

PRIORITIZATION OF PROGRAMS

by

Dr. Cameron R. Peterson
Decisions and Designs, Incorporated

I'm going to restrict my discussion today to the problem of generating an order, a buy, both within a sponsor group and between sponsor groups.

How do you rank-order a set of programs in the most cost-effective manner? We use computer models that are different from the models that we've been discussing the last day or two. We have been discussing operations research models and simulation models, which are models of the environment. They are engineering models of a complicated environment that are designed to simulate that process in such a way that a decision-maker can have a better grasp, a better understanding of that environment. By contrast, I will describe psychological models based upon decision analysis, models designed to fit in and work with a conference.

For the psychological model, the knowledge is in the heads of people rather than in the computer. The computer is used to provide a framework for dialogue and debate among a group of participants. Therefore, the participants must be knowledgeable about the items to be prioritized and, ideally, they should be motivationally involved. Advocates argue and pound the table and fight with each other, and that's how this process works best.

First, some background. Over the last 10 years, we have gradually evolved a process that we call a "decision conference." It's a highly structured process designed to use computer modeling to aid people who have to make decisions. Typically, a general manager will bring 10 or 15 of his managers into the conference room, and they sit around the table to develop priorities. If he's going to build a factory, the priorities are about amount of robotics versus amount of computer-aided design versus size of building, etc. If it's a POM working group, there will be priorities among tanks, helicopters, radios, trucks, and so on. But the assumption of this modeling process is that the people have the knowledge, and the computer model is used to develop a framework so that knowledge can be exposed and the participants can communicate with each other in order to achieve an effective consensus on the priorities.

We use a procedure called triangulation, which means that we ask the same question different ways, with the aim of

finding inconsistencies. We then focus discussion on these inconsistent answers so that the participants can understand why the differences are what they are, and thereby clear up knowledge gaps on the way to effective prioritization.

Here is an example of triangulation. First, the participants intuitively come up with a decision, which might be a rank order of the programs under consideration. Second, we ask questions about costs and benefits, which are typed into a computerized cost-benefit model. The computer then generates an order of programs which is compared with the intuitive rank order. Where the two rankings differ, there has been a mistake, either in the intuitive ordering process or in the costs and benefits that led to the computer ordering. The search for the reasons for this inconsistency between the two procedures leads the participants to a better understanding and an ability to communicate the rationale behind their priorities. They then reconcile the inconsistency by changing either the intuitive rank order or the order that came from the cost-benefit model.

Three analysts typically assist the participants in a decision conference. One asks questions about the items to be prioritized: What are the costs, what are the benefits, what are the reasons for having one item higher than another one on the list. While this dialogue is taking place, a second analyst enters the costs and benefits into a computer model. Typically, the participants are on the edge of their chairs, waiting to compare the computer answers with their own answers. A third analyst captures the rationale. Sometimes a group can argue for a half-hour about why the priority should be higher for one package than for another one; then, 3 weeks later, they forget which one had the highest priority. They forget the arguments. So, while the group argues, thinks hard, and debates, this third analyst captures the reasons for why one program is rated higher than another program. These reasons may be intuitive, qualitative judgments made by a group, or they may be quantitative arguments derived from simulation models.

These psychological models require a scaling of the benefits of the programs. There are two important reasons for going beyond ordinal or ranking judgments and using scales. The first is to permit assessments of cost to be combined with assessments of benefit for a cost-benefit analysis. The second purpose of the benefit scale is to facilitate prioritization across sponsors.

This benefit scale has the property that if program A is worth two points and program B

is worth three points and program C is worth five points, then it must be true that the participants are indifferent between the benefit from program C and the combined benefit from the other two programs. People make serious mistakes when they attempt to generate such a scale. Psychological research demonstrates why it is difficult. A psychologist, George Miller, once wrote a paper called the "Magic Number 7 Plus or Minus 2" where he demonstrated over a wide variety of tasks, like judging weight, or brightness, or musical pitch, or loudness, that there is not a good internal scale against which a person can compare the thing being evaluated.

The following sort of judgment facilitates the generation of a scale. Consider this imaginary experiment. I have five rocks sitting on a table, and I ask you to judge the number of pounds of each rock. You lift one rock and say that weighs 6 pounds, another and say it weighs 5 pounds, a third and say that weighs 3 pounds, and so on. Next, you compare rocks 6, 5, and 3 with each other. You lift rocks 5 and 3 in one hand and rock 6 in the other. Assume that you judge that there is no difference between the two hands, that the two rocks weigh the same as the third. Those two judgments are inconsistent; so, you must reconcile them. You must either change your estimates of pounds or change your judgment that the heavier rock weighs the same as the sum of the two lighter rocks. There is no doubt about which you would do; you would change your estimates of pounds so that they agree with your indifference judgment.

People are much better as null instruments for making indifference judgments than they are at generating scale numbers. Therefore, it is useful to use indifference judgments for creating a benefit scale. The question may take the following form: is the overall mission of the organization enhanced more by Program A or by both Programs B and C? Such judgments are then used to modify the numerical weighting of relative benefit.

The next step is to use the benefit scale in a cost-benefit analysis. Even though these measures of benefit are highly subjective, it has been our experience that participants are more confident about the veridicality of the benefit scale than about the estimates of average annual life cycle cost used in the cost-benefit analysis. The conclusion is almost always that there is more uncertainty about the costs than there is about the benefits.

A cost-benefit order of buy is generated by selecting programs according to the ratio of cost to benefit. When people intuitively

select an order of buy, however, they tend to rank order primarily by benefit--they essentially ignore cost. The program with the highest benefit is selected first, then the program with the next highest benefit, and so on.

Because of this systematic bias, it is enlightening to contrast an intuitive order of buy (typically according to benefit) with an order of buy generated by cost-benefit ratios. Consider the following example. The benefit order is A-B-C-D-E and the cost-benefit order is B-D-E-C-A. Assume that the cost of A is equal to the combined cost C-D-E. Participants invariably prefer the benefit order. But now further assume that A, even though it has more benefit than any other single program, has less benefit than the combination of C-D-E. Substantial internal conflict is created when participants prefer the order A-B-C-D-E but prefer the package B-C-D-E to the package A-B, which costs the same amount. The benefit scale is useful in identifying this kind of internal conflict, and it forces participants to consider cost as well as benefit and, usually, leads to a substantial modification in their intuitively prioritized order of buy.

Benefit scales are also useful for prioritizing across sponsors, which is more difficult because the programs tend to serve more diverse functions. The procedure for scaling across sponsors is as follows. First, participants within each sponsor use the procedure described above to scale their programs in terms of benefit. Then a group of "honest brokers" is selected for going across sponsors. Their task is to find a program or a set of programs for Sponsor A that yields the same level of benefit as another program for Sponsor B, and as yet another program for Sponsor C, and so on. The objective is to develop a benefit scale for a subset of programs from each sponsor. Then the scales within and between sponsors can be combined to yield a common scale for all programs. This procedure requires additional effort within sponsors in exchange for a more effective prioritization across sponsors.

This procedure for prioritizing programs assumes that the programs are independent, i.e., that neither the cost nor the benefit of one program changes as a function of whether any other program is selected. Linear models, which assume no interaction, are simple and transparent as frameworks for discussion. Unfortunately, the convenient assumption of universal independence is incorrect. Many programs interact with other programs.

A large body of research, however, has shown that only a slight amount of inaccuracy results from ignoring mild-to-moderate interactions. A linear model typically accounts for over 90 percent of the variance of an interacting environment. Consequently, we have adopted the strategy of selectively modeling interactions in these prioritization models used in conferences. The procedure is to identify the relatively small percentage of interactions that are very important, model those interactions, and then treat the

relatively high percentage of remaining programs as independent.

Thus, this approach of applying computer models to conferences is designed to yield results that are precisely wrong but approximately right. It applies the Pareto principle of achieving 80 percent of the value with 30 percent of the resources to the process of analysis, as well as to the cost-benefit of the programs that are the subject of analysis.

