has been developed, two main components are revealed in the decision process. These are 1) the likelihood ratio of the composite score for a new candidate, and 2) the costs of a wrong decision and the probability of success or failure before the selection system was implemented. These elements are labeled by $\mathcal{L}(x)$ for the likelihood ratio and by β for the costs and prior probabilities. The likelihood ratio reflects the extent to which the test battery can discriminate between successes and failures on the basis of test scores, while β represents a variable cut-off score whose placement is determined by those costs and prior probabilities.

The prediction algorithm is based on these two components as well as the fundamental principle of SDT.

Make the prediction that results in the minimum average loss.

The prediction algorithm as based on these precepts is quite simple. It says

If $\mathcal{L}(x) \geq \beta$

Then PREDICT PASS

Else PREDICT FAIL

The likelihood ratio is the ratio of the probability density function for the composite score obtained by the candidate on the test battery under the two possible training outcomes, pass or fail. In symbols, this likelihood ratio is

$$\mathcal{L}(x) = \frac{P(D_{N+1}|\text{Pass})}{P(D_{N+1}|\text{Fail})}. \quad (2)$$

The subscript $N + 1$ on the data symbol is used to indicate that there exists N cases where test data and training outcome are known. These are the subjects who provided the validation sample. The $N + 1$ st candidate provides test data, but his or her training outcome is unknown. Therefore, the subscript $N + 1$ refers to all new candidates that the test battery is classifying as a predicted Pass or Fail. Pass or Fail appears to the right of the conditioned upon symbol in Equation (2) to indicate that one of these is assumed true. Thus, the composite score obtained by this candidate could come from one of two different probability density functions.

The ratio of these two numbers is compared to β . Imagine that these two numbers are the y -ordinates from two Normal probability density functions (pdf). For an efficient selection system, these two probability density functions are separated by a large margin. If the means of these two distributions are far apart (and they have the

same standard deviations), wrong decisions are kept to a minimum. We will see shortly how these two distributions arise.

β determines where the response threshold, or cut-off score, is placed. It is placed according to SDT to minimize the average loss. As already mentioned, β is a function of the costs assigned to the wrong decisions and the prior probabilities. A prior probability is meant to refer to the probability of passing or failing before the selection system was implemented. C_1 is the cost of the first wrong decision to predict fail when the candidate actually would have passed training. C_2 is the cost of the second wrong decision to predict pass when the candidate actually fails training.

In symbols, then, we can write out the definition of β as

$$\beta = \frac{C_2}{C_1} \times \frac{P(\text{Fail})}{P(\text{Pass})}. \quad (3)$$

Because C_2 is usually more expensive than C_1 , the first term in Equation (3) is greater than 1. $P(\text{Fail})$ is usually less than $P(\text{Pass})$ even without a selection system so that the second term in Equation (3) is less than 1. It might then turn out that these two terms conspire to produce a β equal to 1. For example, if it were true in the LCAC training community that $C_2 = \$150,000$, $C_1 = \$50,000$, $P(\text{Pass}) = .75$, and $P(\text{Fail}) = .25$, then

$$\begin{aligned} \beta &= \frac{\$150,000}{\$50,000} \times \frac{.25}{.75} \\ &= 1. \end{aligned}$$

In this case, the cut-off score is placed where $\mathcal{L}(x)$ happens to equal 1, that is, where the two pdf curves intersect. If a candidate obtains a composite score such that $\mathcal{L}(x)$ is less than 1, then he/she is a predicted fail and will be rejected by the selection system. On the other hand, if the composite score is high enough such that $\mathcal{L}(x)$ is greater than 1, then the candidate is a predicted pass and enters training. It is important at this point not to confuse the value $\beta = 1$ with the value of the composite score at this β . The composite score for $\beta = 1$ is not +1.00. We will later calculate the value of the composite score for $\beta = 1$.

Theoretically then, β is determined by the costs given to C_1 and C_2 and the known attrition rates before the selection system was implemented. In practice, however, it can be difficult to obtain estimates for these costs. This is also true for other training communities where selection systems are used. It would be extremely useful to us if LCAC policy makers and administrators would provide estimates of C_1 and C_2 . (It may be difficult because no one has tried to gather together the policy makers to assign costs after SDT has been explained to them.)

What seems to be done as an alternative to the theoretical calculation of β is to place the cut-off score at various locations and then gauge the reaction to the resulting frequencies of the wrong decisions. Eventually, the cut-off score will be maneuvered into a position where the trade-offs seem palatable. That is, the number of candidates wrongly rejected is balanced by the number of failing candidates correctly filtered from training. Of course, this method of setting the cut-off score is completely equivalent to the theoretical method. After having located an acceptable threshold by the pragmatic means just described, one can go back and fill in the values for C_1 , C_2 , $P(\text{Fail})$, and $P(\text{Pass})$ that would have given the same cut-off score.

Perhaps the best method of all is some combination of the theoretical and pragmatic approaches. It is relatively easy to construct a "what-if" spreadsheet-like program that would allow policy makers to shift back and forth between the two views (or show them simultaneously).¹ Interacting with such a program should make the task of establishing a cut-off score a more rational undertaking.

¹The author has written a couple of prototype programs in Visual Basic that do this.

THE PREDICTIVE DENSITY FROM BAYESIAN THEORY

After substituting for the values of $\mathcal{L}(x)$ and β , the prediction algorithm now looks like the following:

If $\frac{P(D_{N+1} | \text{Pass})}{P(D_{N+1} | \text{Fail})} \geq \frac{C_2}{C_1} \times \frac{P(\text{Fail})}{P(\text{Pass})}$

Then PREDICT PASS

Else PREDICT FAIL

The right-hand side of the algorithm does not present any mathematical difficulties. The only difficulties are ones of policy. The left-hand side, however, presents some interesting statistical issues.

The selection system is trying to make a decision about a candidate whose composite score on the test battery is known but whose training outcome is unknown. As explained earlier, this is what is meant by D_{N+1} . There is then, by definition, some uncertainty about the training outcome for this candidate. Uncertainty and the problem of making decisions in an uncertain world is the province of probability theory and SDT. The only mathematically self-consistent way that has been found to treat uncertainty is through probability theory. From the basic axioms of probability theory, one is able to construct a generally accepted framework for uncertainty called the Bayesian approach. One of the core concepts within the Bayesian approach is the predictive probability density function. The ratio on the left-hand side of the prediction algorithm is exactly the ratio of two predictive densities.

The Bayesian formula for the predictive density of the composite score conditioned on a pass in training is given by

$$P(D_{N+1} | D_N, \text{Pass}) = \int_R P(D_{N+1} | \theta, D_N, \text{Pass}) P(\theta | D_N, \text{Pass}) d\theta. \quad (4)$$

An entirely analogous equation can be written for the predictive density of the composite score conditioned on failing in training. θ stands for a set of parameters governing the likelihood of the data. In our case, the set of parameters consists of μ and σ , the mean and standard deviation, from the Normal distribution. When the discriminant analysis forms the composite scores, they are forced to go into a Normal distribution.

Within the formalism of the Bayesian approach, that is, by keeping strictly to the formal symbol manipulations of probability theory and not worrying about how actual numbers are arrived at, Equation (4) is not too difficult to derive. It is merely 1) the marginalization of the joint distribution of the new composite score and the set of parameters

$$P(D_{N+1} | D_N, \text{Pass}) = \int_R P(D_{N+1}, \theta | D_N, \text{Pass}) d\theta \quad (5)$$

and 2) the use of the product rule on the term inside the integral in Equation (5)

$$P(D_{N+1}, \theta | D_N, \text{Pass}) = P(D_{N+1} | \theta, D_N, \text{Pass}) P(\theta | D_N, \text{Pass}). \quad (6)$$

Because subjects provide independent data, Equation (4) can be shortened to

$$P(D_{N+1} | \text{Pass}) = \int_R P(D_{N+1} | \theta, \text{Pass}) P(\theta | D_N, \text{Pass}) d\theta. \quad (7)$$

The integral in Equation (7) consists of two terms. The first term is the likelihood of the new data, the composite score, for the candidate who is to be classified. The second term is the *posterior probability* of the set

of parameters as conditioned on all the available data from the N subjects who do have both composite scores and training outcomes. These two terms are multiplied for every value that θ can assume and then summed (integrated in the general case where the parameter can take on continuous values). One could regard Equation (7) as just an average. It is the average of all the likelihoods for the composite score as weighted by the posterior distribution of the parameters that determine the likelihood.

The specifics of Equation (7) for the particular set of circumstances that present themselves for the LCAC Navigator Selection System need explanation. As mentioned above, the likelihood for the composite scores as they are derived from the discriminant analysis follows the Normal distribution. Therefore,

$$P(D_{N+1} = x | \mu, \sigma, \text{Pass}) = \frac{1}{\sqrt{2\pi}\sigma_{Pass}} \exp \left\{ -\frac{1}{2} \left(\frac{x - \mu_{Pass}}{\sigma_{Pass}} \right)^2 \right\}. \quad (8)$$

Notice that the parameters of the Normal distribution have been substituted for θ and that x stands for the composite score obtained by the candidate after completing the test battery. Equation (8) has been conditioned on the assumption that the candidate came from the Pass group. An analogous equation can be written for the other situation where it is assumed that the candidate belongs to the Fail group.

The second term is the posterior distribution of μ and σ . This is available from Box and Tiao [4] as

$$P(\mu, \sigma | D_N, \text{Pass}) = k \sigma_{Pass}^{-(N+1)} \exp \left\{ -\frac{1}{2\sigma_{Pass}^2} (\nu s^2 + N_{Pass}(\mu_{Pass} - \bar{y}_{Pass})^2) \right\} \quad (9)$$

where k is a constant, $\nu = N_{Pass} - 1$, s^2 is the sample sum of squares, and $\bar{y}$ is the sample mean of the discriminant scores. Again, the posterior probability distribution conditioned on FAIL is entirely analogous.

A computer program was written to carry out a numerical integration of Equation (7) using the results of Equations (8) and (9). The region R for the integration was taken to be -4 to $+4$ for μ and $.7$ to 1.3 for σ . The computer iterated through the double loop for μ and σ in steps of $.01$.

NUMERICAL EXAMPLE

Suppose that a LCAC Navigator candidate obtains a composite score of $+0.60$ on the test battery. This is a score near the mean of the PASS group. Will this candidate be a predicted pass or predicted fail? To answer this question, first calculate

$$\mathcal{L}(x) = \frac{P(D_{N+1} = 0.60 | \text{Pass})}{P(D_{N+1} = 0.60 | \text{Fail})}$$

using the numerical approximation to Equation (7). The computer program provide the answers

$$P(D_{N+1} = 0.60 | \text{Pass}) = .4257$$

and

$$P(D_{N+1} = 0.60 | \text{Fail}) = .1354.$$

(Remember that these numbers are probability densities, the value of the y -ordinate for the predictive density function, and not probabilities.) Therefore,

$$\begin{aligned} \mathcal{L}(x) &= \frac{.4257}{.1354} \\ &= 3.14. \end{aligned}$$

To determine whether this candidate is a predicted pass or predicted fail, we need to know the cut-off score, or β .

If, for the sake of this numerical example, β can be calculated as in the previous section, then $\beta = 1$. We now apply the prediction algorithm.

If $\mathcal{L}(x) \geq \beta$

Then PREDICT PASS

Else PREDICT FAIL

Because $\mathcal{L}(x) = 3.14 \geq 1$ is true, the selection system predicts a pass for this candidate and would recommend entry into training.

THE NUMERICAL APPROXIMATION FOR THE PREDICTIVE DENSITIES

Having seen the prediction algorithm working at the global level, we now delve into a little more detail about the computation of the predictive densities

$$P(D_{N+1} = 0.60 | \text{Pass}) = .4257$$

and

$$P(D_{N+1} = 0.60 | \text{Fail}) = .1354.$$

Table 8 illustrates the workings of Equation (7) to arrive at the predictive density function. Table 8 lends some

Table 8: *Some values from the computer program to numerically approximate the predictive density function at a composite score of +0.60 assuming the score comes from the PASS group.*

μ	σ	Likelihood	Posterior density
-4.00	0.7	2.39×10^{-10}	1.25×10^{-268}
-1.00	1.0	.1109	3.78×10^{-21}
+0.60	1.0	.3989	1.32×10^{-5}
+0.60	0.7	.5699	3.42×10^{-6}
+0.60	1.3	.3069	6.43×10^{-7}
+0.60	0.9	.4433	2.01×10^{-5}

insight into how the value for the predictive density function conditioned on PASS,

$$P(D_{N+1} = 0.60 | \text{Pass}) = .4257$$

was arrived at. The first two columns give the values of the two parameters of the Normal distribution, μ and σ . Within the program, μ varies from $\mu = -4.0$ to $\mu = +4.0$ in increments of .01; σ varies between the limits of .7 and 1.3, also in increments of .01. The third column is the likelihood of $D_{N+1} = +0.60$ given whatever parameter values are specified in the first two columns. For the first row, then, column 3 contains the value of the likelihood,

$$P(D_{N+1} = +0.60 | \mu = -4.0, \sigma = .7, \text{Pass}).$$

The final column lists the posterior density function for the specified value of the parameters in the first two columns. For the first row, column 4 contains the value of

$$P(\mu = -4.0, \sigma = .7 | D_N, \text{Pass})$$

and is computed from Equation (9).

The likelihood value for a composite score of +0.60 if μ were to equal -4.0 and σ were to equal .7 is an extremely low value of 2.39×10^{-10} . This is simply the same calculation that one would make in computing a z -score. If the true mean is at -4.0 with a standard deviation of .7, then a score of +0.60 is over 6 standard deviations to the right of the mean. The y -ordinate (pdf value) of the Normal curve at 6 standard deviations from the mean must be very small.

The subsequent rows in Table 8 offer an illustration of the change in the likelihood and the posterior density as μ and σ are altered. Row 2 shows what happens when μ is eventually changed within the computer program to -1.0 and σ , as well, is eventually iterated to a value of 1. Now, intuitively we know that a composite score of +0.60 must be closer to a true mean of -1.0 with a standard deviation of 1 as compared to the conditions in row 1. It is now less than 2 standard deviations away from the mean. Therefore, the y -ordinate (pdf value) of the Normal curve must be larger at this value of μ and σ . The value of .1109 confirms this increase.

In row 3, μ is eventually changed to $\mu = +0.60$ and σ will recycle to 1 at some point at this setting of μ . If the composite score exactly matches the mean, then the y -ordinate must take on its maximum value for a given σ . The mean is the highest point of the Normal curve. For $\sigma = 1$, this pdf value is easy to calculate:

$$\begin{aligned} \text{pdf}(x = +0.60) &= \frac{1}{\sqrt{2\pi}\sigma} \exp \left\{ -\frac{1}{2} \left(\frac{x - \mu}{\sigma} \right)^2 \right\} \\ &= \frac{1}{\sqrt{2\pi}} \exp \left\{ -\frac{1}{2} (.60 - .60)^2 \right\} \\ &= \frac{1}{\sqrt{2\pi}} \exp(0) \\ &= \frac{1}{\sqrt{2\pi}} \\ &= .3989. \end{aligned}$$

This calculation verifies the output of the computer program for the likelihood given in row 3.

Row 4 shows that it is possible to obtain an even higher value for the likelihood. If σ is smaller than 1, ($\sigma = .7$ in Row 4), then the likelihood increases to a value of .5969. Two more typically high values for the likelihood, and the parameter values associated with these likelihoods, are shown in the fifth and sixth rows. In Table 8, only six computations are shown. In the computer program, $800 \times 60 = 48,000$ separate likelihood and posterior density computations were performed.

Given that we have all these likelihoods, what do we do with them? According to the predictive density formula of Equation (7), we have an integration involving these likelihoods. An integration is the same as calculating an average with respect to a probability distribution. In the predictive density formula, that probability distribution is the posterior distribution

$$P(\mu, \sigma | D_N, \text{Pass}).$$

That is why we have also calculated the posterior pdf for every value of μ and σ according to Equation (9). These are the values shown in the last column of the table. These numbers indicate how much to weight the likelihood in calculating the average. Parameters values like $\mu = -4.0$ and $\sigma = .7$ have an astronomically low posterior pdf (2.01×10^{-268}) because they are extremely unlikely to have arisen from the PASS group as based on the past data. Therefore, these parameter values lend an almost infinitesimal weight to their particular likelihoods. In essence, likelihoods like that in row 1 are excluded from the average.

The Bayesian approach always includes all the information we have available. The posterior density function is

where the prior data containing both composite scores and training outcome are included into the prediction for the composite score of a new candidate. We know how likely, or unlikely, various values of μ and σ are based on the past data, D_N . Equation (9) dictates that the posterior density depend upon both the distance of the sample mean, $\bar{y}$, from μ and the sample standard deviation, s , from σ . Because

$$\bar{y}_{Pass} = .5989 \text{ and } s_{Pass} = .9123,$$

values of the posterior density at parameter values like $\mu = -4.0$ and $\sigma = .7$ must be very small.

The final column of Table 8 shows that the posterior density becomes progressively larger for the values of μ and σ supported by the past data. The parameter values like $\mu = -1.0$ and $\sigma = 1$ become more likely given the past data, and parameter values like $\mu = +0.60$ and $\sigma = .9$ become the most likely of all because they match the past data very well. These parameter values weight the average of the likelihood so that likelihood values such as 2.39×10^{-10} are not counted at all, but likelihood values such .3989 are weighted heavily. When all 48,000 weights multiply all 48,000 likelihoods, they are summed and then divided by the sum of the weights. The resulting average value of the likelihoods is .4257 as given above.

Table 9 illustrates exactly the same points for $P(D_{N+1} = +0.60|\text{Fail})$. The first row lists the same parameter

Table 9: Some values from the computer program to numerically approximate the predictive density function at a composite score of +0.60 assuming the score comes from the FAIL group.

μ	σ	Likelihood	Posterior density
-4.00	0.7	2.39×10^{-10}	2.01×10^{-82}
-1.00	1.0	.1109	2.04×10^{-5}
-0.90	1.1	.1413	2.23×10^{-5}
-0.95	1.2	.1444	1.78×10^{-5}
-0.96	0.9	.0987	1.25×10^{-5}
-0.93	1.1	.1379	2.25×10^{-5}

values as row 1 of Table 8 to illustrate that the likelihood value listed in Table 8 remains the same for the same parameters. It must remain the same because the composite score for the candidate does not change, and the likelihood only depends upon the composite score's relationship to the given μ and σ . What does change is the posterior distribution of μ and σ given the different sample data from the FAIL group. For the FAIL group, the sample mean and sample standard deviation are

$$\bar{y}_{Fail} = -.9316 \text{ and } s_{Fail} = 1.1252.$$

Therefore, the weights for the likelihoods over all the parameter values will change and result in a different average likelihood. The next five rows of Table 9 concentrate on the more probable parameter values for the FAIL group. This small sample of highly probable likelihoods is seen to be in the range of about .10 to .15 with the actual average as computed by the program of

$$P(D_{N+1} = +0.60|\text{Fail}) = .1354.$$

CALCULATING THE CUT-OFF SCORE

What does $\beta = 1$ correspond to in terms of the composite score? In other words, what is the value of the cut-off score so that we can make predictions about candidates directly without having to compute $\mathcal{L}(x)$? What value of D_{N+1} will make $\mathcal{L}(x) = 1$? Or equivalently,

$$P(D_{N+1} = ?|\text{Pass}) = P(D_{N+1} = ?|\text{Fail}).$$

Through an iterative process, the computer program for predictive densities found that composite score where the densities conditioned on PASS and FAIL were exactly equal. If the composite score is equal to -0.218 , then

$$P(D_{N+1} = -0.218|\text{Pass}) = P(D_{N+1} = -0.218|\text{Fail})$$

and $\mathcal{L}(x) = 1$.

Therefore, the cut-off score can be placed at -0.218 for $\beta = 1$. This was, in fact, the cut-off score used to construct Table 5. We want to re-emphasize the point that this cut-off score is appropriate only for a specific set of costs and prior probabilities. If these should change, then the cut-off score would change as well.

For example, we arrived at $\beta = 1$ by arbitrarily letting $C_2/C_1 = 3$ and $P(\text{Fail})/P(\text{Pass}) = 1/3$. If, contrary to these assumptions, the ratio of prior probabilities is closer to

$$\begin{aligned} \frac{P(\text{Fail})}{P(\text{Pass})} &= \frac{.40}{.60} \\ &= \frac{2}{3} \end{aligned}$$

as we commented on earlier when discussing the data in Table 3, then this would argue for moving the cut-off score upwards to $\beta = 2$, assuming that the costs have remained invariant.

What has been the impact of allowing for uncertainty in the parameters? If we had just calculated a cut-off score based on the sample means of the PASS and FAIL distributions, then the cutoff-score would have been placed midway between these two means for $\beta = 1$ (Blower [3]).

$$\begin{aligned} \text{Cut-off score} &= \frac{.5989 + (-.9316)}{2} \\ &= -.1664. \end{aligned}$$

allowing for uncertainty in the parameters μ and σ has displaced the cut-off score slightly downwards to -0.218 . The predictive density curves are a little bit flatter than Normal curves erected over the sample means. They are flatter because they had to take account of many parameter values when constructing the likelihood for a new score, while the Normal curves just assumed a true likelihood at the sample values. Thus, where the predictive curves intersect will be at a composite score a little lower than for the Normal curves.

SUMMARY

A firm statistical foundation has been laid down for the LCAC Navigator Selection System. It is based on Statistical Decision Theory with the probabilities required by SDT provided by the Bayesian approach. The ratio of likelihoods within SDT turn out to be a ratio of Bayesian predictive densities. The predictive density was derived from the basic axioms of probability theory, the sum rule and the product rule. For the specific case treated here, this involved the integration of a normal density for composite scores times the posterior probability density for the two parameters of the normal density. A computer program was written to numerically integrate the predictive density conditioned on the Pass group and the predictive density conditioned on the Fail group. When the costs for the wrong decisions are specified and the prior probabilities for passing and failing are filled in, a decision to admit into training can be made for every new LCAC candidate. These decisions are optimal in the sense that they result in the lowest possible average cost.

REFERENCES

1. Biggerstaff, S., Blower, D. J., Portman, C. A., and Chapman, A. *Landing Craft Air Cushion (LCAC) Navigator Selection System: Initial Model Development*. NAMRL-1399, Naval Aerospace Medical Research Laboratory, Pensacola, FL, 1998.
2. Blower, D. J. *A Cost-Benefit Analysis of the Landing Craft Air Cushion (LCAC) Selection System*. NAMRL Special Report 00-2, Naval Aerospace Medical Research Laboratory, Pensacola, FL, 2000.
3. Blower, D. J. *The Statistical Foundations of the Pilot Prediction System*. NAMRL Technical Report in review, Naval Aerospace Medical Research Laboratory, Pensacola, FL, 2000.
4. Box, G.E.P. and Tiao, G.C. *Bayesian Inference in Statistical Analysis*. John Wiley & Sons, New York, NY, 1973.

Reviewed and approved 12/15/00



R. R. STANNY, Ph.D.
Technical Director



This research was sponsored by the Naval Sea Systems Command under grant number N0002401WR00588.

The views expressed in this article are those of the author and do not necessarily reflect the official policy or position of the Department of the Navy, Department of Defense, nor the U.S. Government.

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

REPORT DOCUMENTATION PAGE			Form Approved OMB No. 0704-0188	
Public reporting burden for this collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington, VA 22202-4302, and to the Office of Management and Budget, Paperwork Reduction Project (0704-0188), Washington, DC 20503.				
1. AGENCY USE ONLY (Leave blank)		2. REPORT DATE 15 DEC 2000		3. REPORT TYPE AND DATES COVERED
4. TITLE AND SUBTITLE An Update to the Landing Air Craft Cushion (LCAC) Navigator Selection System Prediction Algorithm			5. FUNDING NUMBERS Naval Sea Systems Command grant number N0002401WR00588	
6. AUTHOR(S) D. J. Blower			Code PMS 377 Work Unit 7212	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Naval Aerospace Medical Research Laboratory 51 Hovey Road Pensacola Fl 32508-1046			8. PERFORMING ORGANIZATION REPORT NUMBER NAMRL Special Report 00-3	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES) Naval Sea Systems Command PMS 377C 2531 Jefferson Davis Hwy Arlington, VA 22242-5160			10. SPONSORING/MONITORING AGENCY REPORT NUMBER	
11. SUPPLEMENTARY NOTES Approved for public release; distribution unlimited.				
12a. DISTRIBUTION / AVAILABILITY STATEMENT			12b. DISTRIBUTION CODE	
13. ABSTRACT (Maximum 200 words) This report updates the status of the Landing Craft Air Cushion (LCAC) Navigator Selection System prediction algorithm. The last revision took place at the end of 1997. With the receipt of new training data, the prediction algorithm was changed in June 2000 to take advantage of this new data. Some rough estimates of the new attrition rate and the rejection rate due to the selection system may be made on the basis of the new data. The attrition rate is estimated as 17.24% and the rejection rate as 36.96%. The failure rate during training appears to be about 39.13%. If the rejection rate of about 37% is acceptable, then the LCAC Navigator Selection System can reduce the attrition rate from around 40% to around 17%. Of course, these estimates are based on rather small numbers and therefore are subject to substantial revision as more data accrues. The prediction of a success or failure for any given candidate by the LCAC. Navigator Selection System is built on the foundation of statistical decision theory (SDT). The selection system is trying to make a decision about a candidate whose composite score on the test battery is known, but whose training outcome is unknown. There is then, by definition, some uncertainty about the training outcome for this candidate. Uncertainty and the problem of making decisions in an uncertain world is the province of probability theory and SDT. The only mathematically self-consistent way that has been found to treat uncertainty is through probability theory. From the basic axioms of probability theory one is able to construct an appealing approach to uncertainty called the Bayesian approach. One of the core concepts within the Bayesian approach is the predictive probability density function. This paper presents some numerical examples of how a Bayesian predictive density can be calculated for the LCAC Navigator Selection System.				
14. SUBJECT TERMS Personnel selection, statistical decision theory, Bayesian approach, predictive density functions			15. NUMBER OF PAGES 18	
			16. PRICE CODE	
17. SECURITY CLASSIFICATION OF REPORT UNCLASSIFIED	18. SECURITY CLASSIFICATION OF THIS PAGE UNCLASSIFIED	19. SECURITY CLASSIFICATION OF ABSTRACT UNCLASSIFIED	20. LIMITATION OF ABSTRACT SAR	