

Forecasting the Reliability and Safety of Future Space Transportation Systems

Joseph R. Fragola<PK>, SAIC New York

Key Words: crew safety, forecasting, launch vehicle, reliability, spacecraft design

1. INTRODUCTION

The Space Shuttle, despite the dramatic accident, which occurred on the 51L mission, has proven to be a reliable launch vehicle. In fact, the demonstrated reliability of the Space Shuttle system exceeds that of all other launch vehicles in the US inventory and that of other western nations. With the recently incorporated Phase II upgrades to the system the mission reliability is expected to be well beyond that of even the most reliable of previous launch vehicles, the Russian Soyuz.

Despite its unparalleled reliability the space shuttle remains extremely expensive to operate as compared to conventional expendable launch vehicles. Also its four-vehicle fleet, which was inaugurated in 1981, is aging. In addition, the limited options in the current design for crew survival and recovery given an accident, severely limits the safety of the crew aboard to the level of safety provided by the launch system alone. Given these constraints it is difficult to see how the current Space Shuttle system will ever make earth to orbit travel as safe as has been envisioned in a recent speech by the NASA Administrator [1] and as has been contained in a recently published NASA document. [2]

These observations suggest that sooner or later in the next century the US will be required to replace the current Space Shuttle system. However, although it may be clear that eventual replacement seems inevitable what is not clear at all is what system should be chosen as the replacement nor when should this replacement take place? The options for replacement abound from modifications to the current system, to extending current expendable systems to allow for human transport, though the development of new "human rated" expendables, to building an entirely new and revolutionary single or multiple stage completely reusable launch vehicles.

Although not the only consideration in the difficult decision among alternative approaches, Crew Safety and Mission Reliability have been clearly identified as NASA's top priority [1]. For this reason it is important that each of the alternative options is evaluated in terms of these two characteristics on a comparative basis. But how can a well known and existing option with known, but perhaps unacceptable, levels of safety and reliability be compared to non-existing proposed options that promise higher levels at significantly higher risk of attainment?

Co-equal comparison requires that the measures of safety and reliability be well defined and that the elements, which contribute to each of these measures, reflect the nature of the individual design option. Further co-equal comparison

requires that demonstrated performance, as in the case of the current Space Shuttle system, be given sway as compared to the forecasted potential of future designs. Additionally, to be credible, any forecast must be soundly based not only on an understanding of the development of space systems, but also on an understanding of the historical elements of transportation systems which have tended to drive safety and mission reliability. Finally forecasts, which include extensions to the historical database, which is the case for all new designs, must properly and explicitly account for the uncertainty involved in these extensions.

2 HISTORICAL CREW SAFETY

Although it may be nostalgic to think back to the days of the early space program and, in this 30th anniversary year of Apollo 11, to reminisce about the "Good old days" of early manned space flight it is well to remember that they were not all that good. As has been mentioned [3], at the time of Kennedy's Moon announcement in January of 1961, only 50% of US launched payloads had achieved orbit, and the best US launcher, the Redstone, had barely demonstrated a 75% reliability. In fact one of the reasons why a suborbital mission was chosen first over an orbital mission might have been that the manned orbital launcher, the Atlas, had barely demonstrated 50% reliability. The fact is that, despite the stellar record of the pre-shuttle manned space program, there were a relatively few flights. For this reason even with no failures in the program* the pre-Challenger shuttle had **demonstrated** a reliability, in terms of the number of successful flights, greater than that demonstrated by all of the Apollo lunar landing missions. In fact, as is shown in Figure 1, even including the Skylab and Apollo-Soyuz successes produces such a large uncertainty that the pre-Challenger shuttle performance lies well within the bounds of uncertainty. Moreover, recent estimates place the current shuttle at the mean of the entire Apollo era with far less uncertainty in the estimate.

*One could argue that the Apollo 1 fire should be counted as a crew safety failure, but including this event would only strengthen the point being made.

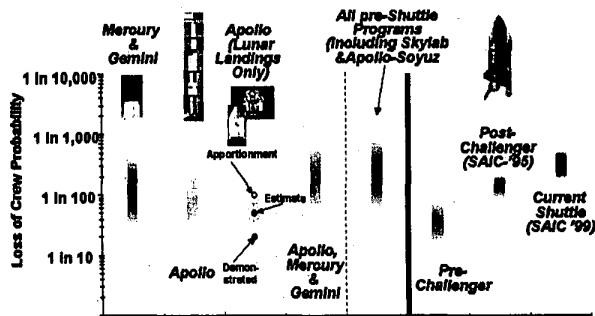


Figure 1: Historical Crew Safety

3. Current Launch Vehicle Reliability

While the demonstrated safety level of the shuttle compares well with history this comparison is not altogether fair. There are no US manned launch vehicles other than the shuttle. Therefore to compare the shuttle with vehicles which are no longer flying does not allow for the possibility that, if they had continued flying, they might have improved their demonstrated performance. In fact expendable launch vehicles have been shown to demonstrate significant reliability growth throughout their program life [4]. In one case, the Soyuz, where the design was allowed to mature in over 1000 flights, a crew safety record better than the shuttle has been demonstrated. This is because there has been only one admitted failure of the Soyuz on return, which would give the crew safety nod to Soyuz on the basis of manned program history. However not all of the Soyuz missions were manned and the overall mean missions between failure reliability for all missions as of 1994 of about 1 in 51. This performance is still very good but about a factor of two less than the shuttle using flight history alone.

Direct comparisons of the shuttle, as a launcher with either the Soyuz or with other launchers using total mission reliability is not fair because the shuttle has a descent and return mission as well as a launch mission on every flight. If the shuttle is compared as a launcher to other launchers, using only the assessed launch portion of the shuttle risk, the shuttle demonstrates a far superior in launch reliability in comparison to other alternatives as is shown in Figure 2.

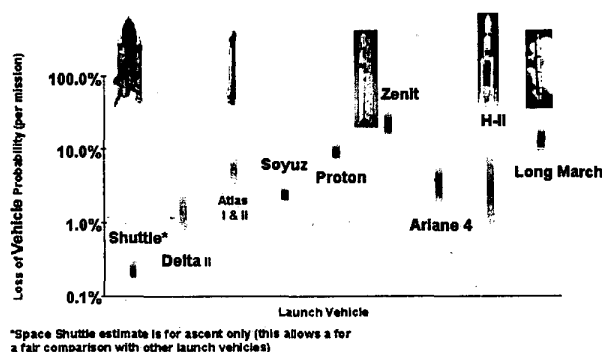


Figure 2: Current Launch Vehicle Safety/Reliability

4. THE IMPACT OF CREW SURVIVAL AND RECOVERY SYSTEMS

There are four traditional ways of preserving the life of a crewmember or passenger. First, and foremost, is by the inherent reliability of the vehicle. Second, is by intact abort with an accident initiated. Third, is a rapid separation from the effects of an initiated accident and recovery. Fourth, is the protection of the crew in an accident from the effects by employing a survivable structure and recovery. The first two methods are used exclusively for highly reliable and highly survivable vehicles like commercial aircraft, passenger rail, and commercial and passenger automobiles. The third method is used in military fixed wing aircraft and for traditional expendable launch vehicles with manned spacecraft, and the fourth method is used for motor racecars and helicopters. All of these systems are designed to have a high probability of avoiding or surviving the effects of the initial insult and all have been shown to provide a significant measure of protection. In a probabilistic sense it is said that these systems are designed to be **independent** of the initiation and conditions of the accident or, alternatively, that they are designed to minimize their "correlated" or "common cause" failure potential with the initiated accident.

Analytically this implies, for complete independence, that the probability of crew survival (or crew safety) equals the probability that no accident is initiated, **plus** the product of the probability that the accident is initiated and that the crew survives and is rescued. When the probability of an accident is small this additional probability does not add very much to an already high probability of survival. However when the probability of an accident is reasonably high, such as in higher risk systems like racecars, high performance military aircraft, and launch vehicles, this additional probability becomes important. To see how important consider the late Alan B. Shepard's initial suborbital flight on the Redstone. The Mercury system included a launch escape tower system which was extensively tested and, although never used it might be assumed to have been proven by test to have been successful 90% of the time in the event of a booster failure during launch. If this were true then the risk to Shepard's life from a launch accident would change from a considerable 1 in 4, to a much more acceptable, but still sporty, 1 in 40; a factor of 10 increase in safety. This might explain Shepard's lack of concern, as quoted elsewhere, [5] despite his knowing the risk.

In general the factor of safety improvement provided by such systems is given by 1 minus their reliability; presuming that they are independent of the initiated accident. But the important message is not the factor of safety provided by these systems alone, but rather it is the safety provided by the inherent reliability of the launcher **combined** with these systems. For example the risk of air crash as reported by the FAA is about 1 in 2 million departures. Changing this to 1 in 4 million or even 1 in 10 million by providing an escape system for passengers and crew would certainly not be worth it in terms of the cost of the risk averted in this case. However the addition of a factor of two improvement might well be worth it on a space launch system with a 1 in 10 or 1 in 20 inherent reliability as would be the case if the current stable of US

unmanned expendable launchers were considered for manned flight.

5. WHAT SHOULD THE ACCEPTABLE RISK OF SPACEFLIGHT BE?

This question is not a new one but goes back to the beginning of the US space program at least [5]. More recently the NASA administrator has indicated a preference for a goal of 1 in a million flights consistent with current US commercial air carrier performance [1]. NASA Johnson Space Center (JSC) has proposed [2] a spectrum of goals which, quite reasonably, vary with the nature and the frequency of the mission. The JSC formula is somewhat complex so the reader is asked to refer to the reference for specifics, but consider the simple case of transferring crew back and forth to the planned International Space Station. In this case the Earth to Orbit (ETO) launch is designed for pure, routine transportation of the crew to the facility at which they are to perform research. It would not be unreasonable to suggest that their risk of transport, on a yearly basis for example, should be the same as that of a typical US active business traveler. If the business traveler takes 50 business trips a year, and requires one stop over on average per trip then he accepts the risk of 100 departures per year or a 1 in 10,000 loss of life risk. To impose any less risk on a space traveler would seem unreasonable. Therefore, in the equal risk per year case, a space crewmember will travel to the ISS at most once per year on average so a 1 in 10,000 goal would seem reasonable for such transportation systems. As it turns out this goal is consistent with the JSC requirements, although they are arrived at differently, for an ETO per flight risk for a program with 100 flights.

What about other missions? The JSC document suggests lower goals depending on the nature of the mission. Specifically, it is suggested that lower goals be applied to non-routine missions, such as test flights of new vehicles or pioneering missions to the moon or to Mars. Suppose it is further agreed that test flights of vehicles intended for routine use should have higher goals than pioneering flights, and finally that pioneering flights should have goals consistent with the 1 in 100 crew safety goal set for the Lunar Module, for example. [5] In this case a consistent set of crew safety goals would be:

1. ETO flights: 1 in 10,000 missions
2. Experimental or developmental flights for ETO vehicles: 1 in 1,000 missions
3. Pioneering flights: 1 in 100 missions

6. WHAT WOULD IT TAKE TO ACHIEVE THE ETO GOAL?

One way to answer this question would be to be to consider what would be the reasonably demonstrated capabilities of a crew return system, that is either a survival and recovery or escape and recovery system. As has been indicated above, a suitable and comprehensive test program should be able to test these systems so as to demonstrate a 90% success rate (These systems have been used only once,

in Russia, under actual launcher failure conditions. That one time the system worked perfectly). A higher level might be possible, but the uncertainties in the test to flight failure conditions might make this difficult. In any case given a 1 in 10 chance of failure of the crew return system the launcher would require a 1 in 1,000 inherent reliability for the attainment of a 1 in 10,000 as is indicated by Figure 3. Since the launcher goal must include both immediate catastrophic failure as well as aborts which are not successful this implies that the goal for immediate catastrophic failure must be higher than 1 in 1,000, but this, for simplicity, will be the goal used here. (If it is assumed that intact aborts are highly successful, for example if they only fail once in 6,000 attempts, then the immediate catastrophic failure goal would be close to 1 in 1000, actually 1 in 1200 as the figure indicates.)

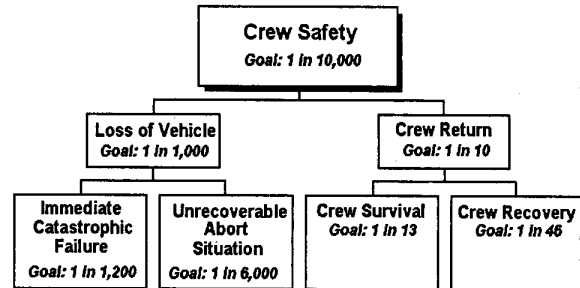


Figure 3: Example Risk Allocation - 1:10,000

As Figure 1 shows the inherent reliability of the current shuttle is about 1 in 200 to 1 in 400, depending upon how much credit is given for current and future upgrades to the system. At the 1 in 400 level the further 63% reduction in risk (about a factor of 2 or 3 increase in crew safety) would be challenging but possible. One reasonable apportionment of the risk is shown in Table 1. As the table indicates the large risk reductions would be expected in the Orbiter vehicle overwhelmingly in areas related to the descent and landing phases. If these were possible the orbiter risk would be reduced to 1 in 3000, and the propulsion-dominated portion would be the risk driver at 1 in 1300. Some might argue that these levels of further reduction are unreasonable to expect out of a system as mature as the shuttle, and many would rightfully argue the specifics of where these reductions, if any, are likely to come. However, despite these arguments, it is fair to say that these reductions are at least in the realm of possibility.

Pathways to 1 in 1,000

Element	Upgrade Risk	Next Generation Risk	Percent Reduction	Justification
SSME	1693	3000	47	More benign cycle, materials improvement
Propulsion	1382	2700	50	Reusable first stage
ET	6844	8000	17	Design maturity
MSFC Total	419	1300	68	
TPS	3600	40,000	90	Assured TPS Status On-orbit
Landing	2012	10,000	60	Landing Gear, Landing commit aids, Landing aids, Crew Interface
Other	1100	6000	80	Growth in hardware reliability
Orbiter Total	433	3000	86	
Shuttle	mean	336	910	
	median	370	1000	

Table 1: Pathways to 1 in 1,000

What this implies is that it is possible, even if not likely, for the current space shuttle system to achieve the ETO Crew Safety Goal of 1 in 10,000 if it had a 90% crew return system. However the crew return system would have to be 90% viable throughout all of ascent and descent to landing or at least proportionally so. Unfortunately even if the current escape pole bailout system were considered 100% reliable over the flight regimes over which it is applicable its flight regime is extremely limited so as not to afford much improvement over the inherent reliability of the shuttle.

To understand how limiting this lack of a comprehensive crew return system is to the shuttle let us suppose the US were to build a new expendable launcher and suppose it would be able to demonstrate the current launch reliability levels of the Soyuz, say 1 in 50. Even at this level it would be inherently less reliable than the current shuttle by a factor of 4 for the entire mission, and by a factor closer to 10 of what the shuttle might be able to achieve with all the upgrades incorporated. However suppose we add a crew return system with 90% reliability. Then the crew safety level of this launcher would be 1 in 500 or equal to that of the shuttle assuming that all the upgrades are incorporated and they capture all the risk they are directed at. If the new launcher would be only a factor of 2 better than the Soyuz, which seems feasible, then the system would be 10 times better than the shuttle even with all the enhanced upgrades taking full risk reduction effect.

7. THE CREW SAFETY DILEMMA FOR REUSABLE SPACECRAFT DESIGNERS

The above discussion points out how past decisions on the shuttle have had a lasting effect on the ability to achieve high levels of crew safety. That is in the case of the shuttle crew safety depends solely on launch reliability and intact abort capability much as is the case for a commercial airliner. As was mentioned this design approach is quite reasonable for designs with 1 in a million reliability levels, and is even reasonable for an ETO vehicle with a 1 in 10,000 inherent reliability as some in NASA believed of the shuttle even after the 51L accident [5]. However this design approach is unreasonable if inherent reliabilities are limited, even under the most optimistic design scenarios, to 1 in 1000 when high levels of crew safety are desired.

Further the mistaken optimism by some in NASA as to the shuttle's reliability led to the early abandonment of crew return design development for the shuttle. In fact after the first few "developmental" flights the commander and pilot ejection seats were pinned so that they could no longer be activated, and then eventually removed, so as not to discriminate against other crew members who did not have this means of escape. This design abandonment was critical to future manned spacecraft design for many reasons, but two of the most important are: 1. The lack of design consideration at the initial design makes the addition of comprehensive crew escape to the existing design expensive in cost, weight, and probably most importantly center of gravity (cg), and 2. The design of a crew return system for any tandem design (i.e. any design which configures the crew astride the fuel

sources) requires different design considerations than those of the conventional in-line stack design.

The first reason was demonstrated by the study completed shortly after 51L. This study [6] concluded that the current, severely limited, pole system was the only viable design configuration that could be reasonably incorporated in the shuttle. The second reason has only recently come to light by studying the details of the 51L event. Detailed investigation into the 51L-accident timeline indicated that, at most, there was 7 seconds available between the first detectable indications of a problem and the deflagration. Further, additional analysis has indicated, that the 51L type of event which progresses very rapidly from initiator to crew involvement are among the most probable catastrophic accident scenarios for tandem vehicles, like the shuttle, because of the foreshortened geometrical distance between the crew cabin and the fuel sources. The bottom line is that it is very likely that, even if there had been a conventional crew escape system installed on the shuttle prior to 51L, **it would not have been activated in time to prevent crew cabin involvement in the accident.** That is, the foreshortened distance between the fuel and the crew reduces the time of response for a conventional escape system to be activated to remove the crew from the accident effects. This problem is true not only for the shuttle, but also for all tandem arrangements whether external, as in the case of the shuttle, or internal, as in the case of some single stage to orbit (SSTO) designs.

8. SOLID ROCKET MOTORS (SRMS)** AND CREW SAFETY

As is readily apparent the shuttle uses solid rockets to provide thrust augmentation assistance throughout the initial portions of its ascent trajectory. These SRBs have been much maligned because they cannot be shutdown once started and because they are perceived to be less reliable than their liquid fueled alternatives. Although the actual reliability comparison might be open to question [7] there are two things which are less controversial about SRMs. These are, that they were the cause of the 51L accident and that they still represent, on a per second of burn time basis, the largest risk contributor to the shuttle during ascent. Despite these undeniable facts there are other facts which have recently been brought to light which may reveal hidden crew safety benefits of SRMs. That is, even though they are the most risk intense contributor to launch their failure history indicates an extremely high percentage of failures will not lead to SRM explosion. Rather there is a very high probability that if a failure occurs in an SRM it will likely be a burn through of some sort [8] and lead to a thrust imbalance, as in the 51L case. This is an extremely important piece of information from the crew safety perspective because it indicates, if shown to be true, that even though the SRMs would be the most likely initiator of a future

** The correct shuttle program terminology is reusable solid rocket motors (RSRMs). These are combined with the solid rocket booster (SRB) elements to form what some in the program call the integrated solid rocket motors (ISRMs) which is what is referred to in this paper simply as the SRMs.

accident the solid rocket fuel would be unlikely to be involved in the accident development. This in turn implies that any fuel involvement would be likely to involve only the hydrogen fuel, as was the case with 51L.

9. THE BENIGN NATURE OF HYDROGEN

This previous point is important due to the now suspected behavior of cryogenic hydrogen fuel. It has been known since the work of the post accident explosion working group report [9] that the Challenger did not explode. In fact what was observed was a deflagration or burning of the hydrogen fuel. Whatever initial flash explosion did occur was limited to the tank interstage and even there to at most tens of pounds of TNT equivalent as opposed to the thousands of equivalent pounds available. This implies that there was no explosive pressure wave that might have caused the breakup of the Challenger, but rather the vehicle stack and then the orbiter itself broke up as a result of aerodynamic forces and not explosive forces. That is it broke up with only the normal ascent aerodynamic forces impacting the crew cabin. This fact was well known since post accident investigation showed that the crew cabin survived the accident environment essentially unscathed until ocean impact.

As important as these facts are two further pieces of information which have only recently been made available are even more important from the perspective of crew safety. First recent testing performed at the NASA WSTF [10] seems to indicate that the maximum explosive effect of any hydrogen fuel explosion is limited to only the volume contained very close to the maximum contact surface area. In the case of the shuttle this implies that it is most probable that any hydrogen fuel explosive event would be limited in effect to a maximum of 500 pounds of TNT equivalent. Further, recent detailed explosion simulations [11] have implied that the loads placed upon the shuttle crew module from even 500 lbs. of TNT would produce loads at the crew module, for at least some likely initiator points, which are less than the normal aerodynamic loads from 10 sec. after launch through the separation of the SRMs.

All this recently developed information seems to suggest is that taking the rapid separation escape approach to crew safety, which is applicable for in-line vehicles, a military jet aircraft; it might be incorrect for tandem vehicles. In the tandem case, at least for hydrogen involved scenarios^{***}, designing an accident survivable crew cabin, as is the case for helicopters and racecars, might be a better approach.

10. FORECASTING THE SAFETY AND RELIABILITY OF FUTURE SYSTEMS

It is difficult enough to forecast the safety and reliability risk of current launcher systems from their historical data set, because of the small number of samples and the developing nature of the designs in most cases. It is even more difficult to attempt to forecast the future safety and reliability risk potential of future designs. This is because the tolerance

^{***} This might not be the case for hydrocarbon fueled reusable boosters, for example, because of the different behavior of hydrocarbon fuel.

uncertainties are larger due to the lack of any design specific experience. However on existing systems^{*}, as in the case of the shuttle, it is also unlikely that there would be enough flight level specific data available to be able to forecast the future reliability within the required level of uncertainty. Even if there were, the continuing design development characteristic of such systems would make such a forecast on this data alone questionable. In both existing and future space systems extensions to the data set are required. The question is how to do so, so as to allow for reasonable forecasts to be made.

The first thing required to be understood is that there is always some "heritage" to new designs. Further although new design can definitely eliminate failure mechanisms and thereby increase the overall level of reliability it is unlikely to eliminate the core risk drivers as long as the same physical principles are applied. In the case of launch vehicles this implies that unless a dramatically different method of launch physics it used, such as external propulsion, the propulsion systems will still be the drivers of the risk and therefore that ascent portion of the mission is likely to be the most risky. If the launch vehicle concept includes reuse then the descent risk drivers are thermal protection and landing unless the capability for go around is added.

For these reasons the rational way to approach forecasting the safety and reliability risk of future designs is to identify the core risk drivers from the heritage designs and to identify why and where the new design includes features that are intended to reduce the risk contribution of these drivers. The more revolutionary the design feature the more of the risk it may be considered to potentially abate, but with correspondingly more uncertainty in its performance because of its deviation from the existing knowledge base. These features may exist at a high level in the design, such as the addition of engines to allow a more robust capability of mission continuance with an engine failure. They might also exist at the mid level such as the incorporation of a more benign engine cycle to decrease the stress on the engine. They might also exist at lower levels for particular lower level contributors such as the use of hydrostatic rather than friction bearings. The design feature changes can occur at all levels but the key is to capture those that offer a significant potential for risk reduction.

Once this is completed a risk model of the design concept is developed from the risk model of the current concept in this case either from the shuttle or from the heritage expendable launchers. The model is constructed taking care to apply the portions of the design which are similar, and to represent the risk beneficial deviations in a manner which highlights their potential for risk reduction at the appropriate level of the model. The mean estimate of the failure performance expected above the current implementation should be estimated in proportion to the level of risk intend to be captured by the design change. As was mentioned the greater potential for risk capture usually implies a more revolutionary design with less reliance on heritage and therefore with greater uncertainty of achievement. This greater uncertainty in safety or reliability performance should also be proportionally

^{*} The Soyuz/Molinya system might be considered an exception.

represented according to the level of current technology in the applied are. One technique to represent this uncertainty is to use a standard distribution with an increasing level of uncertainty about the promised enhanced mean according to the level of readiness of the incorporated technology, that is its Technological Readiness Level (TRL) [12].

The potential gain in safety and reliability is then estimated and the associated uncertainty in attainment represented in the portions of the model, which have been modified. These portions are in turn integrated into the overall design model and the risk and uncertainty evaluated as would be done for the evaluation of the risk of an operational design. The resulting evaluation "discounts", in the same sense as in financial discounting, the promised benefits of new designs according to the uncertainty that they will be forthcoming. This discounting effect more fairly treats comparisons of risks among existing designs with known, but perhaps unacceptable risk, and future designs which promise more acceptable risks, but with less certainty of attainment.

11. IMPLICATIONS FOR FUTURE LAUNCHERS

When this approach is applied across the board in the comparison of the crew safety and reliability risks for a spectrum of designs a result such as that represented in Figure 4 would be expected (Note: the figure is notional but reasonable). The reader will note that the areas surrounding the mean points are the 50%iles of the uncertainty distributions. Other percentiles would have similar symmetry but would encompass larger or smaller areas as appropriate. The expected mean point around which the uncertainty is estimated is in the geometric center of each represented area.

1 in 10,000

50th Percentile Probability

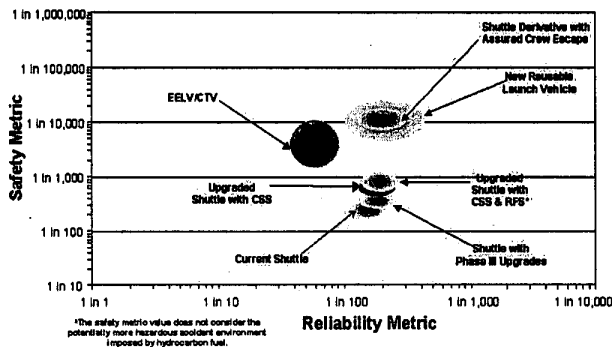


Figure 4: Safety/Reliability with Uncertainties
50th Percentile Probability

Beginning with the current shuttle the uncertainty in both the safety and reliability risk is the same because both are driven by the uncertainty in its inherent reliability. This proportional uncertainty would hold for an upgraded shuttle with a slight increase (since the upgrades would presumably be at a high TRL) in uncertainty expected due to the fact that the impact of these upgrades is not known precisely.

In the case of a new US manned expendable launcher the reliability is estimated at the level demonstrated by the Soyuz and the uncertainty in reliability is estimated to be about what it would be for the upgraded shuttle because again

it is known technology. The mean estimate in the safety level is given by the discussed potential 90% capability of the escape system and the uncertainty associated with incorporating such a system into a new launcher considering that both the launcher and the escape system are relatively mature technologies.

Advanced reusable launch systems cover many types. All promise greater reliability because of more robust propulsion systems and all incorporate crew return systems, but these systems have not been even designed no less tested. Further for most, if not all, of these reusable concepts the crew module is geometrically tandem to the fuel tanks so the in-line escape concept would not be applicable. Still the incorporation of integral escape suggests a mean safety potential well above the current shuttle and in line with that of the expendable. As far as reliability is concerned, the robustness of the propulsion systems in most proposed designs, and the more benign engine design characteristics suggest mean reliability levels well above the promise of the new expendables as well as proportionate gains above the shuttle. However the revolutionary nature of these designs suggest that there is considerable uncertainty in the achievement of both the promised safety and reliability gains. In fact there is a significant probability that these designs would end up less reliable than the current shuttle even though there is little probability that they would be less safe. Compared to a new expendable there is a significant probability that it would be both less safe and less reliable because of the relative maturity of the technology in both cases.

This comparison suggests an interesting alternative design approach. That is what if a design could be proposed that would have the reliability and certainty in the reliability of the shuttle but would include an integrated crew return system to enhance crew safety. In this case this hypothetical design would promise the reliability of the shuttle and the crew safety level of the new vehicles. Whether such a vehicle could be built is open to question, and if it could, because of its tandem nature, it would still have the uncertainty associated with unproven survival and return concept. However if it could be built it would offer crew safety levels consistent with those suggested by the NASA Administrator, by NASA JSC and by this paper at perhaps a fraction of the cost of a new vehicle.

12. CONCLUSION

Forecasting the future is can never be considered as an exact science. However if the tools of risk assessment are used to develop such forecasts in a systematic and coherent fashion, and if they properly take into account the uncertainty in the achievement of the promised benefits of future designs. Then, not only can the decision make process made more tractable and orderly, but also design alternatives, which are not immediately apparent, may come to light.

13. ACKNOWLEDGEMENT

The author wishes to acknowledge the contribution of Mr. Gaspare Maggio of SAIC who pioneered the application of the heritage based evaluation approach in its application to future propulsion system designs. He also wishes to acknowledge the contribution of Mr. Michael Gaunce of NASA ARC, who persistence in the requirement for an orderly process for safety and reliability evaluation caused the advancement of the approach discussed here, and Mr. Steven Cook, of NASA MSFC, who led the charge in applying the approach in a detailed evaluation of space launch systems.

14. REFERENCES

- [1] Goldin, D.S., "AIAA Luncheon Speech Transcript", 3 May 1999.
- [2] Jenkins, M., "Human-Rating Requirements" JSC-28354, June 1998.
- [3] Moody, J.W., "Reliability of Ballistic Missiles and Space Vehicles", Working Paper, Reliability Office, George C. Marshall Space Center, Huntsville, AL, 15 April 1961.
- [4] Fragola, J.R., and Perkins, David R., "Propulsion Reliability: An Historical Perspective", AIAA/SAE Joint Propulsion Conference, Monterey, CA, 10-13 July 1989.
- [5] Fragola, J.R., "Risk Management in US Manned Spacecraft: From *Apollo* to *Alpha* and Beyond", Proceedings of ESA 1996 Product Assurance Symposium and Software Product Assurance Workshop, Noordwijk, The Netherlands, 19-22 March 1996.
- [6] Louviere, A.J. and O'Connor, B.D., "STS Crew Egress and Escape Study, Volume 1 - Analysis", JSC 22275, August 1986.
- [7] Fragola, J.R., "Launch System Reliability: A Second Look", *Aerospace America*, November 1991.
- [8] Gunn, C., "Failures are Not an Option", *Launchspace*, January/February 1999.
- [9] Reynolds, J.R., "Explosion Working Group Final Report of the NASA/DOE/INSRP STS 51-L and Titan 34D-9", KSC-00232, June 1989.
- [10] Bunker, R.; Starritt, L.; Flint, Q.; Eck, M. and Taylor J., "Joint NASA/NASDA Hydrogen/Oxygen Vertical Impact (HOVI) Test Program", TR-820-001R, 31 March 1995.
- [11] Baum, J., unpublished data, SAIC, April 1999.
- [12] Fragola, J.R. and Maggio G., "Application of Probabilistic Program Planning to an Advanced Radiographic Hydrodynamics Test Facility", European Space Agency Risk Management Workshop, 30 March - 2 April 1998, ESA/ESTEC Noordwijk, The Netherlands.

BIOGRAPHIES

Joseph R. Fragola
Science Applications International Corporation
7 West 36th Street, 10th Floor
New York, NY 10018

Internet (e-mail) joseph.fragola@cwixmail.com

Mr. Fragola has almost 30 years of experience working in reliability and risk technology in both the aerospace and nuclear industries and for other industries in the US and throughout Europe. Mr. Fragola received his B.S. and M.S. degrees from the Polytechnic Institute of New York and in the past has worked for Grumman Aerospace Corporation, and IEEE at their headquarters in New York. He is currently a Vice President and manager of the Advanced Technology Division at SAIC. He has published almost 50 papers and two books. He has been awarded the P.K. McElroy RAMS best paper award, the IEEE Region I award, and has been named an IEEE Fellow for his contributions to the theory and practice of risk, safety, and reliability. In March of 1996 he was awarded a SAIC publication prize for peer-reviewed journal publication.