

Who's Eating Your Lunch?

A Practical Guide to Determining the Weak Points of Any System.

Kenneth E. Murphy; ARINC, Albuquerque

Charles M. Carter; ARINC, Albuquerque

Robert H. Gass; ARINC, Annapolis

Key Words: Ranking Methods, RAPTOR, Monte Carlo Simulation, V-22 Osprey

SUMMARY AND CONCLUSIONS

This paper describes a technique for focusing system-level improvement efforts on those areas that are contributing most to undesirable RAM performance. Monte Carlo simulation that incorporates localized performance reporting is used to identify the key culprits affecting a system's availability or reliability. This ranking technique accounts for issues such as distributional variability, logistical constraints, complex redundancy, internal dependencies, and dynamic components that are often overlooked using traditional mean life ranking methods. We applied this methodology to the V-22 Osprey tilt-rotor aircraft and produced some insightful results.

1.0 INTRODUCTION

The vast majority of systems in existence today will undergo some manner of modification to improve their effectiveness, their capacity to generate products or to minimize their maintenance costs. Often just a handful of components will disproportionately contribute to a system's poor operation. In some cases, component-level logistics issues have the most significant affect on a system's performance. A ranking method that finds these troublesome components and indicates potential improvements could greatly enhance the usefulness of any system.

Techniques for finding these "weak link" components have been around since the dawn of reliability science and are based on either an expert's "sense" or a mathematical ranking. In general, mathematical ranking techniques are preferred over expert opinion as systems become more complex and expensive. Unfortunately, traditional ranking techniques only deal with a component's mean life. Real life reliability concerns such as component dependency, repair rates, sparing and resource competition are rarely considered by ranking techniques because these issues are often too difficult to incorporate.

Furthermore, large complex systems that cannot be resolved into simple series or parallel subsystems cannot use the traditional ranking techniques because of enumeration

limitations. Components whose mean life varies during different phases of a mission profile are also not easily handled by traditional methods.

Simulation offers a better approach by considering all these issues simultaneously and producing a true ranking of components that are causing poor system performance, or as we say, "*eating your lunch*." This paper provides a practical example of our use of the Raptor (Ref. 1) simulation software tool to overcome the limits of traditional ranking techniques. We will demonstrate the use of this simulation ranking process through a case study based on actual analysis performed on the U.S. military's V-22 Osprey aircraft.

Nomenclature and Notations

A_o	Operational Availability
GUI	Graphical User Interface
k -out-of- n	k paths of n are needed to operate successfully
θ	Mean Life
μ	Mean Repair Time
RAM	Reliability, Availability and Maintainability
RBD	Reliability Block Diagram
$R(t)$	Reliability at time t

2.0 RANKING FACTORS DEFINED

2.1 Distributions

To accurately determine which components are the dominant causes of a system's unavailability or poor reliability, numerous factors must be investigated. The two most commonly considered are the mean life (θ) and the mean repair (μ) times for each component of a system. Unfortunately, the actual failure and repair times are often not recorded and this makes determining the underlying failure or repair distribution impossible. Any ranking method based on simple point estimates that does not account for distributional variability is irrelevant and does not deserve discussion. To overcome a lack of data, the statistical distributions representing a component's operating and repair cycles are often assumed to be exponential in nature. This can introduce

complications as delineated by the Murphy, Carter and Brown (Ref 2) paper on the misuse of the exponential distribution. For the purpose of this paper, we will assume that relevant failure and repair distributions for each component of a system are known.

2.2 Logistics

The next family of factors that is routinely overlooked in traditional ranking methods is the logistical constraints of a system. These factors generally separate into two distinct classes of delays that inhibit components from being fixed or replaced. We classify these as *material* and *temporal*. Examples of material delays include the lack of spare parts or the lack of repair personnel to conduct repairs. Examples of temporal delays include the time associated with: traveling to a repair site, performing checkout procedures, and documenting repairs. There are numerous other logistical factors, but for the purpose of this paper we will consider only those mentioned above.

2.3 Topology, Interactions and Dynamic Components

Factors that are rarely considered in traditional ranking methods are the complex topology of a system, the interactions between components, and how the RAM characteristics of components change over time.

Series and parallel topological (i.e., *k-out-of-n*) structures are frequently included in ranking methods, but systems that cannot be expressed using only these configuration types are seldom considered. Non-series and non-parallel structures are denoted as *complex* and include such structures as bridge networks, delta and star formations (Grosh Ref 3), and cascading or adjacency arrangements.

Interactions between components generally precipitate their effects when one or more components stop operating because the component they are dependent upon has stop functioning. Considering these types of dependency aspects within a ranking methodology is extremely difficult due to the conditional mathematics that are generated.

Dynamic components are defined as those that change their attributes as time progresses. In a simulation environment, this is known as *phasing* and is often used to mimic components that are subjected to varying stresses. Unless every component considered with a ranking method fails and repairs exponentially (a Markovian fantasy), considering phasing is difficult due to the interactions between adjacent and discontinuous phases. For example, a component that fails according to a Weibull distribution and repairs lognormally does not possess the exponential memory-less property, and the time already exhausted on a component in a particular phase must be considered in subsequent phases to properly rank that component.

The topology, interactions and dynamic component factors are generally not included in ranking methods because of one

underlying attribute. These factors all introduce complex conditional mathematics that in the past has led to the use of weaker traditional ranking methods. In the old days (i.e., last century), the use of simplifying mathematics was necessary and the inherent inaccuracies were acceptable. With the widespread availability of simulation tools today, ignoring the significant affects discussed in the previous paragraphs should be considered nothing less than poor quality analysis.

With the major ranking factors having been discussed, we will turn our attention towards two very different methods employed to produce a ranked list of weak link components.

3.0 METHODS

3.1 The Traditional Approach

The most frequently used method of ranking components of a system to determine which are the major players in unavailability or poor reliability is what we call the *mean life rank method*. The mean life of each component is sorted in ascending order and the first few listed are the ones that are considered the worst. Resources are then committed to improve these weak components through new technology initiatives, redundancy or logistical improvements. We have found that the magnitude of system-level improvements advertised for availability or reliability is rarely achieved, however, when this type of ranking method is implemented.

Occasionally, two enhancements are added to this traditional technique to advance this method from a poor to a mediocre methodology. The first is to incorporate the effects of simple redundancy or *k-out-of-n* structures. This improvement recognizes that weaker components with relatively smaller mean lives can be ganged together in a redundant fashion to improve their overall combined mean life. When this feature is added to the traditional method, components that were once considered weak when viewed as a single entity, are now viewed as stronger when consider as part of redundant subsystem.

A second improvement to the traditional method is the consideration of component repair rates. For example, if a component fails often but requires very little time to repair (i.e., loss of production is not significant) then ranking that component as a system weak point may not make sense. When the effects of simple redundancy and component repair rates are both integrated into the traditional ranking method, a more sophisticated assessment can be made of those components that are eating a system's lunch. This is especially true if availability is the key parameter of concern.

Traditional ranking methods, however, are grossly deficient in their ability to consider distributional variability, logistics, complex topology, interactions and the effects of dynamic components. When only the mean life of a component is considered, the spread of the underlying failure distribution is disregarded. The spread of the failure data has a significant

affect on system performance if reliability is the principal parameter of interest. A component with a low mean life but very tight tolerances (very small distributional spread) is usually not something that should be considered significant in a ranking assessment since it can be eliminated with the appropriate frequency of preventive maintenance. A high mean life component with a large spread between its upper and lower bounds could warrant more consideration. The traditional ranking method cannot distinguish properly between these two types of components.

Furthermore, by not considering logistical constraints, complex topology, component interactions and dynamic aspects of components, the traditional ranking methodology has to be considered at best an inadequate estimation of those components that are significantly contributing to a system's poor availability or reliability. On the bright side, the traditional ranking method is fast. Nonetheless, with the current availability of RAM simulators, there is no reason to continue using these archaic and inaccurate ranking methods.

3.2 Simulation: A More Accurate Approach

Simulation can produce an accurate and relevant ranking of components if the simulator can sufficiently model the factors we have previously discussed. A good simulator will allow all these factors to be characterized simultaneously as well as representing the effects of their interactions. Discrete event Monte Carlo simulations that have been simulated to the appropriate stopping conditions and for enough trials, are the best methodology for determining a ranking of a system's weak link components. Proper simulation theory is a vast subject (too large to discuss in this paper); we would recommend our readers review Murphy and Carter's (Ref. 4) paper on the appropriate length and number of trials a simulation requires to generate high-quality results.

We chose to use the simulator known as "Raptor" for our case study for two reasons. First, Raptor can model all the factors previously mentioned, and second, Raptor possesses an ergonomically designed graphical user interface (GUI) that allows even a novice user to produce results quickly. Raptor has a built-in function called "node analysis" that automatically generates a ranking of weak components with respect to both availability and reliability. Raptor defines a component's availability as a ratio of time the component is operational to the total time a simulation trial is conducted. Raptor defines a component's reliability as a ratio of the number of trials a component operated without failing to the total number of trials simulated.

The Raptor node analysis function is based on placing measuring devices known as "nodes" wherever local availability or reliability values are desired. In general, a node in Raptor is an element that controls a local subsystem. This local subsystem can be a single component, a collection of components, a set of redundant components, or any complex topology that can be envisioned. These nodes not only control their local subsystems, but also monitor and collect statistics

on their status. Raptor defines a node's status as being in an up or down condition. While a simulation is in progress, a component failure will affect the status of its node. If a node's *k-out-of-n* condition is not achieved, that node will change to a down status. A failed component will continue to negatively affect its node throughout any delays such as waiting for spares, waiting for resources or post repair delays, in addition to the actual repair effort.

Statistics are collected for all nodes and their availability or reliability results reflect the true impact failed components would have on a system. Once a simulation is completed, a ranking of the nodes' availability or reliability is provided. These simulation ranking results incorporate all the relevant factors and thus provide much greater insight than traditional rankings.

Raptor provides a weak link visualization feature that allows a user to designate three levels of availability or reliability tolerances. Those tolerances are displayed on nodes in a red, yellow and green color scheme. The red tolerance defines a node that is unacceptably poor with respect to its availability or reliability. The yellow and green tolerances define a debatable and acceptable level of availability or reliability, respectively.

With an easy to use simulator available that can model all the relevant ranking factors simultaneously, their interactions, and produce an accurate component ranking, it is a wonder that any other method would be used in today's time.

4.0 A CASE STUDY

4.1 The V-22 Osprey Aircraft

To demonstrate the versatility of the node or weak link analysis ranking method, we chose to relate an actual initial assessment effort that we completed for the U.S. military. We conducted this ranking analysis to support the Navy and Air Force's joint venture to procure the V-22 Osprey tilt-rotor aircraft (Figure 1).



Figure 1. V-22 Osprey

The V-22 aircraft incorporates advanced yet mature technologies proven in the XV-15 tilt-rotor demonstrators. A tilt-rotor platform combines the speed, range and fuel efficiency normally associated with turboprop aircraft with the vertical take-off, landing and hover capabilities of helicopters. This tilt-rotor aircraft represents a technological breakthrough in aviation that is designed to meet the future needs of both military and commercial industries.

Of interest to the V-22 program office was the weapon system reliability (i.e., aircraft reliability for a four-hour sortie) and identification of components that were most frequently causing aborted missions. Development of the V-22 RBD involved the participation of numerous organizations including HQ Air Force Special Operations Command in Florida, the System Program Office in Maryland, and cooperation from major developmental contractors such as Boeing and Bell Helicopter. This RBD was infused with additional nodes at each location where we desired a localized measurement of a component's mission reliability. Figure 2 represents the actual RBD we used to conduct the weak link analysis.

It was determined by analysis, based on simulation theory, that 2,000 trials were sufficient to provide adequate results in comparison to the errors introduced by the simulation input data. To mimic reality, analysts chose not to allow the components to be repaired during a four-hour simulated flight. Additionally, many of the logistical aspects of the components were disregarded for the same reason. The V-22 RBD did however possess complex topology, dependency and dynamic

components. Furthermore, as is the case with most aircraft simulations, *phasing* was added to the RBD. The V-22 four-hour sortie was delineated into eight distinct phases representing portions of its flight profile. The reason for this is that the requirements for certain components varied across these phases. For example, some components were only needed for the first half of a sortie, and a failure of these components late in the mission did not constitute a system-level failure. As discussed earlier, the traditional ranking method does not take into account this operational phasing concept.

Although the V-22 analysis effort did not require simulation of all the relevant ranking factors, Raptor could easily have taken into consideration the missing factors and produced similar results just as quickly.

4.2 Raptor Node Analysis Results

The simulation was then run to produce ranking results with respect to mission reliability, that is, the probability of surviving a four-hour sortie. These results provided a realistic ranking of components based on some of the factors previously described. The top twenty components or subsystems contributing to unreliability are displayed in Table 1. When ranking ties occurred, both components were given the same rank but the next ranking number was then skipped. The percentage of mission aborts that each component contributed to the total number was derived from the Raptor weak link analysis data and is also provided in Table 1.

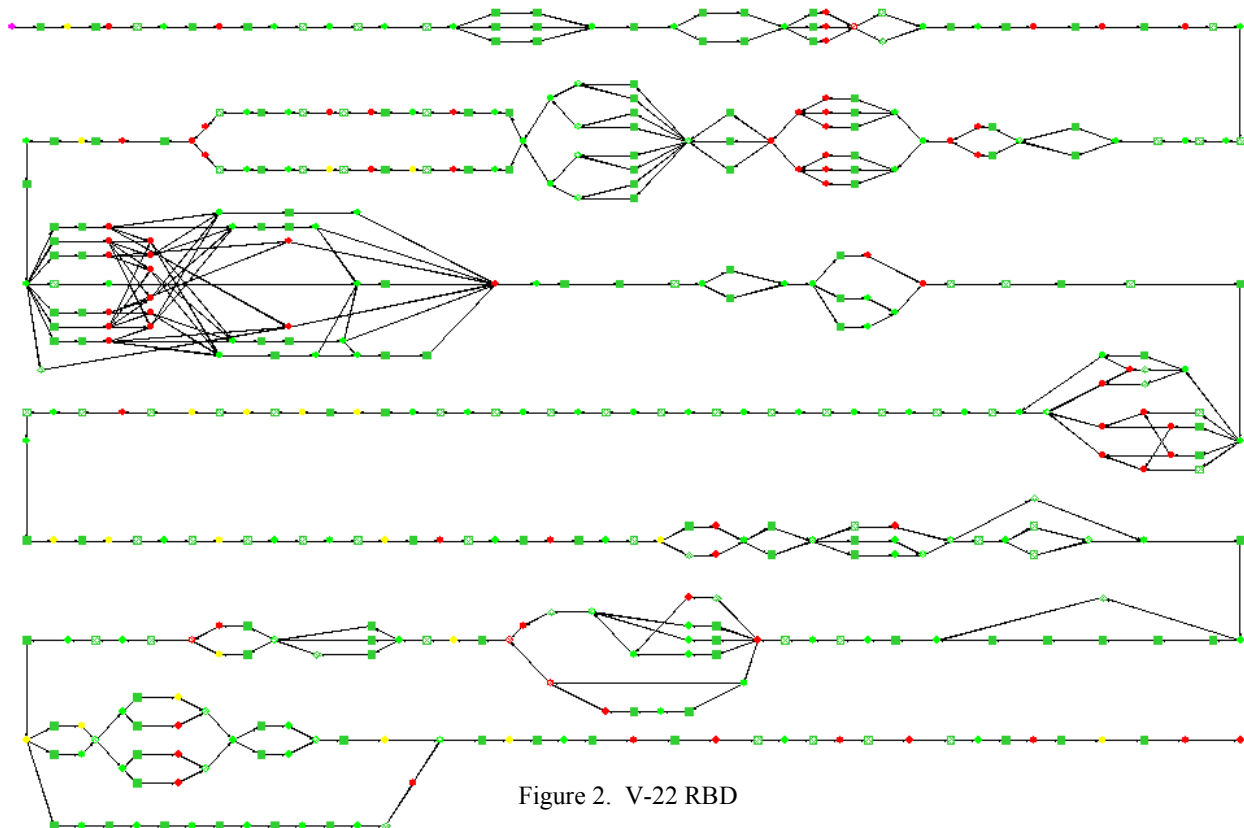


Figure 2. V-22 RBD

Rank	R(4)	%	Subsystem
1	0.9490	17.03	Propulsion
2	0.9710	9.68	Conversion Actuators
3	0.9715	9.52	Electrical Power
4	0.9735	8.85	Swash Plate Actuators
5	0.9790	7.01	Nacelles
6	0.9815	6.18	Fire Suppression
6	0.9815	6.18	Radar Counter Measures
8	0.9890	3.67	Air Conditioning SDC
9	0.9915	2.84	Access Doors
10	0.9935	2.17	Mid Wing Gearbox
11	0.9940	2.00	Flight Control Computer
12	0.9945	1.84	Wiring
13	0.9950	1.67	Hydraulics
13	0.9950	1.67	Interference Cancellor
15	0.9960	1.34	Mechanical Flight Controls
15	0.9960	1.34	Nacelle Interface Unit-1
17	0.9965	1.17	Drive (less gearboxes)
18	0.9970	1.00	Digital Mapping
18	0.9970	1.00	Flight Software
18	0.9970	1.00	Navigation

Table 1. Subsystem Rankings

These top twenty ranked components and component subsystems comprised 87.15% of all failures that would result in a mission abort. It was not surprising to discover that the engines, electrical subsystem and actuators (i.e., Propulsion, Conversion Actuators, Electrical Power, Swash Plate Actuators, and the Nacelles) contributed to more than 52% of all mission abort failures. Most aircraft historically display similar results since these components either contain the most parts or are tasked to perform in the most stressful environments. What was surprising was that the Fire Suppression and Air Conditioning subsystems combined to contribute to nearly 10% of mission abort failures. While modifying any of the top five drivers would be costly at this stage of the acquisition cycle (i.e., low rate production and

field testing), changes to the Fire Suppression and Air Conditioning subsystems would be economically feasible.

To give readers a sense of how the Raptor software displays the node analysis ranking information and how the color tolerances help with visualization of the weak links of a system, Figures 3 and 4 are provided. Figure 3 portrays a subset of the Display subsystem with its localized node analysis ranking values. (Note: Figure 2 displays the color tolerance scheme for all subsystems.)

The color tolerances used for the V-22 case study were allocated as follows:

$$0.9995 \leq \text{Green} < 1.0000$$

$$0.9980 \leq \text{Yellow} < 0.9995$$

$$\text{zero} \leq \text{Red} < 0.9980$$

Figure 4, on the next page, provides a listing of all component reliabilities within a Raptor dialog box, which allows greater precision than can be displayed within the area of a node.

4.3 Comparison to the Traditional Ranking Methods

Conducting the traditional mean life ranking method for the same V-22 data that was used to construct the Raptor simulation resulted in several misleading conclusions. The traditional method indicated that the Multi Functional Displays and Refueling subsystems were among the top twenty drivers. When the dynamic natures of these subsystems are considered, however, they do not appear in Raptor's top twenty driver list. For example, the refueling system is only occasionally used during a flight and this dynamic aspect cannot be properly accounted for with the traditional method. Finally, the traditional method failed to identify three subsystems (Nacelle Interface Unit-1, Drive system less gearboxes, and Flight Software) that accounted for 3½% of all mission abort failures.

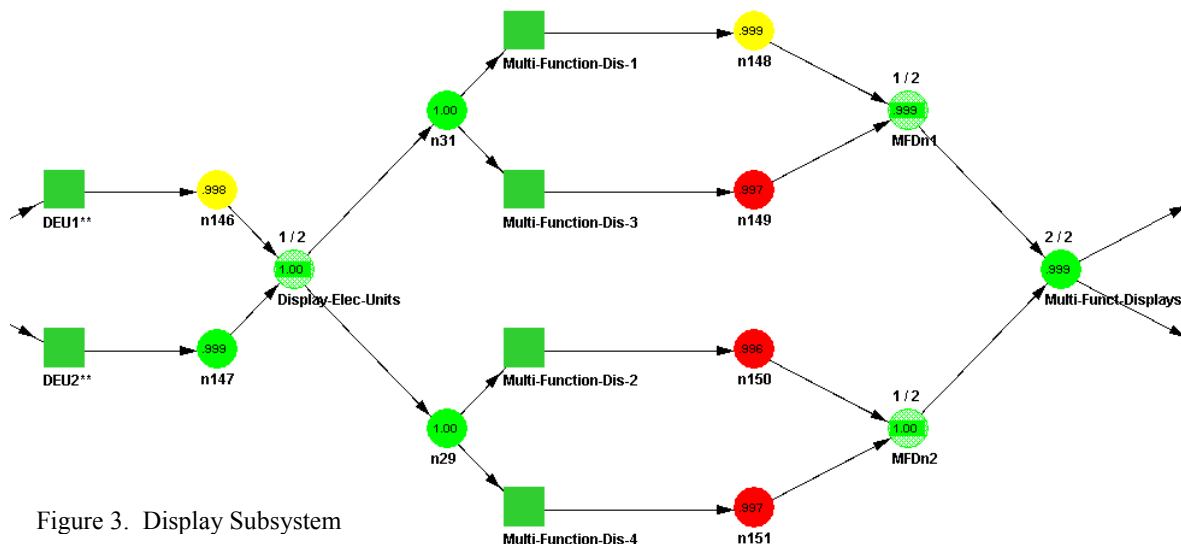


Figure 3. Display Subsystem

Node Name	Reliability	Node Color	Green	Yellow	Red	Blue	Dk Grn	Grey
Propulsion	0.949000000	Red	85.40	0.00	0.00	0.00	11.97	2.62
Conversion-Actuators	0.971000000	Red	85.40	0.00	0.00	0.00	13.05	1.55
Electrical-Power	0.971500000	Red	85.40	0.00	0.10	0.00	13.38	1.12
SwashPlateActuators	0.973500000	Red	85.40	0.00	0.00	0.00	13.21	1.39
Nacelles	0.979000000	Red	85.40	0.00	0.00	0.00	13.62	0.98
Radar-CounterM	0.981500000	Red	90.31	0.00	0.73	0.00	8.96	0.00
Fire-Suppression	0.981500000	Red	85.40	0.00	0.00	0.00	13.53	1.07
Air-Conditioning-SDC	0.989000000	Red	85.40	0.00	0.00	0.00	13.99	0.60
AccessDoors	0.991500000	Red	85.40	0.00	0.00	0.00	14.20	0.39
MidWing-Gearbox	0.993500000	Red	85.40	0.00	0.00	0.00	14.16	0.44
Flg-Control-Comp	0.994000000	Red	85.07	0.33	0.16	0.00	14.44	0.00
Wiring	0.994500000	Red	85.40	0.00	0.00	0.00	14.32	0.27
Interfer-Cancel	0.995000000	Red	85.40	0.00	0.00	0.00	14.31	0.29

Figure 4

5.0 SYNOPSIS

It should be clear that using the traditional ranking method in lieu of the more accurate simulation method could lead to misappropriations of resources and efforts to fix components that will not significantly improve the system's overall availability or reliability. Had component repair distributions, logistical constraints, and more complex dependencies been incorporated in the case study, the gaps between the simulation ranking method and the errant traditional method would have been greater still.

With readily available simulators, there is no reason to employ mathematical techniques that are full of simplifying assumptions. For real life systems, the factors overlooked by these abridged ranking techniques can significantly affect system performance and lead to poor economic decisions. Choosing to ignore the relevant ranking factors can only be considered inadequate analysis.

REFERENCES

1. RAPTOR 6.0 Software, ARINC Engineering Services, LLC, 2001.
2. Kenneth E. Murphy, Charles M. Carter and Steve O. Brown, *The Exponential Distribution: the Good, the Bad and the Ugly. A Practical Guide to its Implementation.* Annual Reliability and Maintainability Symposium, 2002 Proceedings, pp. 550-555.
3. Doris L. Grosh, *A Primer of Reliability Theory*, John Wiley & Sons, New York, 1989, pp 43.
4. Kenneth E. Murphy and Charles M. Carter, *How long Should I Simulate, and for How Many Trials? A Practical Guide to Reliability Simulations.* Annual Reliability and Maintainability Symposium, 2001 Proceedings, pp. 207-212.

BIOGRAPHY

Kenneth E. Murphy
ARINC Engineering Services, LLC
2155 Louisiana Blvd. NE, Suite 4200
Albuquerque, NM, 87110, USA

kmurphy@arinc.com

Kenneth E. Murphy graduated from the University of Colorado with a Bachelor of Science degree in Aerospace Engineering in 1989. In 1990, he graduated from the University of Alabama with a Masters of Science in Systems Engineering and a minor in Reliability Engineering. Ken was a reliability engineer for the United States Air Force for eleven years where he introduced his visions of the way reliability analysis and testing should be conducted. Ken developed the theory of a quick but robust reliability simulator in 1992 and three years later his idea became a reality when the free reliability software tool called RAPTOR was distributed to the world. In 1996 Ken became the chief reliability engineer for the Air Force's operational testing agency where he continued to lead the development of RAPTOR versions 3.0 and 4.0. Ken is currently the Simulation & Analysis Manager with the ARINC Corporation where he is building a reliability skunk-works shop responsible for the development of new versions of RAPTOR, reliability analysis consulting, and RAPTOR training.

Charles M. Carter
ARINC Engineering Services, LLC
2155 Louisiana Blvd. NE, Suite 4200
Albuquerque, NM, 87110, USA

ccarter@arinc.com

Charles M. Carter has a B.S. in Aeronautical and Astronautical Engineering from the University of Illinois and a M.S. in Systems Engineering from the Air Force Institute of Technology. As an officer in the USAF, Chuck was assigned to Vandenberg AFB, CA, where he coordinated ground processing of Titan II and Titan IV payloads and served on the launch crew. He was also assigned to Kirtland AFB, NM, where he worked as a reliability engineer evaluating space and electronics systems. He later worked for SAIC at Johnson Space Center as a payload safety engineer, performing safety analysis on Space Shuttle and International Space Station payloads. Chuck is currently a principal engineer with ARINC and is responsible for development of the RAPTOR simulation engine. He is the original author of the RAPTOR simulation engine and developed the algorithms RAPTOR uses to track all system state changes.

Robert H. Gass
ARINC Engineering Services, LLC
2551 Riva Road
Annapolis, MD, 21401, USA

rhgass@arinc.com

Robert H. Gass has a B.S. in Electrical Engineering from the New Jersey Institute of Technology and a M.S. in Systems Management from the Air Force Institute of Technology. During a 20 year Air Force career he has supported reliability, maintainability (R&M) and logistics efforts on Minuteman, AWACS, B-1A, B-52 Upgrade, cruise missiles, and satellite programs. He is currently a principal engineer with ARINC providing R&M consultant support to the Navy and Air Force V-22 Joint Development Program; providing R&M engineering and R&M flight data analysis. Bob has also provided support to Headquarters Department of Energy (DOE), and National Aeronautics and Space Administration (NASA) Space Station Freedom Level I program office, and Air Force (AF) F-16 program. His DOE support included: developing Monte Carlo cost simulation computer software programs, environmental cleanup cost databases, quality assurance (QA)-R&M programs, and conducting QA-safety appraisals for both nuclear and non-nuclear projects. Additionally his consultant support included: defining requirements for NASA on logistics, safety, R&M, and QA and providing a detailed R&M analysis to support an AF study on modification versus new design decisions for the F-16 Stores Management System Test Set.