

Reliability-Related Safety Analyses for Satellite Navigation Systems

Meng-Lai Yin • Raytheon Systems Company • Fullerton

Craig L. Hyde • Raytheon Systems Company • Fullerton

Larry E. James • Raytheon Systems Company • Fullerton

Key Words: Safety analyses, Reliability analyses, Loss-of-Functionality hazard, Satellite navigation systems, Markov modeling mechanisms.

SUMMARY & CONCLUSIONS

Safety and reliability are two interrelated attributes for safety-critical systems. While the typical safety analysis focuses on preventing hazards associated with erroneous safety critical outputs, this paper introduces an equally important hazard for the loss of critical functionality, referred to as the “loss-of-function” hazard. Tradeoffs are studied among three safety/reliability measures, i.e., the probability of working correctly, the probability of generating erroneous outputs and the probability of losing critical functionality. One of the goals for this study is to assist system engineers in making correct and timely design decisions. A major problem encountered in computing the probabilities of the various safety hazards is the initial condition consideration. This is because a fault-tolerant system can have various operational conditions and a hazard can occur under any of the working conditions, each with different probabilities. To provide a reasonable estimation, a measuring method that incorporates all possible initial conditions is proposed.

1. INTRODUCTION

Safety analyses typically concentrate on preventing a system from unknowingly generating erroneous safety critical output (Ref.7,8). Reliability analyses normally focus on how to make a system work correctly (Ref.3,9). When a safety-critical system requires high reliability, these two attributes interact. Figure 1 shows the two-dimensional framework that we use to address safety and reliability. Initially, a system is constructed with certain levels of reliability and safety features, shown as *design point 1* in the figure. Reliability and safety assessments on the system are performed to estimate these two attributes. If the estimation results are not satisfactory, the design needs to be changed. Design changes can go several directions. In some cases, tradeoffs are feasible. For example, change from design 1 to design 2 (refer to Figure 1) trades reduced reliability for higher level of safety, while change from design 1 to design 3 does the opposite. However, in some situations, tradeoffs are not acceptable. In these cases, enhanced design is needed so that

both safety and reliability requirements can be met, as the change from design 1 to design 4 illustrates.

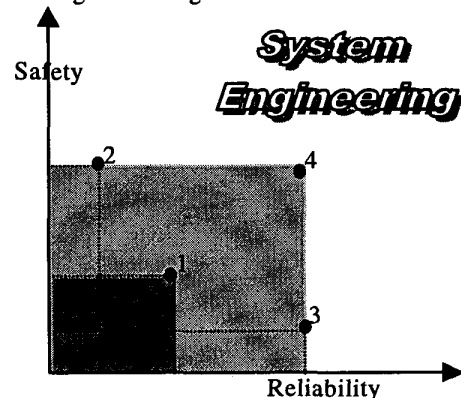


Figure 1. System Safety/Reliability Engineering Framework

This paper addresses issues regarding the safety features that are related to reliability. In particular, the “loss-of-function” hazard is explored as the starting point. Tradeoffs are studied between safety and reliability measures. The topic of fault tolerance design is brought up when a design change such as from design 1 to design 4 in Figure 1 is required. A measuring method that provides a more realistic estimation is introduced.

2. PRELIMINARIES

2.1 The Loss-Of-Function Hazard

Besides the hazard analyses for preventing the generation of erroneous safety critical output the hazard that represents the loss of critical functionality, referred to as the “loss-of-function” (LOF) hazard, is recognized as important for satellite navigation systems. This is because the satellite navigation system provides sole means for navigation. Note that, in general, losing mission-critical functionality is not necessarily “unsafe”. For example, in the nuclear industry, the “fail-safe” design or the “protection system” is widely applied, where an unsafe state of the system will cause the system to intentionally stop functioning for safety’s sake.

2.2 Measures, Tradeoffs and the Basic Model

Three time-dependent measures are used here. They are reliability (or the probability of working correctly), denoted as R , the probability of generating erroneous safety-critical outputs, denoted as P_{HMI} ¹, and the probability of losing critical functionality, denoted as P_{LOF} . These three measures are the cumulative probability of an event occurring during the time interval $[t_0, t]$. Moreover, they are inter-related.

Since safety and reliability have been involved in tradeoff design studies. A typical case of this is demonstrated in the work by Vaidya and Pradhan 1993 (Ref.6), where safety is defined as the probability that the system output is either correct, or the error in the output is detectable. Note that the assumption made is that the system is safe when the error is detected, which is different than what we have for the satellite navigation system. As usual, reliability has been defined as the probability that a system is working correctly in their study.

To address the interrelationship among the three measures described above, this paper employs a simple state diagram, as shown in Figure 2. Here, the fault-error-failure scenario described below is applied. Initially, the system is assumed to be working correctly. However, faults exist and they can cause errors. With coverage C and rate λ , an error occurs and is detected, then the system is brought to the failed state. There is a probability, $1-C$, that the error is not detected. Under this situation, the system continues working, however, incorrectly. When the system is in the undetected state, another error might occur (with rate λ')². With coverage C' , this error is detected and the system enters into the failed state. Otherwise, the system stays in the undetected state³.

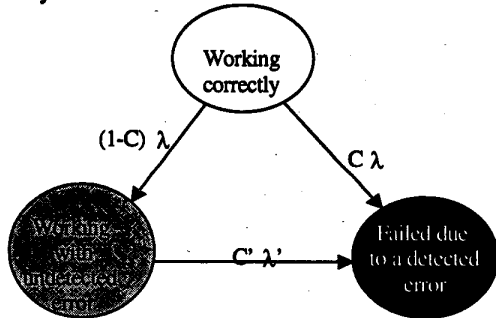


Figure 2. State Diagram for the Fault-Error-Failure Scenario

For systems that require high reliability or availability, reducing failure rates and increasing fault detection coverage are the two commonly used technologies. This paper explores the impacts of these two on the considered measures.

Fault-tolerant systems design, which is applied widely to improve a system's reliability and/or availability, is usually seen in safety-critical systems. In this paper, we will show that fault tolerance techniques not only improve a system's reliability, but also enhance its safety. An impact of the fault

tolerance design that is of particular interest considered here is the measuring method. This is because a fault-tolerant system has more than one "working" state, where each state has different combinations of operating components. Moreover, a satellite navigation system is constantly operating and being repaired, as opposed to mission-oriented systems that can be assumed to be in a perfect state to begin the mission. Since a hazard can occur under any of the working conditions, the assessment of the probability of hazard occurrences should consider all the possible initial conditions. However, for most safety analyses, only the fully operating state is considered as the initial state. The safety measures based on this measuring method *cannot* represent real-life situations. Hence, a method that incorporates all the possible initial conditions is proposed in this paper.

3. THE INTERRELATIONSHIPS

3.1 Time-Dependent Interrelationships

The three measures discussed here are time-dependent functions. In this section we first show that the interrelationships are also time-dependent, as shown in Figure 3. Note that at any point in time, $P_{HMI}(t) + P_{LOF}(t) + R(t) = 1$. Figure 3 illustrates the interrelationships. In this example, the MTBF (Mean Time Between Failure) is one year (8760 hours), and the coverage is 80%.⁴ As shown in the figure, reliability decreases as time increases. Since no repair is considered, reliability will drop to zero when time goes to infinity, theoretically. As predicted, the probability of losing critical functionality increases as time elapses. The probability of generating erroneous outputs increases at first, then decreases due to the subsequent detected errors. Generally speaking, the loss-of-function hazard has a higher probability of occurrence than that of the hazardous-misleading-information hazard when the detection rate is above 80%.

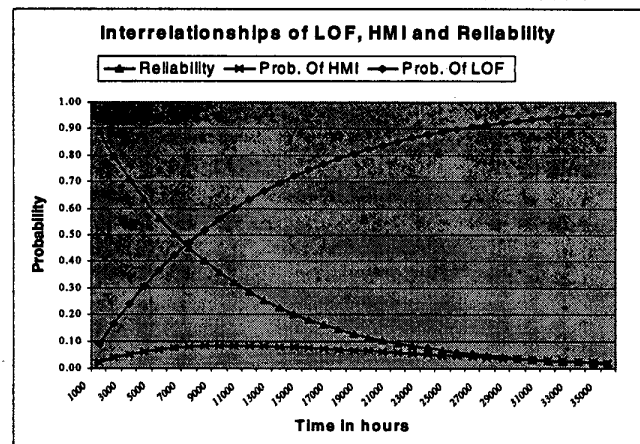


Figure 3. Interrelationships as Functions of Time (MTBF=8760 hours, C=80%)

Note that the probability of HMI initially increases but then decreases with time. This is because the probability of HMI

¹ HMI stands for *Hazardously Misleading Information*.

² The value of λ' may be different than λ , as C' may be different than C .

³ Once a system is in the undetected state, it is assumed that it will not go back to the "working correctly" state for the rest of the modeling time. No repair is considered in this model.

⁴ The results are obtained from using the SHARPE (Symbolic Hierarchical Automated Reliability and Performance Evaluator) tool (Ref.11). In these experiments, λ' is assumed to be the same as λ , and C' is assumed to be the same as C .

events naturally increases with time in an operational system. However, since HMI cannot be produced when the system is not operating, the probability of HMI decreases as the probability of LOF becomes high.

3.2 Impacts from System Designs

For the model shown in Figure 2, *reducing* the probability of being in the state "Working with undetected error" means enhancing safety in terms of preventing erroneous outputs. On the other hand, *reducing* the probability of being in the "Failed due to a detected error" state also enhance the safety, because of the decrease in the probability of the "loss-of-function" hazard. Note that reducing the probability of staying in both the "Working with undetected error" and the "Failed due to a detected error" states actually *increases* the probability of being in the state "working correctly", which is the reliability. Therefore, in general, enhancing safety will improve reliability; there is no conflict between the two.

There are two efforts applied to improve system reliability: reducing failure rates and increasing error detection probability (coverage). When the design or implementation of a system has decreased the value of λ (failure rate), the probability of being in the "working correctly" state is increased, which makes the sum of P_{HMI} and P_{LOF} is decreased. Hence, reducing failure rate enhances both safety and reliability.

For the detection coverage C , note that the changes on the value of C will not change the total failure rate. That is, the rate coming out of the state "working correctly" remains the same, i.e., $(1-C)\lambda + C\lambda = \lambda$. Therefore, theoretically, the reliability will not be affected by the value of C , as it will be for the other two attributes. Specifically, the impacts from the changes of C are on the tradeoffs between P_{HMI} and P_{LOF} .

Let us first examine what will happen when C increases. Usually, when a hazard is identified, mechanism(s) are developed to mitigate the hazard by detecting the errors and providing protections against it. In this case, coverage C will be increased due to this safety activity. The probability of being in the state "working with undetected error" will decrease, while the probability of loss-of-function hazard will increase. Note that now there is a tradeoff between the two properties.

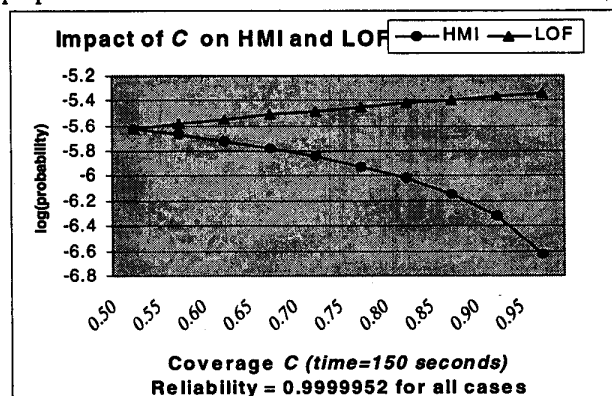


Figure 4. The Impacts of Coverage C on HMI and LOF

Figure 4 shows the changes on P_{HMI} , P_{LOF} due to the increased C . The same parameters are used as in Figure 3, except that the time is now fixed at 150 seconds. As expected,

reliability remains the same, i.e., 0.9999952, throughout all the cases. Thus, only P_{HMI} and P_{LOF} are displayed. As C increases, the probability of HMI decreases, while the probability of LOF increases. This example shows that the improvement of detection capability can reduce the probability of HMI occurrences. However, the effect is the increased probability for LOF.

An implication of the tradeoffs due to the changes of C is the safety requirements applied to P_{HMI} and P_{LOF} . As described above, a system's safety concern is not only limited to the undetected errors, but also the loss of its functionality. In particular, for the satellite navigation system considered, the user's normal operation relies heavily on the information provided by the system, so loss of its function will cause major hazardous results. The *degree of severity* of the loss-of-function hazard depends on the overall navigation environment. If the satellite navigation system is the sole means of aircraft navigation, loss of its function will be dramatic. However, if other navigation methods are also provided, the loss of the functionality will not be as severe as in the "sole-means" case.

Usually, the "failure condition categories" are used to distinguish the severity of hazards. In particular, the levels specified in (AC25-1309-1) categorize five safety hazard results levels for aviation systems. These 5 levels are *catastrophic*, *hazardous/severe-major*, *major*, *minor* and *no effect*. One can apply this categorization to quantitative specification for hazards. For example, a hazard categorized as having major safety effects will have the requirement that its occurrences shall not exceed 10^{-9} per hour; while another hazard categorized as having minor effects will have the requirement saying that its occurrences shall not exceed 10^{-7} per hour.

This quantitative specification based on the severity level has impacts on systems design. As explained before, the increment of C will decrease P_{HMI} but increase P_{LOF} . When the safety effect for the loss-of-function hazard becomes more severe, P_{LOF} requirement will be changed (to a smaller probability). However, the requirement on P_{HMI} remains the same. Under this circumstance, improving the error detection coverage is not enough. Other mechanisms such as fault tolerance technologies that can improve the overall system robustness are required.

3.3 Safety Marginal Errors

In some cases, a safety-critical system will tune its functionality so that certain safety marginal errors can be detected. For example, the satellite navigation system can tune its functionality so that when the calculated position falls outside of an allowed range, the output is declared erroneous. In some cases, higher accuracy is desirable for safety's sake. Thus the allowed range becomes smaller, and a higher failure rate will result because now there is a higher probability that the calculated position will fall outside of the allowed range. In terms of the model, the value of λ is increased, which will decrease the reliability. The errors due to this safety concern are referred to as *safety marginal errors*.

When we redefine the conditions for "working correctly" due to the safety marginal errors, the failure rate of the overall system increases. This is undesirable from the reliability

standpoint. However, if the architecture of the system has already been determined, or the system has already been built, tuning the system to protect against certain hazards seems to be a more feasible approach than changing the design. Since the price paid for taking care of the safety marginal errors is lower reliability, the tradeoff between safety and reliability becomes an issue. Usually this causes the debates between safety and reliability. Nevertheless, we have to keep in mind that this tradeoff occurs because of the “fixed” system architecture. When no matter how the tradeoff is made, neither the safety nor the reliability requirement can be satisfied, changing the system architecture becomes inevitable.

The safety marginal errors normally have detection coverage higher than the original coverage C . Figure 5 shows a simplified model reflecting the effects of safety marginal errors. Note that the safety marginal error phenomenon can be much more complicated. The model used here is just for demonstration. As shown in the figure, the total failure rate coming out of the “working correctly” state is increased by λ_M . The detection coverage for the safety marginal errors is denoted as C_M , which is anticipated to be higher than C . In Figure 6 and 7, the impacts of the safety marginal errors on the measures are shown, with λ_M ranging from 10% of the original failure rate to 100%, C_M is set to 0.99999999.

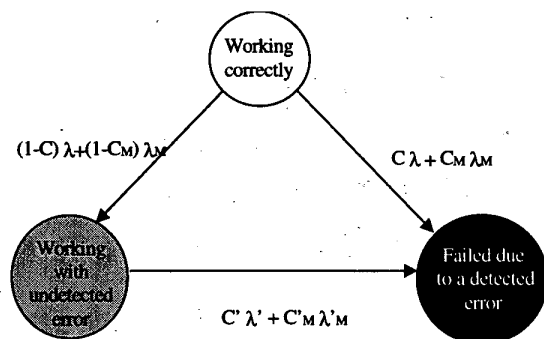


Figure 5. State Diagram with Safety Marginal Errors

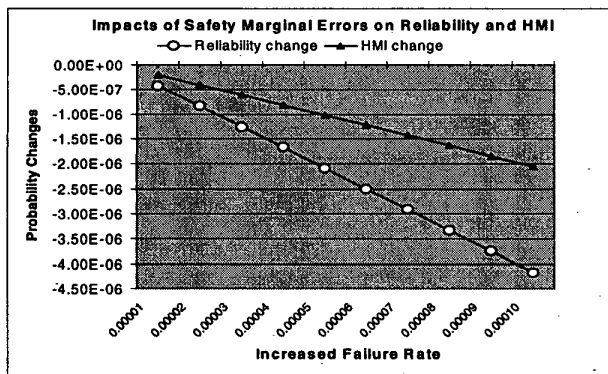


Figure 6. The Impacts of Safety Marginal Errors ($t=150$ seconds, $MTBF=8760$ hours, $C_M=0.99999999$, $C=0.8$)

Note that the probability P_{LOF} changes are on the order of 10^{-1} (as shown in Figure 7), while the changes for reliability and P_{HMI} are much smaller (on the order of $10^{-6} \sim 10^{-7}$, as shown in Figure 6), with the increased failure rate ranges from

10% to 100%. This is based on the situation where the detection coverage for safety marginal errors is extremely high, i.e. 0.99999999. This example shows that a safety conservative approach that considers safety marginal errors will have more severe impacts on the probability of loss-of-function hazard than to the other two measures. Moreover, as shown in Figure 6, the decrease of P_{HMI} is slower than that of the reliability, which means the rate of decreasing P_{HMI} is lower than that of decreasing reliability. Thus, this specific conservative approach to detect more safety marginal errors for the purpose of decreasing P_{HMI} has to pay the price of more-rapidly-decreased reliability. Moreover, the probability of LOF increases with an even higher rate.

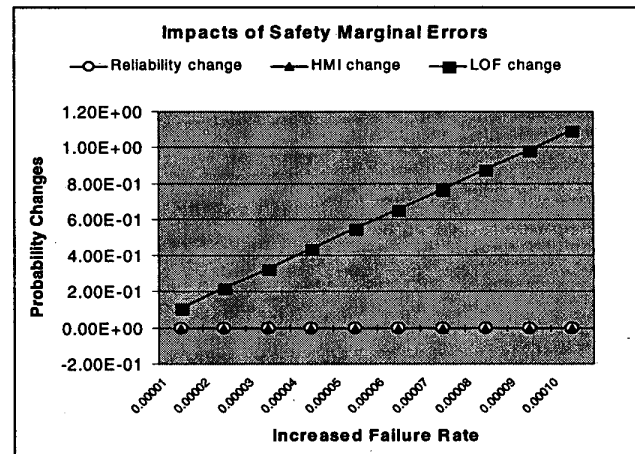


Figure 7. The Impacts of Safety Marginal Errors – with LOF probability ($t=150s$, $MTBF=8760hrs$, $C_M=0.99999999$, $C=0.8$)

3.4 Fault Tolerance Design and Its Impacts on Safety/Reliability Measures

Fault tolerance design has been applied to improve reliability and/or availability for decades (Ref.9). The basic concept of fault tolerance design is that a system can continue working with the existence of faults. A common technique used is the triple-modular-redundancy (TMR) scheme. In this scheme, three modules are working simultaneously, where the outputs of the three modules are voted so that the faults due to a faulty module can be detected and the erroneous output can be masked. When a faulty module is identified, the remaining two modules still can check each other's outputs. Therefore, the fault detection capability still exists when there are only two modules left. This fault tolerance technique improves the detection coverage C due to the voting and checking mechanism.

Fault tolerance techniques increase the reliability or availability, according to their redundancy nature. To see this, consider a single module having reliability p equals to 0.999. TMR with three such modules will have reliability $p^3 + 3 \times p^2 \times (1-p)$, calculated as 0.999997. Thus, in general, fault tolerance techniques will decrease the overall failure probability and increase the detection coverage C . As indicated previously, these are the characteristics that can

enhance safety. Thus, fault tolerance design can enhance reliability and also improve safety.

One problem with the fault tolerance design on the issue of safety measurement is the initial condition consideration. Traditional reliability analysis considers only two states for a system; that is, a system is either working or failed. However, a fault tolerance system has more than one working state, denoted as the *degraded mode states*. When a system is in a degraded mode state, it still provides service, however, some components are not functioning. Note that the degraded-mode phenomenon for performance measurement have been studied within the area of "performability" (Ref.4).

4. RELIABILITY-RELATED SAFETY MEASURES

4.1 Initial Working Conditions

The safety measurements discussed are time-dependent functions. For aircraft navigation applications, time-dependent assessments are important because of the nature of a flight mission. The probability of a hazard occurrence for a time period, such as en-route or precision approach, with the given initial working condition, is of primary concern for an aircraft. If the system were not degradable, then the initial condition will be either that the system is fully operational, or the system is not working at all. However, the system considered is, indeed, fault-tolerant. Which means the initial condition can be any state from the non-working state to the fully operational states. The point is that the estimation may be made at any point in time. At that particular moment, the system can be in any of the working conditions. The key to reasonably access the safety attributes of the system is to consider all possible initial conditions.

For the TMR example described above, initially all 3 modules are working and the outputs are checked via a voter. A faulty module is identified if its output is different than that of the others. Once an error is detected, the faulty module is removed and the TMR becomes a Dual-Redundancy module where the 2 modules continue working and checking their outputs with each other. A subsequent error can occur and may be detected when only 2 modules are working. However, when a mismatch occurs among the two remaining modules, the checking mechanism cannot identify which module is faulty. Hence, the second detected error will cause the whole TMR to be failed, referring to the state diagram shown in Figure 8.

In this particular example, an undetected error will bring the TMR module to the undetected state. When the system is in the undetected state with 3 modules working, a subsequent detected error can happen (with rate $3C\lambda$), which will bring the system to the 2 modules working state. A subsequent error that is not detected will keep the system in the same state. From the 2 modules working state, an undetected error can occur, which transitions the system into the "2 modules with undetected errors" state. Similarly, a subsequent detected error will bring the system to the failed state. Note that at any point in time, the initial condition can be (1) all 3 modules are working without undetected errors, (2) only two modules are working without undetected errors, (3) 3 modules are working with undetected errors or (4) 2 modules are working with undetected errors. Note that it is possible that the system is in

one of the undetected states when the safety measurement is being assessed. The reason for distinguishing these initial conditions is that the probabilities of a hazard occurrence are different for different initial conditions.

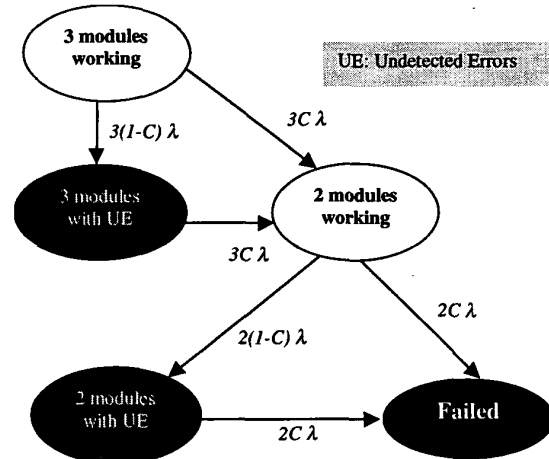


Figure 8. State Diagram for the TMR

4.2 Safety Measure as an Attribute of an Operational State

A safety measure is a property of an operational state⁵. In other words, as long as the system is operational (either fully working or partially working, or even worse, with undetected errors), there is a safety concern. Given that, at time t , the system is in an operational state S_i , S_i can be the fully operational state or a degraded state. The probability that the system will reach to the undetected state during $[t, t+\Delta t]$ is denoted as $P_{HMI}[S_i, t, \Delta t]$. As described before, $P_{HMI}[S_i, t, \Delta t]$ is different than $P_{HMI}[S_j, t, \Delta t]$, where $i \neq j$.

From the discussion above, the probability of a hazard occurrence for a time period depends on the initial condition when the estimation is made. However, it is unfeasible to present the hazard occurrence probabilities for every single initial condition, especially for a large system like the satellite navigation system considered. To see this, suppose the system has 10 receiver stations to receive the satellite information and 3 out of the 10 receiver stations need to be operating in order for the system to be considered "working". Hence, a state with m working stations, where $3 \leq m \leq 10$, will be considered a "working" state. That means, even if every station is either working or failed (without considering the undetected errors states), then totally $2^{10} - \left(\binom{10}{0} + \binom{10}{1} + \binom{10}{2} \right) = 968$ different

degraded modes exist and all these 968 states will have a safety attribute! Due to this complexity, we need a more feasible way of presenting safety measures that take different initial conditions into account.

Denote the probability of HMI with the initial condition being the fully working state as P_{HMI-F} . If the initial condition is a degraded state, the probability of HMI is then denoted as P_{HMI-D} . Moreover, when the initial condition is an "undetected errors" state, the probability of HMI is denoted as P_{HMI-U} .

⁵ Assuming the system is not repairable. For repairable systems, the failed state can have safety concerns, too.

Although there could be more than one degraded state or more than one undetected state, for simplicity, we consider only one degraded state and one undetected state. When the safety measurement is made, there is a probability that the system is in the fully working state, denoted as p_F . There is also a probability that the system is in the degraded state, denoted as p_D . Similarly, p_U is the probability that the system is in the undetected state. Note that normally when such a safety measurement is being made, the system has been operating for quite some time, and thus should be in the steady state condition. Therefore, p_F , p_D and p_U are the steady-state probabilities. When the model is irreducible, steady-state solutions can be yielded.

Thus, a hazard occurrence can be described by a discrete random variable X . The value of X is defined as the probability of a hazard occurrence. For the above example, the values of X can be P_{HMI-F} , P_{HMI-D} or P_{HMI-U} . Thus, the probability that $X = P_{HMI-F}$, denoted as $P\{X = P_{HMI-F}\}$, is p_F . The probability that $X = P_{HMI-D}$, denoted as $P\{X = P_{HMI-D}\}$, is p_D . Similarly, the probability that $X = P_{HMI-U}$, denoted as $P\{X = P_{HMI-U}\}$, is p_U . Since X is a discrete random variable, the expected value of X , denoted as $E[X]$, is a weighted average of the possible values that X can take on. Thus, the expected probability of HMI, given the system is operating initially, is

$$E[P_{HMI}] = \frac{p_F}{p_{Total}} P_{HMI-F} + \frac{p_D}{p_{Total}} P_{HMI-D} + \frac{p_U}{p_{Total}} P_{HMI-U}$$

where $p_{Total} = p_F + p_D + p_U$.

The result from this calculation will be a higher hazard occurrence probability, i.e. $E[P_{HMI}]$, than the result without considering the degraded mode, i.e., P_{HMI-F} , as shown in Figure 9. For the steady-state probabilities, a repair time of one-hour is used in the model. Since a fault tolerance scheme is used, a higher coverage $C=0.999$ is applied. As shown in Figure 9, considering all the possible initial conditions gives a much higher estimation on the probability of HMI. The difference in this example is on the 10^{-4} order of magnitude, which is considered to be significant for safety-critical systems. This example demonstrates that the estimation based only on the fully operational initial condition is optimistic and unrealistic.

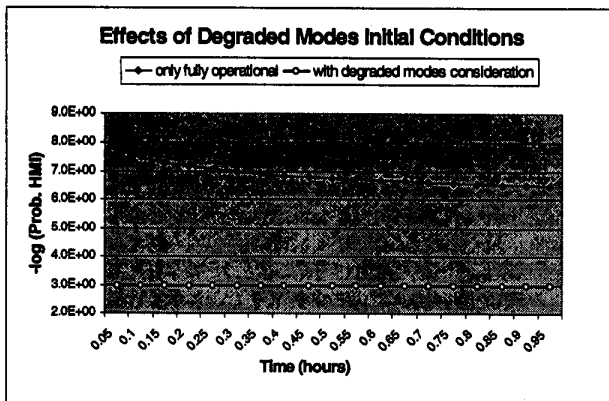


Figure 9. The Effects of Considering All Possible Initial States

REFERENCES

- [1] Beaudry, M. "Performance-Related Reliability Measures for Computing Systems," *IEEE Trans. On Computers*, Vol. C-27, No.6, June 1978, pp.540-547.
- [2] Choi, C.Y., Johnson, B.W. Profeta, J.A., "Safety Issues in the Comparative Analysis of Dependable Architectures," *IEEE Transactions on Reliability*, Vol.46, No.3, Sep. 1997.
- [3] Johnson, B.W., *Design and Analysis of Fault Tolerant Digital Systems*, Addison-Wesley Publishing Company, Inc., 1989.
- [4] Meyer, J.F., "On Evaluating the Performability of Degradable Computing Systems," *IEEE Transactions on Computers*, Vol. C-29, No. 8, 1980.
- [5] Smith, R.M., Trivedi, K.S. Ramesh, A.V., "Performability Analysis: Measures, an Algorithm, and a Case Study," *IEEE Transactions on Computers*, Vol. 37, No. 4, 1988.
- [6] Vaidya, N.H., Pradhan, D.K., "Fault-Tolerant Design Strategies for High Reliability and Safety," *IEEE Transactions on Computers*, Vol.42, No.10, Oct. 1993.
- [7] *Global Positioning System: Theory and Applications*, edited by Bradford W. Parkinson and James J. Spilker Jr., American Institute of Aeronautics and Astronautics, Inc., 1996.
- [8] *Safety-Critical Systems*, edited by Felix Redmill and Tom Anderson, Chapman & Hall Publisher Inc., 1993.
- [9] *Reliable Computer Systems --- Design and Evaluation*, Daniel P. Siewiorek, Robert S. Swarz, Digital Press, 1992.
- [10] *Applicable Regulations and Guidance Material*, Federal Aviation Administration AC 25-1309-1A.
- [11] *Performance and Reliability Analysis of Computer Systems --- An Example-Based Approach Using the SHARPE Software Package*, by R.A. Sahner, K.S. Trivedi, and A. Puliafito, Kluwer Academic Publishers, 1996.

BIOGRAPHY

Meng-Lai Yin, PhD

Raytheon Systems Company
Loc. FU, Bldg. 675, M/S AA341
Hughes Drive, Fullerton, CA 92834 USA
email: mlvin@west.raytheon.com

Meng-Lai Yin received her MS degree in electrical engineering from National Cheng-Kung University, Taiwan, in 1983, the MS degree and the PhD degree in information and computer science from University of California, Irvine, in 1989 and 1995, respectively. Since 1990, she has been working at the Fullerton site (first, for Hughes Aircraft Company and now for Raytheon Systems Company).

Craig L. Hyde, PhD

Raytheon Systems Company
Loc. FU, Bldg. 675, M/S AA341
Hughes Drive, Fullerton, CA 92834 USA
email: clhyde1@west.raytheon.com

Craig L. Hyde received his MS degree in physics from University of Michigan, and Ph.D. degree in applied math from University of Arizona. He is currently working for Raytheon Systems Company, on reliability modeling tasks for the Wide Area Augmentation system.

Larry E. James

Raytheon Systems Company
Loc. FU, Bldg. 675, M/S AA341
Hughes Drive, Fullerton, CA 92834 USA
email: lejames@west.raytheon.com

Mr. James has been with Hughes/Raytheon for 28 years and is responsible for directing systems effectiveness technical activities (reliability, maintainability, fault tolerance, systems safety and life cycle cost) for technology studies and advanced projects within the Command and Control Systems Division. He holds an BS in Mathematics from St. Mary's Univ. and an MA in Mathematics from UCI.