



Academy Sharing Knowledge

ask

The NASA Source for Project Management and Engineering Excellence | APPEL

SPRING | 2009



Inside

IS SOFTWARE BROKEN?

FILLING THE KNOWLEDGE GAPS

PLAN, TRAIN, AND FLY

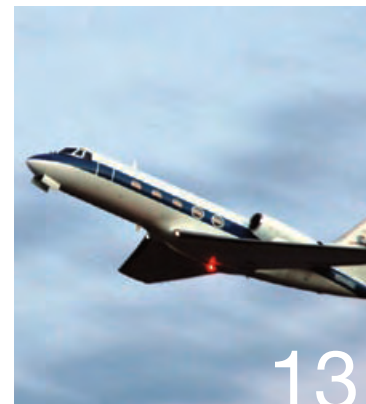


Photo Credit: NASA

ON THE COVER

The Hubble Space Telescope (HST) is shown sporting new and modified solar arrays stowed against its barrel after its first servicing mission in 1993. An astronaut begins other repairs of the HST while perched atop a foot restraint on shuttle *Endeavour*'s remote manipulator system arm. NASA is preparing for its final HST servicing mission in 2009, when astronauts will install two new instruments, repair two inactive ones, and perform component replacements to keep the telescope functioning at least into 2014.

Contents



DEPARTMENTS

- 3**
In This Issue
BY DON COHEN
- 4**
From the APPEL Director
BY ED HOFFMAN
- 58**
The Knowledge Notebook
BY LAURENCE PRUSAK
- 60**
ASK Interactive

INSIGHTS

- 10**
Integrating Risk and Knowledge Management for the Exploration Systems Mission Directorate
BY DAVE LENGYEL The ESMD uses a variety of media and methods to make important knowledge accessible.
- 17**
Interview with Kenneth Szalai
BY DON COHEN The former director of the Dryden Flight Research Center explains why flight programs matter.
- 22**
Is Software Broken?
BY STEVE JOLLY Increasingly complex and central to spacecraft design, software makes new demands on systems engineers.
- 30**
Leave No Stone Unturned, Ideas for Innovations Are Everywhere
BY COLIN ANGLE To keep innovating, iRobot listens to customers, partners with other organizations, and welcomes new ideas from its own employees.
- 43**
Viewpoint: Building a National Capability for On-Orbit Servicing
BY FRANK J. CEPOLLINA AND JILL MCGUIRE The authors argue that NASA and the Department of Defense should collaborate on satellite-servicing technology.
- 50**
Things I Learned on My Way to Mars
BY ANDREW CHAIKIN The author of *A Passion for Mars* discusses the perils and rewards of exploring the red planet.

STORIES

- 6**
From Generation to Generation: Filling the Knowledge Gaps
BY JIM HODGES Langley engineers working on an Ares test vehicle draw on the expertise of their Apollo-era counterparts.
- 13**
Plan, Train, and Fly: Mission Operations from Apollo to Shuttle
BY JOHN O'NEILL The technology of preparing for human space flight has changed over the years, but the basic principles of careful planning and exhaustive training remain the same.
- 26**
Project Lessons from Code Breakers and Code Makers
BY JOHN EMOND Faced with the life-and-death intelligence issues of WWII, the Allies enlisted the talents of surprisingly diverse individuals; NASA's challenges call for a similar approach.
- 34**
Mars Science Laboratory: Integrating Science and Engineering Teams
BY ASHWIN R. VASAVADA Close collaboration between engineers and scientists is helping MSL meet its ambitious goals.
- 38**
ESA, NASA, and the International Space Station
BY ALAN THIRKETTLE International cooperation works best when the partners make essential, interdependent contributions to clear common objectives.
- 46**
Building the Team: The Ares I-X Upper-Stage Simulator
BY MATTHEW KOHUT An in-house development project at Glenn calls for team members who are flexible and ready to learn by doing.
- 54**
Apollo 12: A Detective Story
BY GENE MEIERAN How one contractor team helped decide that Apollo 12 was safe to fly.

ask

Issue No. 34

Staff

APPEL DIRECTOR AND PUBLISHER

Dr. Edward Hoffman
ehoffman@nasa.gov

EDITOR-IN-CHIEF

Laurence Prusak
lprusak@msn.com

MANAGING EDITOR

Don Cohen
doncohen@rcn.com

EDITOR

Kerry Ellis
kerry.ellis@asrcms.com

CONTRIBUTING EDITOR

Matt Kohut
mattkohut@infactcommunications.com

SENIOR KNOWLEDGE SHARING CONSULTANT

Jon Boyle
jon.boyle@asrcms.com

KNOWLEDGE SHARING ANALYSTS

Ben Bruneau
ben.bruneau@asrcms.com

Katherine Thomas
katherine.thomas@asrcms.com

Mai Ebert
mai.ebert@asrcms.com

APPEL PROGRAM MANAGER

Yvonne Massaquoi
yvonne.massaquoi@asrcms.com

DESIGN

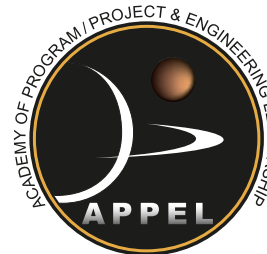
Hirshorn Zuckerman Design Group, Inc.
www.hzdg.com

PRINTING SPECIALIST

Hanta Ralay
hanta.ralay 1@nasa.gov

PRINTING

GraphTec



The Academy of Program/Project and Engineering Leadership (APPEL) and *ASK Magazine* help NASA managers and project teams accomplish today's missions and meet tomorrow's challenges by sponsoring knowledge-sharing events and publications, providing performance enhancement services and tools, supporting career development programs, and creating opportunities for project management and engineering collaboration with universities, professional associations, industry partners, and other government agencies.

ASK Magazine grew out of the Academy and its Knowledge Sharing Initiative, designed for program/project managers and engineers to share expertise and lessons learned with fellow practitioners across the Agency. Reflecting the Academy's responsibility for project management and engineering development and the challenges of NASA's new mission, *ASK* includes articles about meeting the technical and managerial demands of complex projects, as well as insights into organizational knowledge, learning, collaboration, performance measurement and evaluation, and scheduling. We at APPEL Knowledge Sharing believe that stories recounting the real-life experiences of practitioners communicate important practical wisdom and best practices that readers can apply to their own projects and environments. By telling their stories, NASA managers, scientists, and engineers share valuable experience-based knowledge and foster a community of reflective practitioners. The stories that appear in *ASK* are written by the "best of the best" project managers and engineers, primarily from NASA, but also from other government agencies, academia, and industry. Who better than a project manager or engineer to help a colleague address a critical issue on a project? Big projects, small projects—they're all here in *ASK*.

You can help *ASK* provide the stories you need and want by letting our editors know what you think about what you read here and by sharing your own stories. To submit stories or ask questions about editorial policy, contact Don Cohen, Managing Editor, doncohen@rcn.com, 781-860-5270.

For inquiries about APPEL Knowledge Sharing programs and products, please contact Katherine Thomas, ASRC Management Services, 6303 Ivy Lane, Suite 130, Greenbelt, MD 20770; katherine.thomas@asrcms.com; 301-793-9973.

To subscribe to *ASK*, please send your full name and preferred mailing address (including mail stop, if applicable) to ASKmagazine@asrcms.com.

In This Issue



To accomplish its mission of developing new launch vehicles and manned spacecraft, NASA must excel at learning. We need to learn lessons from the extraordinary technical advances that culminated in the moon landings of the sixties. We must share what we know effectively within and between projects. And, as we work on programs that will establish our space exploration capabilities for decades to come, we are obliged to make sure that knowledge we are developing now will be available to the engineers, scientists, and managers who will face new challenges in the future. Many of the articles in this issue of ASK consider these learning issues.

In “The Knowledge Notebook,” Laurence Prusak reflects on how much of our knowledge we owe to those who came before us, and Jim Hodges offers a vivid contemporary example of learning from the past in “From Generation to Generation.” The Langley team developing an Ares I-X test vehicle turned to the people who carried out a similar project for the Saturn V in the sixties because documents from that period did not say enough about the *how* and *why* of those earlier tests and test results. The retirees who did that work provided detail that brought the documents to life as useful guides to the Ares team.

Dave Lengyel’s update on the Exploration Systems Mission Directorate (ESMD) efforts to capture and share essential knowledge (“Integrating Risk and Knowledge Management”) looks at the same issue in relation to recent and current work. By focusing on knowledge about recognized risks, the ESMD ensures that it is preserving expertise that matters, but the knowledge-based risk team recognizes that without context—the *how* and *why* of decisions and technical information—current and future project teams will not be able to put that knowledge to use. They are using case studies, group discussions, wikis, and other approaches to provide that context.

Learning from people who have been there before is indispensable when there are difficult tasks to be accomplished, but much of what they learned and what current team members learn comes from doing the work. Several articles here are about the irreplaceable value of learning by doing. In the interview, Kenneth Szalai talks about what can be learned from flight programs, which he describes as “the truth serum and lie detector of what is possible.” John O’Neill’s history of mission operations (“Plan, Train, and Fly”), Matthew Kohut’s discussion of building a team for in-house development at Glenn, and Alan Thirkettle’s “ESA, NASA, and the International Space Station” highlight knowledge that can only be gained from experience.

How much you learn depends in part on whom you work with. “Mars Science Laboratory: Integrating Science and Engineering Teams” by Ashwin R. Vasavada, “Project Lessons from Code Breakers and Code Makers” by John Emond, and Frank J. Cepollina and Jill McGuire’s “Building a National Capability for On-Orbit Servicing” all argue for the value of bringing together diverse expertise. Colin Angle’s “Leave No Stone Unturned” shows that the new learning we call innovation comes from being open to as many sources of knowledge as possible.

Finally, Ed Hoffman’s “From the APPEL Director” column about attending a Flight Readiness Review makes powerful points about the conditions that make learning and sound decision-making possible. Without trust, openness, inclusiveness, and respect, learning doesn’t happen.

Don Cohen
Managing Editor

From the APPEL Director

Good Team Design

BY ED HOFFMAN



The decision to launch a shuttle brings together complex issues of many kinds—issues of engineering, safety, systems, technology, time, pressure, and people. All these elements are important and any of them can loom large. For me, the team dynamics on display in the long, intense Flight Readiness Review meeting are the most stunning.

In classrooms and team-building sessions, someone usually asks, “Why do we need teams?” The answer becomes self-evident during a Flight Readiness Review. The review is filled with project and engineering complexity; every decision is critical, with an impact on mission success and crew safety. The importance of effective teamwork becomes obvious in this situation. Good team design is essential, not a vaguely desirable “extra.” It makes the difference between success and tragedy. Sitting in on the Shuttle Flight Readiness Review, I saw many of the factors that go into good team design in action.

Context and setting matter. The entire team understands the importance of the decisions they make—to mission success and the lives of the crew. The setting supports and emphasizes their joint responsibility. The Flight Readiness Review takes place in a large, open room with a design that focuses attention on discussion and visual evidence. The primary decision makers sit at a long, central table. Anything displayed on the three large screens in front of the table can be seen from any seat in the room. Surrounding the central table are rows of seats in all directions. Everyone in the room can be seen and heard by all. At first glance, the seating may appear haphazard, but closer inspection shows it has the precision of an ant army. Special teams

are organized in different seating areas: teams from the centers, teams from engineering, teams from the program, teams from safety. Experts are gathered and organized to ensure that every relevant question will be posed and answered and every answer will be thoughtfully considered. This is not a place for hiding out or holding back. A big video eye in front records everything.

Size and organization depend on the task. The Shuttle Flight Readiness Review goes against the literature that advises minimizing the number of people on a team. There are more than one hundred people in the room, all of whom contribute at different points. The size of the team reflects the range of technical expertise needed and the interdependence of the systems they understand. There are no simple or isolated decisions in the review. Every decision has an impact on other systems. During the discussion about recently discovered cracks in hydrogen valves, solutions must be understood in the framework of the larger system. A seemingly reasonable solution can cause disaster if the systemwide impacts are not clearly understood and extensively tested. Schedule decisions affect numerous goals and multiple missions. For instance, a decision to reduce risk by delaying a shuttle launch creates additional risk on the International Space Station. The potential problems posed by a team of this size are reduced by organizing members into functional groups, small “communities” of experts that function as teams within the larger team.

These varied teams and sheer number of experts present provide the *diversity of ideas* essential to the complex, interdependent issues involved in Flight Readiness Review decisions. The collective

knowledge, experience, and cross-discipline wisdom are truly amazing and make it a joy and a privilege to watch the team in action. Decision making tends to take place either in the group as a whole or among the communities that work together under the broad headings of engineering, safety, and program.

Leaders—there are several leaders of the review process—must *balance* diversity of individual perspective and collective direction. They must encourage conflict and promote consensus to the appropriate degree at the appropriate times. Analysis and learning must lead to action, but the need to act cannot be allowed to undermine careful analysis.

It is impossible to overstate the amount of skill that goes into making this process work. The necessary expertise is not simply technical, because the right technical answers can only be arrived at with the help of strong project management and interpersonal skills. The project management perspective understands the implications for project cost, schedule, performance, and planning of every technical decision. And the collaboration that defines a successful review would be impossible without the interpersonal skills that build a foundation of trust, openness, inclusion, and respect.

Good team design includes constructive feedback that helps the team evaluate what it has done and adapt to new demands. This is where the relationship between successful leadership and the whole team is most evident. At the end of every critical phase of the shuttle review, the team is asked to provide thoughts. The leaders deliberately pause, visually scanning the room to encourage feedback. Everyone at the central table is asked for specific comments. Then industry leaders speak, taking responsibility for the elements of the system they are accountable for.

Complex, important projects make great demands on leaders and teams. Decision making under the pressures of mission aims, schedule, and life-or-death safety issues is stressful. That stress can help inspire high performance or push a team toward failure. Good team design that brings the right people

and the right processes together in the right setting is essential to ensuring the best possible decisions in such demanding situations. The STS-119 Shuttle Flight Readiness Review is a prime example of how good team design works and how it contributes to a successful outcome. ●

... THE COLLABORATION THAT DEFINES
A SUCCESSFUL REVIEW WOULD BE
IMPOSSIBLE WITHOUT THE INTERPERSONAL
SKILLS THAT BUILD A FOUNDATION OF TRUST,
OPENNESS, INCLUSION, AND RESPECT.

From Generation

An engineer stands next to a 3 percent scale Saturn V model in the Transonic Dynamics Tunnel at NASA's Langley Research Center in 1966.

Photo Credit: NASA



to Generation:

FILLING
THE
KNOWLEDGE
GAPS

BY JIM HODGES



In 2006, when the staff of the Aeroelasticity Branch at NASA's Langley Research Center learned that it would test ground wind loads for the Ares I-X launch test vehicle, Donald Keller and Thomas Ivanco went in search of history.

Thomas Ivanco prepares a model of Ares I X for testing in the Transonic Dynamics Tunnel at NASA's Langley Research Center.

Photo Credit: NASA/Sean Smith

The branch had performed similar tests in the Transonic Dynamics Tunnel (TDT) in Hampton, Virginia, for the Saturn V that was used to carry Apollo aloft but had conducted few such tests since. After all, NASA hasn't built a human-rated launch vehicle designed to send astronauts to the moon and beyond since Saturn V in the 1960s.

What Keller found was discouraging. "A lot of the reports summarized the results of the tests, but there wasn't a lot of detail about steps that were taken from model concept through fabrication and testing, and how things were done or even why, in some instances," said Keller, who—with Ivanco—oversaw TDT testing and preparation of the 4 percent-scale Ares I-X ground wind loads model during fall of 2008.

The same lack of documented knowledge hampered interpretation of the Ares I-X test data for Systems Engineering and Integration, which ordered the tests. That was primarily Ivanco's job, and he struggled to understand what to do with the data from the tests for his report in March. The real difficulty was in translating the model data to a full-scale vehicle.

"Because of the limitations of the data systems in the sixties, a lot of data acquisition involved somebody writing something in a notebook," Ivanco said. "They were sometimes reading an analog gauge and recording it by hand. Portions of that data ended up in the final reports, but where did those notebooks go?" Many of the tests also recorded data on analog tapes, which also were long gone.

While musing over those knowledge gaps, Keller said, "We started seeing things that we didn't understand as we got more seriously into the project." If the answer wasn't on paper, they figured, maybe it was with the people who wrote the paper. "We realized that anybody we could talk to would give us more than we had," Keller added.

Bill Reed and Bob Doggett, both retired from NASA, come to Langley once a week to work on archiving reports, pictures, and data for the TDT. More important to Keller and Ivanco was that Reed and Doggett both worked at the Aeroelasticity Branch during the days of Mercury, Jupiter, Gemini, and Apollo. And, yes, they had read a gauge and recorded data in a notebook while testing Saturn and Apollo. More to the point, they used strain gauges linked to an oscilloscope that measured loads and motion in two directions.

"We discovered that if you put these two signals on an oscilloscope and just let it run, it runs a pattern," said Reed, who was deputy branch manager during the Apollo days. "You'd see an envelope that defines the outer bounds of the loads relative to the wind.

"We would take a piece of tracing paper and put it on the screen and take a pencil and outline this figure. Then we had the loads. Then we graduated from the tracing paper to a Polaroid camera, and we'd just take a picture of it."

A computer that generates colorful charts does that work today, but having the original test background offered insight into what the computer tells the researchers.

"Today engineers are looking at displays of exactly the same sort of data that was acquired and examined all those years ago," said Doggett. "Conceptually, the data are reduced in the same fashion, but the devil is in the details—going from hand tracings to Polaroid photographs and finally to today's computers." Added Keller, "People like Bill Reed and Bob Doggett were good to talk to for more detail than we could glean out of the reports."

And so the search for history broadened. A net was cast for retired NASA employees who could offer insight into how testing was done forty years ago. Those retirees still around came forth eagerly, and four contributed insight into how to scale the test articles, how to conduct the tests themselves, and what to do with the data. "We wanted to help," said Reed. "We didn't want them to have to reinvent the wheel."

Questions came as the project evolved. Take damping, for example, and how it affects oscillations generated by the winds Ares I-X will find on the launchpad at Kennedy Space Center. Pictures from a past report showed a viscous damper that used fluid and weights to help determine how to counter ground wind loads.

"I told him what I knew," Reed said of Keller, then met the query one better. A notebook from his attic produced figures that hadn't been in the test report all those years ago. Those figures could be used to help check formulas today.

A damper was built for the Ares wind-tunnel model, but "it didn't turn out as well as they had hoped," Reed said. Back to Reed's attic: he had a scale model of a damper used on a long-ago test.

"It gave us an idea for a follow-up design for a damper," Keller said. "It led us down the road to what we eventually used."

Said Reed, “Then they developed a more advanced kind ...” Finished Doggett, laughing, “... of super-duper damping.”

Reed pointed at a picture. “This was crude,” he said, “but it got us by in the Saturn program to get the damping needed to see how it affected the vehicle response to winds, and it ultimately was included in the full-scale hardware at the Cape.”

The success of Saturn V and Apollo is a point Reed and Doggett use in talking about the more limited work done four decades ago. Keller and Ivanco tested Ares I-X with thirty-one wind-speed readings, each about 15 seconds long, 500 samples per second. All were taken from seventy-two different wind

... KELLER SAID, “WE STARTED SEEING THINGS THAT WE DIDN’T UNDERSTAND AS WE GOT MORE SERIOUSLY INTO THE PROJECT.” IF THE ANSWER WASN’T ON PAPER, THEY FIGURED, MAYBE IT WAS WITH THE PEOPLE WHO WROTE THE PAPER.

angles, a tedious process. Yes, Doggett and Reed will say, they used fewer measurements with fewer sensors, but Saturn flew and Apollo landed on the moon and returned.

As the testing on Ares I-X evolved, “there were things that were just hard to explain,” Ivanco said. Apparently, they were just as hard to explain forty years ago.

“They were very vague about it in the reports,” Ivanco said. “You would read the reports, and they would mention, for instance, that they applied a piece of tape to the model and the dynamic loads changed by an order of magnitude. But that’s all they would say. They wouldn’t say which way or offer theories as to why or describe the test conditions they were doing when they put the tape on.”

Some answers were provided by people who were there at the time. Many answers simply were no longer available because the people were no longer available. And some questions were

never answered at all. And so, if the wheel wasn’t reinvented, it was certainly modified to fit 2008–2009 test parameters and those likely to be used in the future.

“One of the reasons the reports that Tom and I are writing right now are vastly more detailed than what they did before is that we do not want future researchers to have to try and figure out what steps we took and why we took them,” Keller said. “But again, a lot of these past reports were probably supporting some projects, and they had to get the bottom line out fairly quickly. They summarized what they needed to summarize, but nobody ever went back to provide the details.”

Keller and Ivanco will leave a legacy of data that Reed, Doggett, and their co-workers could not. “Nowadays you can store the data on DVDs or hard drives, something that will last longer,” Keller said. “Back then, they had magnetic tapes that were thrown out or broke down or whatever. You can store DVDs on a bookshelf, but those tapes would have filled cabinets.”

Reed understands. “That’s why we have this archives effort here,” he said.

Both generations lament the knowledge gap between NASA efforts on Apollo–Saturn and Ares and applaud what’s happening now.

“I used to wish I had been here in the heydays of the Apollo era,” said Ivanco, who wasn’t yet born when Neil Armstrong set foot on the moon. “I used to think that must have been something, working for NASA back then, going to the moon, pioneering.”

It was, said the pioneers.

“And it’s a good thing we aren’t doing this ten years from now,” Ivanco added. “If this program happened ten years from now, fifteen years from now, we wouldn’t have had that resource. All we would have had was the reports.”

Added Doggett with a laugh, “And if they want more, they’d better hurry.” ●

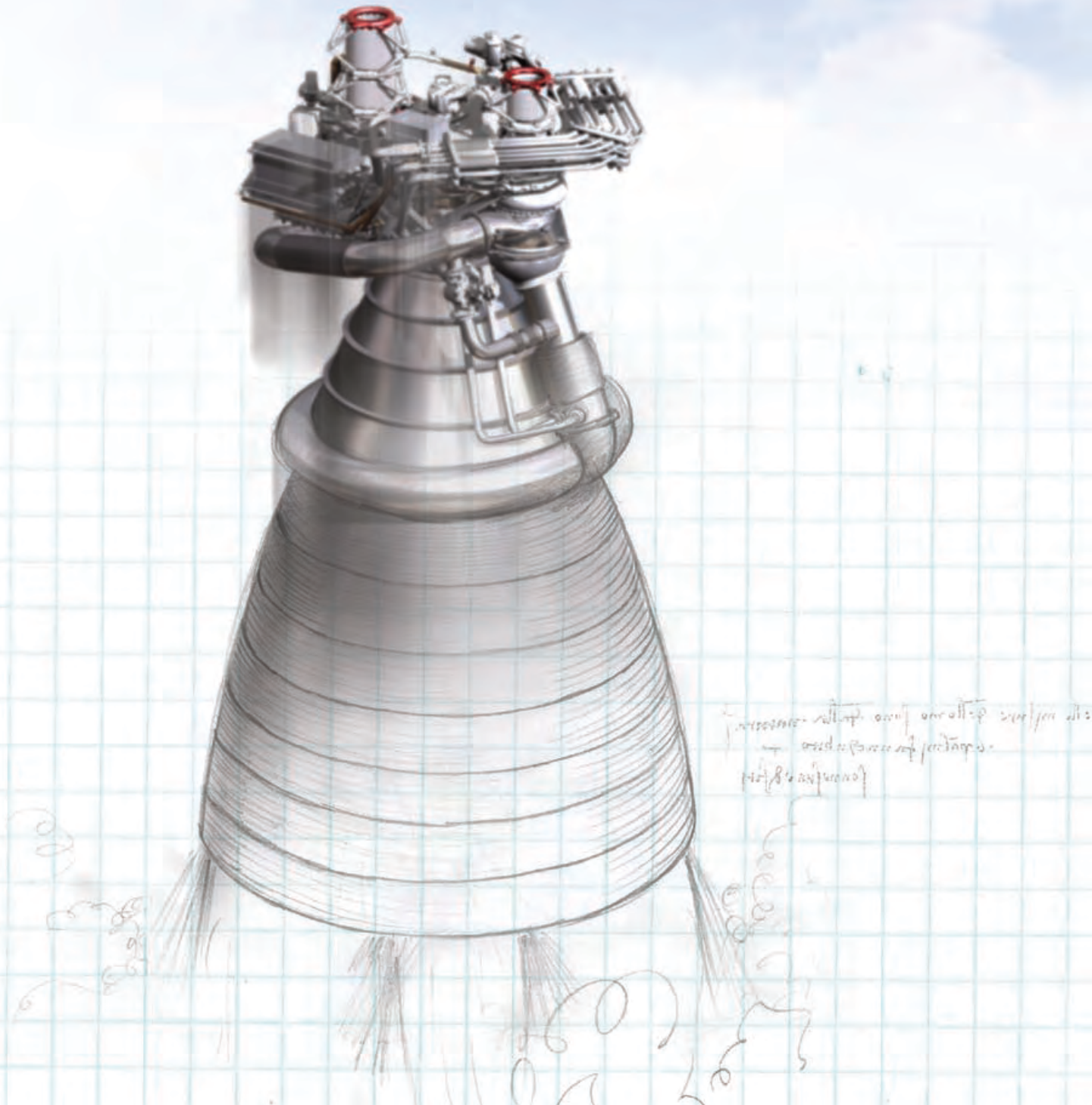
Former *Los Angeles Times* reporter **JIM HODGES** is managing editor/senior writer of the *Researcher News* at NASA’s Langley Research Center.



Photo Credit: NASA/Sean Smith

Integrating Risk and Knowledge Management for the Exploration Systems Mission Directorate

BY DAVE LENGYEL



As NASA undertakes ambitious new programs with a new generation of engineers and managers, it is more important than ever to make sure that valuable experienced-based knowledge gets passed from project to project and from an older generation to a new one. Many organizations try to solve this kind of problem with lessons-learned databases, which, for a variety of reasons, seldom live up to expectations. The databases typically are filled with undifferentiated information, good and bad, relevant and unimportant, so it's hard to find what matters. The content is often fragmented—text or bullets without adequate context, usually lacking the analysis or synthesis needed to make it useful.

In early 2006, we in the Exploration Systems Mission Directorate (ESMD) decided to generate and share engineering/design, operations, and management best practices by combining pre-existing Continuous Risk Management (CRM) work with knowledge management concepts to infuse relevant lessons and best practices into current activities. Three years of doing this work has taught us to focus not on building a so-called learning organization, but on supporting the performance of work. Helping very busy people accomplish the mission is paramount.

Our effort is based on the assumption that risks highlight potential knowledge gaps that might be mitigated through one or more knowledge management practices or artifacts. These risks serve as cues for collecting knowledge, particularly knowledge of technical or programmatic challenges that might recur. We use a variety of modes—text, video, case studies, and classroom activities—to communicate the knowledge while emphasizing “learning through conversation” rather than an IT-centric approach.

Knowledge-Based Risks

When we first looked at integrating risk and knowledge management, we asked ourselves some simple questions:

- How can we fully exploit the risk database?
- Would appending lessons to risk records be accepted as a more effective means of capturing and transferring knowledge?
- Would the risk database be used as a knowledge base over time—not just a risk repository?

Attempting to answer these questions, we developed the concept of knowledge-based risks, or KBRs. ESMD defines a KBR as a risk record, with associated knowledge artifacts, that provides a storytelling narrative of how the risk was mitigated, including what worked or didn't work. A KBR is a means of

transferring knowledge in a risk context. As key risks are mitigated, particularly risks that are likely to reoccur in other ESMD programs, lessons are captured to answer questions such as, “What was the control and mitigation strategy? Did it work? How were cost, schedule, and technical performance affected?”

The lessons are appended to the risk record by program and project risk managers to help identify new risks and develop better plans for dealing with known risks. When new candidate risks are identified, risk owners use related KBRs and other risks as sources to develop their risk mitigation, analysis, and documentation approach. This provides a tight coupling of CRM with lessons learned. Instead of a “collect, store, and ignore” approach, KBRs form an active collection of lessons learned that are continually reused and updated. This approach enhances our existing risk tool functionality as a “knowledge base.”

Topics of KBRs captured and distributed by our design community in the past year from International Space Station (ISS) and Space Shuttle programs include adequate instrumentation, weather factors for ground processing, corrosion prevention, confusing Problem Reporting and Corrective Action (PRACA) codes, overspecification of design margins, critical math models, and factors of safety. We are currently targeting a number of KBRs related to problem solving, anomaly resolution, and the development of flight rationales. While we may not experience engine cutoff sensor or flow valve issues after initial operational capability construction of the Orion/Ares I stack, we will certainly need to revisit the design, test, and systems engineering practices and principals currently used today to keep the ISS and shuttle flying.

Risk Management Case Studies

While KBRs effectively tell a story about a particular risk, our risk management case studies serve as the ultimate multimedia “lessons learned” experience for ESMD work teams. Our first case addresses the project success story of the Space Shuttle Super Lightweight Tank development. This case was selected

because we are currently and will continue to be challenged by mass-related risks for the heavy-lift booster, lunar lander, and habitat modules. Like other cases, it highlights key transferable aspects of risk management, including the identification and analysis of risks, rigorous mitigation planning, and risk trades.

The proper application of risk management principles examined in these cases can help manage life-cycle costs, development schedules, and risk, resulting in safer and more reliable systems for Constellation and other future programs. The risk and systems engineering communities have embraced this approach; many are working toward an organic case study-teaching expertise. In addition to an instructor-led, small-group delivery format for work teams, case studies are available to a wide ESMD audience on the Internet, providing the opportunity for self-study or moderated, “webinar”-based delivery. Current plans call for integrating both KBRs and case studies into annual CRM training and working with the Academy of Program/Project and Engineering Leadership to incorporate both into their course offerings.

Future Goals and Challenges

So what additional progress and improvements are we working to achieve in the coming year? In the area of continuous risk management, we will continue to integrate CRM with cost and schedule risk analysis and earned value management. We also seek to link CRM with systems engineering and systems safety processes more effectively. The KBR management process is being shifted from ESMD to the program/project levels. This will help us derive value-needed solutions for managing risk and knowledge at the most appropriate level.

Last year we piloted two practices new to ESMD. The first was the innovation methodology known as TRIZ, the acronym of a Russian phrase meaning “the theory of inventor’s problem solving.” The second was the Knowledge Café, a knowledge-generation and -transfer technique using small groups and structured and unstructured brainstorming. TRIZ was used

to generate innovative ideas for packaging loose equipment for lunar missions while reducing waste materials. The café approach should be a useful technique for transferring recently captured knowledge from the ISS and shuttle programs to Constellation.

To exploit our continued growth into Web 2.0 technologies, we have embarked on a project affectionately known as the Risk Wizard, which will provide practitioners with risk-identification checklists, risk-analysis techniques, and access to a wealth of information to aid in building better risk-mitigation plans. Later this year, we will begin to share best practices across our 330-plus wiki-enabled teams through an awards program that recognizes participants’ outstanding achievement in the use of wikis across the directorate. Finally—because many risks stem from a failure of process discipline—we will continue to promote an adaptation of the U.S. Army after-action review, called “Process 2.0.”

The beauty of integrated risk and knowledge management is that we can find the best fit for our KBRs, case studies, and other products and use our network of risk managers as the central nervous system for information flow, a very efficient approach to capturing and transferring knowledge. Through it all, we do not want to forget our most important lesson learned to date, which is to maintain focus on supporting the accomplishment of work. That’s what we’re all about. ●



DAVE LENGYEL is the risk and knowledge management officer for the Exploration Systems Mission Directorate. He has held positions in the Shuttle-Mir and International Space Station programs and is retired from the U.S. Marine Corps.

Plan, Train, and Fly:

MISSION OPERATIONS FROM APOLLO TO SHUTTLE

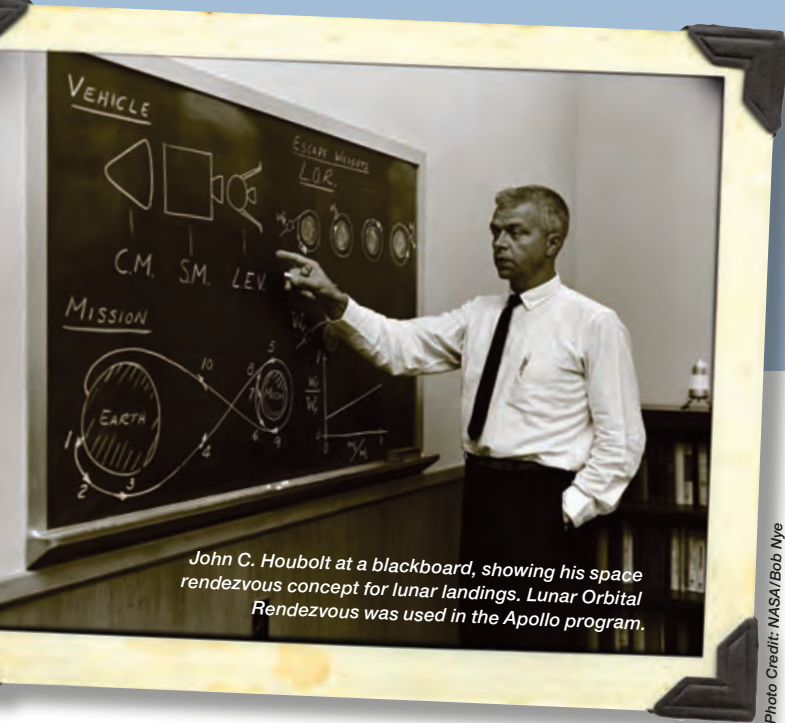
BY JOHN O'NEILL

Personnel at the Mission Operations Directorate at the Johnson Space Center are the final integrators of the planning and execution steps that must occur to get from mission definition and design to flight. Over the years, the technology of some of this essential work has changed, but the general principles and the dedication and skill of those doing it remain the same. A brief look at the history of planning, training, and flying—the three related functions within human space flight mission operations—will make some of the challenges clear and show how we met them in the past and how we meet them today.



On NASA Kennedy Space Center's Shuttle Landing Facility, the Shuttle Training Aircraft (STA) takes to the skies. The STA is a Grumman American Aviation–built Gulf Stream II jet that was modified to simulate an orbiter's cockpit, motion and visual cues, and handling qualities. In flight, the STA duplicates the orbiter's atmospheric descent trajectory from approximately 35,000 ft. altitude to landing on a runway. Because the orbiter is unpowered during reentry and landing, its high-speed glide must be perfectly executed the first time.

Photo Credit: NASA



Planning

President Eisenhower once said, “It has been my experience in a really great crisis that plans were useless but that planning was indispensable.” That is a good guiding principle for the contingency planning that always goes into NASA missions, but, in the complex environment of space, rigorous planning is equally indispensable in accomplishing the defined mission objectives.

Once mission requirements and spacecraft capabilities are in hand, planning essentially begins with the trajectory, navigation, and guidance design. Consider the challenges faced by the Apollo trajectory planners. Activities associated with trajectory control were the largest part of the operational overhead on every Apollo mission. In key mission phases, trajectory control took priority over all other activity and drove the timeline. The trajectory was the framework or skeleton for all subsequent plans and procedures. Before the first moon flights, engineering, trajectory, science, and operations personnel collaborated to develop Design Reference Missions to give “best estimate” guidance to the Apollo mission designers. The Gemini program and its rendezvous missions provided a key source of data that the planners needed to design accurate lunar trajectories.

The flight-planning effort also included developing and refining crew procedures. The flight plan itself was and still is a precise sequence of the interrelated crew and ground support activities. This operations documentation shapes the training that familiarizes crew and controllers with mission procedures and contingencies.

Flight crews also wanted “cue cards”—irregularly shaped cards that fit in available panel space in the crew station—to summarize critical procedures for ready reference. The Apollo 8 ascent cue cards provided an extra bit of excitement in the final launch preparations. The backup crew installed the cards during the last hours of the count so they would be in place for the prime crew. This meant placing the cards with their Velcro backing in position on the mating Velcro on the panel spaces. At that time,

sticky-back Velcro was not yet available; when the cue cards were finalized, an adhesive was used to attach the specially shaped Velcro to the cards. Soon after the backup crew completed the installation, a pad technician discovered that the cards were falling like leaves. The adhesive had failed. In a panic procedure in the Operations and Checkout Building at Kennedy Space Center, the old adhesive was scraped off and fresh adhesive applied. The process took most of the night but was finished in time.

Apollo 11 obviously presented significant new challenges and produced many productive changes. Coordination between the crew and ground support was extremely critical in a mission that included lunar landing and a lunar orbit rendezvous. During the two months before flight, approximately 1,100 changes were made to the flight plan and crew checklists. All those changes were vetted by the crew, the flight controllers, and NASA and contractor engineering personnel. Doing that required streamlined and improved information exchange and led to the development of a formal configuration control process similar to that used for hardware and software today.

Consider the technology of the Apollo era for a moment. Much has been written about the limited capacity of the spacecraft and mission control computers. The lack of word processors also made the careful and accurate updating of operations documentation very tedious. And there were none of the tracking and data relay satellites to provide the full communication coverage that the Space Shuttle and International Space Station enjoy today, and no GPS for navigation support.

Those were the limits on communications resources when the entire operations and engineering force mobilized to deal with the Apollo 13 emergency. As the onboard procedures were reworked, reassembled, and modified, the extensive directions from Mission Control to the crew had to be transmitted totally over the air-to-ground voice loops. This included the famous step-by-step instructions for using tape and covers from the flight plan to adapt a command module lithium hydroxide canister for lunar module use.

Planning, reviewing, and carefully revising the plans have been the cornerstone of NASA’s mission success. Operations planning must start during the requirements phase of a program and be an integral part of the design, development, and testing phases. Two of the most important questions are, “Can the systems be operated in a normal and contingency mode that

FLIGHT CONTROLLERS QUICKLY ANALYZED THE ALARMS AND ADVISED THE CREW THAT THEY WERE NOT SERIOUS AND THE LANDING COULD CONTINUE. THEY KNEW THAT BECAUSE OF THEIR TRAINING EXPERIENCE.

satisfies the mission requirements?” and, “Have the flight crew and ground support been given the systems intelligence and controls necessary to operate the systems?”

Training

A primary goal of intensive flight-specific training is integrating the flight crew and flight controller team. Even given the extensive experience of crews and their Mission Control Center support, the complexity and unique requirements of each flight demand intensive training. Basic training methods have not changed significantly from Apollo to shuttle, but the training tools have evolved tremendously.

Since Mercury, the core of training has been simulations that bring the crew and controllers together in as realistic a manner as possible. Normal flight phases are repeated to polish the performance and interaction of the whole team, but simulation personnel are well known for their ability to introduce problems that test documented procedures and mission rules. The simulations sometimes lead to changes and improvements, as well as to intimate knowledge of how systems operate.

An Apollo 11 example shows how important training can be to flight experience. Apollo 11 almost did not land on the moon because the crew kept receiving a series of computer alarms during the lunar module descent. But flight controllers quickly analyzed the alarms and advised the crew that they were not serious and the landing could continue. They knew that because of their training experience.

The training teams had discovered that some computer alarms intended only for ground testing could be triggered and give the crew and flight controllers a tough interpretation challenge. One of these internal computer alarms was triggered during an Apollo 10 simulation. In the process of determining that these alarms were not serious, the flight controllers investigated every alarm that the computer might display, and which were important.

Then the Apollo 10 flight introduced another factor. Problems in tracking the lunar module in lunar orbit after separation from the command module led to a decision to turn on the rendezvous radar in addition to the landing radar. Knowing

the lunar module position relative to the command module provided the information needed, but this also effectively doubled the work the lunar module computer had to handle, especially when shifting to the higher computation cycles during Apollo 11 descent. As the machine began to be overloaded, it started shedding less important tasks and sending alarm codes at an increasing frequency. With their in-depth understanding of the alarms, the flight controllers could determine that the critical tasks were being accomplished and gave the “go ahead” to continue the landing.

The evolution in training has been driven by improvements in the supporting computer technology. Basic spacecraft systems familiarization and operations procedures instruction is workstation-based and extremely realistic. But the major steps forward in the realism of the crew training with an accurate interface to the Mission Control Center have been in the mission simulators.

For both Apollo and shuttle, mockups and part-task trainers were important components of overall crew training. The Shuttle Training Aircraft covers the orbiter approach and landing phase, but the spacecraft simulators provide the mission environment. For the Apollo command module and lunar module training, both spacecraft required simulator crew stations that could operate in concert to cover mission operations,

Three of the four Apollo 13 flight directors applaud the successful splashdown of the command module “Odyssey” while Dr. Robert R. Gilruth, director, Manned Spacecraft Center (MSC), and Dr. Christopher C. Kraft Jr., MSC deputy director, light up cigars (upper left). The flight directors are (from left to right) Gerald D. Griffin, Eugene F. Kranz, and Glynn S. Lunney.

Photo Credit: NASA



This “fish-eye” view shows NASA’s Multifunction Electronic Display Subsystem (MEDS), otherwise known as the “glass cockpit.” The fixed-base Space Shuttle mission simulator in the Johnson Space Center’s Mission Simulation and Training Facility was outfitted with MEDS to be used by flight crews for training.



Photo Credit: NASA

including the lunar module’s lunar surface approach and landing. The software had to reproduce the actual flight systems with total accuracy. Because virtual image technology was not available then, the simulator out-the-window views were produced by a camera moving over a 3-D model of the lunar landing site. Based on robotic spacecraft imagery, the models were produced by the Department of Defense mapping facility in St. Louis.

SINCE MERCURY, THE CORE OF TRAINING HAS BEEN SIMULATIONS THAT BRING THE CREW AND CONTROLLERS TOGETHER IN AS REALISTIC A MANNER AS POSSIBLE.

The Shuttle Mission Simulator is the primary system for training shuttle crews. This high-fidelity simulator can train crews in all mission phases. There are two orbiter crew cockpits, both representative of an actual orbiter. A fixed-base crew station (FBCS), used for orbital training, accommodates the commander, pilot, mission specialist, and payload positions and has navigation, rendezvous, remote manipulator, and payload support systems so payload operations can be simulated. A motion-based crew station for ascent and entry training features a modified six-degrees-of-freedom motion system to give the commander and pilot the “feel” of mission phases. Digital image-generation systems provide window views in both simulator bases. The landing runway image and the ability to realistically project payload operations are particularly impressive.

During the simulations, system status and crew operations are transmitted to the flight controllers in the Mission Control Center just as they would be in flight. This enables the introduction of scenarios in which the crews and flight controllers must react to emergencies. The goal is to encounter no actual flight situations that have not been trained for in some manner.

Flying

Mission objectives have been finalized. The flight plan, mission rules, and operational procedures have been developed, refined,

and validated in training and then refined, reviewed, and redefined to a final preflight configuration. The flight crew and the flight control team are trained, and the Mission Control Center is configured and ready. Now it is up to the great launch teams at Kennedy and to Mother Nature’s winds and weather. So it has been through the launches before and during Apollo, for intervening programs, and through shuttle.

When the launch vehicle and spacecraft clear the pad, the Mission Control Center takes the handover from the launch team. Occasionally, the mission goes nearly exactly as planned. This has seldom been the case, but most eventualities are handled by the flight crew working with the flight controllers using established procedures, contingency and malfunction procedures refined in training, and the ingenuity of the combined team.

On what have fortunately been rare occasions, the combined team has been challenged by extraordinary issues. That is when the entire flight control team must muster its combined knowledge and experience. Contingencies have also produced individual flight controllers whose decisive actions have established them as icons in the history of human space flight.

I will name just a few. There was Steve Bales and Jack Garman’s response to the computer alarms during the Apollo 11 descent, and John Aaron’s actions after the Apollo 12 lightning strike during liftoff. They avoided an abort in both cases. There was the life-saving response of the entire control team under the leadership of Flight Directors Gene Kranz and Glynn Lunney after the Apollo 13 explosion. The shuttle program experienced two tragedies, but the overwhelming number of safe and successful missions and contributions to science and technology speak to the careful planning, training, and development of people, procedures, and teamwork in mission operations. Finally, any discussion of the contributions to NASA’s programs by operations must recognize the leadership, vision, and operations capability of Chris Kraft, the original flight director and the example to all who have followed. ●

JOHN O’NEILL is a former director of Mission Operations at the Johnson Space Center, where his NASA career covered thirty-four years in the operations organizations on Gemini, Apollo, Skylab, Apollo-Soyuz Test Project, Space Shuttle, and the International Space Station.



INTERVIEW WITH Kenneth Szalai

BY DON COHEN

Dr. Kenneth Szalai led the Dryden Flight Research Center from 1990 to 1998. Earlier in his career at Dryden, he was principal investigator on the F-8 digital fly-by-wire program and participated in many X-plane aircraft programs.

COHEN: I've heard you talk forcefully about the value of flight programs. What makes them so valuable?

SZALAI: I strongly support the flight aspect of the NASA aeronautics program because it is a primary means of discovery in aeronautics, and it is the truth serum and the lie detector of what is possible or not. Flight provides the real answer. You can't say, "It worked," when there is a cloud of black smoke coming up from the desert floor. And, in the process of doing it successfully, you provide a level of confidence to other people. Burt Rutan once said something very profound about SpaceShipOne: "Other people will say, 'If a guy out here on the Mojave Desert can put the equivalent of a three-person spaceship above a hundred kilometers, bring it back safely, and then do it again within a week, I should be able to do that.'"

COHEN: What's an example of that kind of project at Dryden?

SZALAI: When we did the digital fly-by-wire program, a prominent executive of a major aerospace company said, "That gave us the confidence to bid fly-by-wire in our proposal." He said that without having read any of the reports. We gave a three-day seminar to regulatory people in the eighties where we said, "Digital fly-by-wire is coming, and this is to assist you in what you will be dealing with." Many of the people present there said, "This will never happen." But look at the 787 and A380—they're fly-by-wire airplanes.

COHEN: Because you showed it could work.

SZALAI: And also because we showed the enormous benefits of digital fly-by-wire [DFBW]. In some ways, it was reverse technology transfer from space to aeronautics. The technology transfer for human-rated software went from the Apollo moon-landing program to



“

IT'S TRUE THAT YOU **can't start a program** SAYING, 'FUND THIS PROGRAM BECAUSE MAYBE **in the future** SOMETHING **important** WILL COME OUT OF IT,' BUT ANY **leading-edge** PROGRAM **is almost guaranteed** TO PRODUCE NEW UNDERSTANDING, NEW CONCEPTS, NEW IDEAS ...

”

Dryden, and then from Dryden to the aircraft industry. For the first phase of the F-8 DFBW project we used the Apollo lunar guidance computer and inertial navigation system. Because of that, Dryden became the first NASA research center to tap into the tremendous national treasure of processes, techniques, and procedures for flight software development for human-rated vehicles. This had been developed by the Charles Stark Draper Laboratory. The biggest two findings of the first phase of the F-8 program were, first, that digital fly-by-wire is possible and, second, that the task of human-rated software development is so complex and challenging that it will become *the* pacing element for all aircraft digital flight control systems. In the second phase, we essentially invented how to make multiple computers that work together for fault tolerance look like one computer to the pilot. One of the eye-openers in the F-8 program was that you can exercise the software to the *n*th degree in the simulator, test every path, test every function, and test as much as you can until you've reached the point

where you're not finding any more errors. But what happened when we started flying in August of 1976? We started finding software issues. I say "issues" because sometimes it was a specification error, sometimes it was an interpretation error, sometimes it was just something everyone overlooked. None of these issues ever showed up in flight, but they could have. By the way, I was lucky enough to be the chief engineer and principal investigator on the project. I was in the right place at the right time.

COHEN: Another example of the power of flight to convince skeptics?

SZALAI: Early in the Space Shuttle development, designers were still considering air-breathing engines on the Space Shuttle to give it go-around capability or add power if you were a little short, just like an airplane. After all, the crew has one attempt at landing if it has no power. But just imagine, now, if you had to carry a turbojet or turbofan engine or two on the shuttle orbiter with all its systems and all the fuel and all the controls and supporting

avionics. What do you think the payload would be? That might have *been* the payload. At that time, Milt Thompson, who was an X-15 pilot and a brilliant engineer as well as Dryden chief engineer, was saying to a lot of the Johnson Space Center people, “You know, you can do this all unpowered. We have proven that with the X-15 and the lifting bodies. We did unpowered approaches and landings safely and consistently.” But this was a big step to take with the nation’s newest human spacecraft. Many of the space people were a generation beyond the aeronautics group that started the space program. The shuttle program management said, “You’ve been landing the X-15 and the lifting bodies on this enormous lake bed, fifteen miles long and eight miles wide. We’ve got to land on a runway.” So Milt Thompson led an effort to do a convincing demonstration. Dryden launched the X-24B from a B-52 and went out about seventy miles supersonic. The rocket engine burns out, and then they come back from seventy miles and try to hit a new line painted on the runway. That’s a very high-precision task. John Manke, the test pilot who became the director of Dryden in ’81, hit the line. Maybe he missed it by five feet. Then Mike Love, who tragically later lost his life in an F-4 accident, repeated the flight with about the same precision. The shuttle people were amazed. That turned the tide. No one was going to change the shuttle to be a glider from orbit because someone in a simulator says, “I think we can do this.” By the way, an Ames pilot, Fred Drinkwater, was a key player in this program as well, developing low lift-to-drag approach and landing techniques using a large transport aircraft.

Of course, the Dryden lifting body program had its own objectives: exploring aerodynamics and handling qualities for hypersonic reentry vehicles. It produced its own technology and a tremendous amount of data. But arguably one of the most significant contributions to date is what it did for shuttle, which was never envisioned in the lifting body program plan. It’s true that you can’t start a program saying, “Fund this program because maybe in the future something important will come out of it,” but any leading-edge program is almost guaranteed to produce new understanding, new concepts, new ideas, as well as “to separate the real from the imagined and to make known the overlooked and the unexpected problems,” as Dr. Hugh Dryden stated when asked about the reason for flight research.

COHEN: You learn things that you can’t learn through simulation?

SZALAI: You can’t do everything in a simulator; you can’t do it in a lab. We have a tremendous amount of computational capability today, and it plays a dominant role in design and analysis. But given this level of capability in analysis, some thinking goes like this: “With computational tools, simulators, and labs, we can pretty much do everything on the ground. Then, if we have enough money and if there’s interest, and if we have to, we can fly something to validate our concepts at the end.” In all the years I was involved in flight research, we never had a program like that. Flight research is really flights of discovery. You use the flight vehicle as a pioneer and a probe to find out where we are in terms of

understanding and to uncover the gaps in understanding. A major purpose of flight research is to develop, tune, and validate the computational tools and ground facilities for future use in design and analysis. Many of the flight programs of the past produced critical data for both wind tunnel and computational people to develop their capabilities.

COHEN: How, for instance, has flight developed wind tunnel capabilities?

SZALAI: Let me give you a basic example. It’s a big deal in aerodynamics when flow transitions from laminar [smooth] to turbulent. In an aircraft it affects drag; it affects performance. In a spacecraft it dramatically affects heating. One of the things noted in flight is that the natural latent turbulence levels are very low, lower than in most wind tunnels. Why is that? In flight, in smooth air, there is no fan blowing the air across the airplane. In a wind tunnel, you have to make the wind move. Large fans create a flow that goes around corners and bounces off the walls. It’s complicated. NASA and others have spent years learning how to make quiet tunnels—that is, low-turbulence tunnels—so we could design and analyze laminar flow. A stainless steel cone, heavily instrumented to determine when you go from laminar flow to turbulent flow, was traded among wind tunnels to calibrate them. In some tunnels, the flow transitioned at relatively low speed. That meant the tunnel was quite “noisy,” it had a lot of turbulence in its flow. Some were better. Years ago at Dryden we took this cone and put it on the front of an F-15, to fly it in “real life.” The transition

Reynolds number, which indicates when transition occurs, was higher than it had been in virtually all the wind tunnels. In other words, smooth air conditions could not be completely duplicated on the ground. As a result of that flight experiment, wind tunnels were calibrated, tuned, and analyzed to make them better. Knowing how the wind tunnels differed from actual flight meant you could apply a correction factor. There's an example of discovery, not validation. That's what flight programs are for. Most are much more complex than carrying a cone on the nose boom.

For the Space Shuttle, and for next-generation aircraft and spacecraft, we use the same wind tunnels that had been validated over the years with aircraft, so that gives you a high degree of confidence. If you had no aeronautical legacy, you'd have to validate the tunnel for new conditions or environments *while* you were trying to design a new vehicle. The interference effects and the unsteady aerodynamics of the Space Shuttle stack—such large vehicles so close together flying subsonically, transonically, and supersonically—are so complex that they defy prediction to some degree. There was more than 100,000 hours of wind tunnel testing. We had to rely on these tunnels, validated over years of experience and validation in flight, because one could not do simple flight tests to predict the actual shuttle configuration. When you fly transonically—0.9 mach number to maybe 1.1 mach number—early wind tunnels had tremendous problems because the shock wave bounces off the walls and hits the vehicle, whereas in free air it never comes back. So you get erroneous data.

Through a lot of flight research in the early days and flight tests and calibration of wind tunnels, some clever people came up with both perforated and slotted walls so that the shock is swallowed. But for complex configurations in new flight regions, flight research also often finds the terms in the equation that you have left out on the ground.

COHEN: For example?

SZALAI: Sometimes it's something in the interaction between the pilot and the vehicle. Take the Space Shuttle again. The first landing to the runway, flight number five of *Enterprise*, resulted in a serious pilot-induced oscillation [PIO] in both roll and pitch. There was a complicated interaction between the pilot, the vehicle, the flight control system, the visibility out the front, and the configuration of the shuttle that led to a pilot-induced oscillation. The roll PIO was not too much of a technical challenge and was solved by reallocating control functions to the elevons. But the pitch PIO was a bad problem. Even after it was found on the Space Shuttle, even knowing it was there, it couldn't be duplicated on ground simulators. The only duplication of a sort that was done was on the F-8, where we replicated the flight-control system and the time delay and had a real pilot try to land a real airplane on a real runway. That's where the PIO exhibited itself. There's no substitute for the real environment.

COHEN: Dryden took a major role in the SOFIA [Stratospheric Observatory for Infrared Astronomy] program fairly recently.

SZALAI: SOFIA has a flight development phase required to complete the development, integration, and qualification of the overall systems on the aircraft, including a fail-safe system to open an enormous cavity in an airplane above 45,000 ft. at high subsonic speeds. This effort will draw on all the things for which Dryden has a high degree of expertise and experience, namely acoustics, fault-tolerant flight systems, flight controls, dynamics, turbulence, and, above all, safety. There's no book written on how to do this program. Dryden can draw on its deep flight research experience in dozens and dozens of projects to do this work. They know how to do this kind of project. It's a national asset to have that kind of experience within an organization. An aircraft company can't afford to do its own national flight research program.

NASA Dryden, as a national facility, has probably worked on more than a hundred airplanes. After fifteen or twenty years most Dryden people end up working on ten or twenty flight programs. But this does not mean that Dryden works independently of the aerospace industry—it works closely with them and each brings its experience to the program. The best way for the industry to attain the technology developed in our programs is to work in close cooperation with NASA as major contractors. Bell Aerospace was the prime contractor on the X-1, North American Aviation on the X-15, and Grumman on the X-29. They designed the aircraft and fully participated in the flight research programs.

NASA is supposed to do hard things that the industry is not yet ready to undertake as a product or commercial

“YOU USE THE **flight vehicle** AS A PIONEER AND A PROBE TO **find out** WHERE WE ARE IN **terms of understanding** AND TO **uncover the gaps** IN UNDERSTANDING.”

venture. I remember President Kennedy saying, “We do things not because they’re easy but because they’re hard.” Dr. Ken Iliff, who was the chief scientist at Dryden, used to tell me there were three piles you should put your programs into: the easy ones, the things you’re sure you can do; the too-hard ones that you shouldn’t even try; and then there’s the question-mark pile. That’s where Dryden fits in. You shouldn’t try to do something too easy or too hard. If it’s too easy, probably somebody else should be doing it commercially. The impossible may look different some day, but warp drives and anti-gravity machines are not programs for the president to announce or for Dryden to work on yet.

COHEN: What is a Dryden accomplishment few people are likely to know about?

SZALAI: Dryden developed an integrated flight propulsion control system for the SR-71, which showed that you could improve the range of the airplane by 7 percent just by properly integrating the control of the engine and inlets, and the control of the airframe. Dryden and engine companies collaborated on digital engine controls and adaptive engine controls for high-performance aircraft. In another program, Dryden developed

control system concepts that made it possible to fly and land an aircraft with severe damage or massive failures.

COHEN: Where do you see flight testing being important in the future?

SZALAI: One area in the near term will be work on more environmentally friendly airplanes. This includes issues of synthetic fuels, noise, and emissions that contribute to undesirable constituents in the atmosphere. There’s a lot to be learned about alternate fuels. Nobody knows what happens to an engine after years on synthetic fuel. What are the effects on maintenance? What happens to fuel that’s stored in a tank for a long time? What is a long time? Does a synthetic fuel degrade differently from JP [jet propellant]? What happens if it slightly degrades? Can you still use it? If you optimize the aerodynamics for something that doesn’t have a “green” engine and then you put a green engine on it, do you still get the same benefits? There’s an enormous role for flight to explore these things. Not to validate them years downstream, but to be part of the exploration and discovery process.

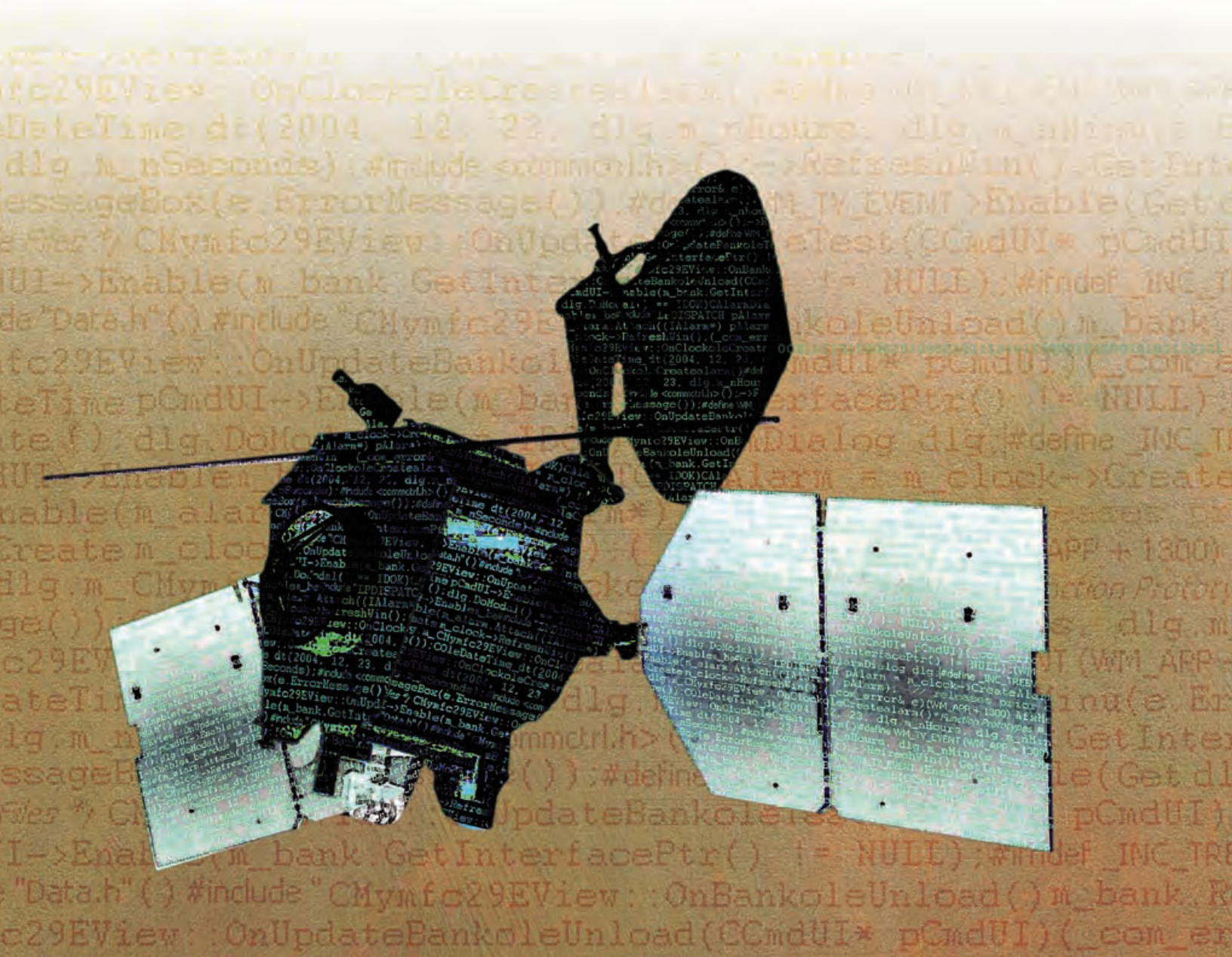
COHEN: Other important work happening at Dryden?

SZALAI: I emphasize flight because it’s often overlooked, but the Orion (the new space launch system) work is a very important thing Dryden is doing now, even at the expense of aeronautics activities. It is crucial that we have a national capability for access to space and a replacement for shuttle for getting to low-Earth orbit and beyond. Dryden is managing the Launch Abort System [LAS] for the Ares–Orion system. The LAS will operate in the atmosphere. It involves the integrated effects of rockets, aerodynamics, control systems, and life-support systems. These are things that Dryden has spent a lot of time on over the past decades. If problems occur during the project, Dryden will know a lot about how to make it a success, drawing on its aerospace flight research legacy. There’s nothing more important than getting that done, just as in the sixties there was nothing more important than NASA Dryden did than the Lunar Landing Research Vehicle which trained the astronauts—principally Neil Armstrong—to make safe manual landings on the moon. It played a pivotal role. There was nothing in aeronautics that would have been more important to do. ●

Is Software Broken?

BY STEVE JOLLY

A few years ago my attitude toward the design and development of space systems fundamentally changed. I participated in a Kaizen event (part of Lockheed Martin's Six Sigma/lean culture) to ascertain contributing factors and root causes of various software overruns and schedule delays that can precipitate cost and schedule problems on both large and small space programs alike, and to propose process improvements to address those causes. I didn't anticipate that software's modern role in spacecraft development could itself be a problem.



A Kaizen event is a tool used to improve processes, a technique popularized by the Toyota Corporation that is now used widely in industry. Stakeholders and process experts come together to analyze or map an existing process, make improvements (like eliminating work that adds no value), and achieve commitment from both the process owners and the users. In this particular Kaizen, software and systems engineering subject-matter experts came together from across our corporation to participate. We had data from several recent spacecraft developments that we could study.

We all suspected some of the cause would be laid at the feet of systems engineering and program management, with the balance of the issues being inadequate adherence to established software development processes or processes that needed improvement. But the Kaizen event is designed to ensure we systemically addressed all the facts, and our large room soon became a jungle of flipcharts covered in a dazzling array of colored sticky notes, each chart representing a different aspect of the problem and each sticky note a potential root cause. The biggest problem was there seemed to be somewhere around 130 root causes.

What we had expected based on other Kaizen events was that this huge number of root causes was really a symptom of perhaps a half-dozen underlying *true* root causes. A small number can be addressed; we can form action plans and attack them. Hundreds of causes cannot be handled. Something was wrong either with the Kaizen approach or with our data.

I came away from the event somewhat puzzled. We resumed the activity several months later, but we did not materially improve upon our initial list and get to a satisfying short list. However, several of us began to notice a pattern. Even though we couldn't definitively link the large majority of causes, we found that problems in requirements issues, development, testing, and validation and verification of the actual code all revolved around interfaces. When viewed from a higher altitude, the preponderance of the causes collectively involved all the spacecraft subsystems. With such systemic coverage of functions on the spacecraft, it was tempting to conclude that software processes were broken. What else could explain what we were seeing?

I couldn't accept that conclusion, however. I knew many of our software engineers personally—had walked many developments with them—and although we had instances of needing better process adherence and revised processes,

something else was clearly at work. I reflected on eight recent deep-space missions that spanned the mid-nineties through 2008, and it became clear that software has become the last refuge for fixing problems that crop up during development. That fact is not profound in itself; what is profound is that software is actually able to solve so many problems, across the entire spacecraft.

For example, while testing the Command and Data Handling (C&DH) subsystem on the Mars Reconnaissance Orbiter, we discovered a strange case of hardware failure deep in a Field-Programmable Gate Array (FPGA) that would result in stale sun-sensor data and a potential loss of power, which could lead to mission failure. To make the interplanetary launch window, there was no time to change or fix the avionics hardware. Instead, we developed additional fault-protection software that was able to interrogate certain FPGA data and precipitate a reset or “side swap” should the failure occur. Indeed, software is usually the only thing that can be fixed in assembly, test, and launch operations and the only viable alternative for flight operations. In fact, close inspection during our Kaizen showed that most of our 130 root causes could be traced to inadequate understanding of the requirements and design of a function or interface, not coding errors. Suddenly the pieces began to fall into place: it's all about the interfaces. Today, software touches everything in modern spacecraft development. Why does software fix hardware problems? Because it can. But there is a flip side.

In the past software could still be viewed as a bounded subsystem—that is, a subset of the spacecraft with few interfaces to the rest of the system. In today's spacecraft there is virtually no part of the system that software does not have an interface with or directly control. This is especially true when considering that firmware is also, in a sense, software. Software (along with avionics) has become the system.

This wasn't the case in the past. For example, Apollo had very few computers and, because of the available technology, very limited computing power. The Gemini flight computer and the later Apollo Guidance Computer (AGC) were limited to 13,000–36,000 words of storage lines.¹ The AGC's interaction with other subsystems was limited to those necessary to carry out its guidance function. Astronauts provided input to the AGC via a keypad interface; other subsystems onboard were

controlled manually, by ground command, or both combined with analog electrical devices. If we created a similar diagram of the Orion subsystems, it would reveal that flight software has interfaces with eleven of fourteen subsystems—only two less than the structure itself. Apollo's original 36,000 words of assembly language have grown to one million lines of high-level code on Orion.

Perhaps comparing the state-of-the-art spacecraft design from the sixties to that of today is not fair. The advent of object-oriented code, the growth in parameterization, and the absolute explosion of the use of firmware in evermore sophisticated devices like FPGAs (now reprogrammable) and application-specific integrated circuits (ASIC) have rapidly changed the art of spacecraft design and amplified flight software to the forefront of development issues. Any resemblance of a modern spacecraft to one forty years ago is merely physical; underneath lurks a different animal, and the development challenges have changed. But what about ten years ago?

Between the Mars Global Surveyor (MGS) era of the mid-nineties to the 2005 Mars Reconnaissance Orbiter (MRO) spacecraft, code growth in logical source lines of code (SLOC) more than doubled from 113,000 logical SLOC to 250,000 logical SLOC (both MGS and MRO had similar Mars orbiter functionality). And this comparison does not include the firmware growth from MGS to MRO, which is likely to be an order of magnitude greater. From Stardust and Mars Odyssey (late 1990s and 2001) to MRO, the parameter databases necessary to make the code fly these missions grew from about 25,000 for Stardust to more than 125,000 for MRO. Mars Odyssey had a few thousand parameters that could be classified as mission critical (that is, if they were wrong the mission was lost); MRO had more than 20,000. Although we now have the advantage of being able to reuse a lot of code design for radically different missions by simply adjusting parameters, we also have the disadvantage of tracking and certifying thousands upon thousands of parameters, and millions of combinations. This is not confined to the Mars program; it is true throughout our industry.

But it doesn't stop there, and this isn't just about software. Avionics (electronics) are hand-in-glove with software. In the late 1980s and early 1990s, a spacecraft would typically have many black boxes that made up its C&DH and power subsystems. As we progressed—generation after generation

of spacecraft avionics developments—we incorporated new electronics and new packaging techniques that increased the physical and functional density of the circuit card assembly. This resulted in several boxes becoming several cards in *one* box; for example, the functionality of twenty-two boxes of the MGS generation was collapsed into one box on Stardust and Mars Odyssey. Then, with the ever-increasing capability of FPGAs and ASICs and simultaneous decrease in power consumption and size, several cards became FPGAs on a *single* card. When you hold a card from a modern C&DH or power subsystem, you are likely holding many black boxes of the past. The system is now on a chip. Together with software, avionics has become the system.

BOTTOM LINE: THE GAME HAS CHANGED IN DEVELOPING SPACE SYSTEMS. SOFTWARE AND AVIONICS HAVE BECOME THE SYSTEM.

So then, there are no magic few underlying root causes for our flight software issues as we'd hoped to find at our Kaizen Event, but the hundreds of issues are unfortunately real. Most revolve around failed interface compatibility due to missing or incorrect requirements, changes on one side or other of the interface, poor documentation and communication, and late revelation of the issues. This indicates our systems engineering process needs to change because software and avionics have changed, and we must focus on *transforming* systems engineering to meet this challenge. This is not as simple as returning to best practices of the past; we need new best practices. The following are a few ways we can begin to address this transformation:

1. Software/firmware can no longer be treated as a subsystem, and systems engineering teams need a healthy amount of gifted and experienced software systems engineers and hardware systems engineers with software backgrounds.
2. We need agile yet thorough systems engineering techniques and tools to define and manage these numerous interfaces. They cannot be handled by the system specification alone or by software subsystem specification attempting horizontal integration with the other subsystems; this includes parameter assurance and management.
3. Using traditional interface control document techniques to accomplish this will likely bring a program to its knees due to the sheer overhead of such techniques (e.g., a 200-page formally adjudicated and signed-off interface control document).
4. Employing early interface validation via exchanged simulators, emulators, breadboards, and engineering development units with the subsystems and payloads is an absolute must.
5. If ignored, interface incompatibility will ultimately manifest itself during assembly, test, and launch operations and flight software changes will be the only viable means of making the launch window, creating an inevitable marching-army effect and huge cost overruns.

Bottom line: the game has changed in developing space systems. Software and avionics have become the system. One way to look at it is that structures, mechanisms, propulsion, etc., are all supporting this new system (apologies to all you mechanical types out there).

Today's avionics components that make up the C&DH and power functions are systems on a chip (many boxes of the past on a chip) and, together with the software and firmware, constitute myriad interfaces to everything on the spacecraft. To be a successful system integrator, whether on something as huge as Orion and Constellation or as small as a student-developed mission, we must engineer and understand the details of these hardware–software interfaces, down to the circuit level or deeper. I am referring to the core avionics that constitute the system, those that handle input-output, command and control, power distribution, and fault protection, not avionics components that

attach to the system with few interfaces (like a star camera). If one merely procures the C&DH and power components as black boxes and does not understand their design, their failure modes, their interaction with the physical spacecraft and its environment, and how software knits the whole story together, then software will inevitably be accused of causing overruns and schedule delays. And, as leaders, we will have missed our opportunity to learn from the past and ensure mission success.

A final note of caution: while providing marvelous capability and flexibility, I think the use of modern electronics and software has actually increased the failure modes and effects that we must deal with in modern space system design. Since we can't go back to the past (and who would want to?), we must transform systems engineering, software, and avionics to meet this challenge. I am ringing the bell; we need a NASA–industry dialogue on this subject. ●

Note: Kaizen is a registered trademark of Kaizen Institute Ltd. Six Sigma is a registered trademark of Motorola, Inc.

STEVE JOLLY is with Lockheed Martin Space Systems. He has development and flight operations experience from many deep-space missions and was the chief systems engineer for the Mars Reconnaissance Orbiter. Most recently he was a member of the NASA integrated design optimization team for Orion and is currently the chief systems engineer for GOES-R.



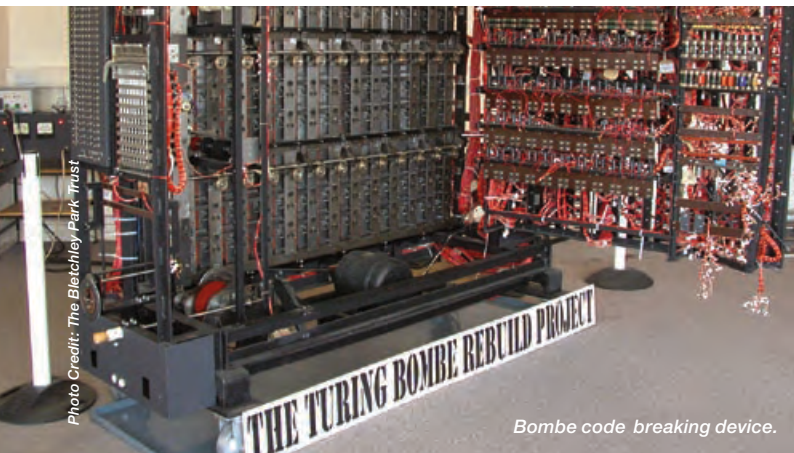
1. James E. Tomayko, *Computers in Spaceflight: The NASA Experience*, NASA Contractor Report 182505 (Washington, DC: NASA History Office, 1988).

Project Lessons from CODE BREAKERS & CODE MAKERS

BY JOHN EMOND

*Navajo Indian Code Talkers Henry
Bake and George Kirk with a marine
signal unit in December 1943.*





Bombe code breaking device.

You may wonder what on Earth World War II code has to do with NASA in the twenty-first century. It's a fair question. My response is that the challenges involved in creating communities of special talents to successfully decipher or protect vital information at a time of global conflict provide insight into effective project management practices as NASA addresses its own challenges.

The goals of code makers and code breakers are to protect valuable information on the one hand and to penetrate that shield of protection to uncover vital information on the other. Codes existed long before the current world of software encryption versus computer hackers.

American hero turned traitor Major General Benedict Arnold used a coded message dated July 12, 1780, to tell his British contacts he was to command the fort at West Point, New York, and in that capacity could surrender the fort to the Crown, doing great damage to the American cause. During the American Civil War, both Union and Confederate forces used ciphers.

A World at War

In the spring of 1940, the German blitzkrieg crossed what had been the static battlefields of World War I trench warfare, covering in two months territory they had failed to capture between 1914 and 1918. In Asia, following the December 7, 1941, attack on Pearl Harbor, the Japanese expanded their empire into the Asian continent as well as the Philippines, which surrendered May 6, 1942.

While armies clashed, a battle of minds was fought over vital military information: troop strength, deployment, intended lines of defense, and targets. Codes protected battle plans and strategies; code breakers tried to unlock the keys to such plans.

Code Breakers: Device, Counter-Device

The German Enigma machine was an electromechanical device with rotating "wheels" that scrambled plain text messages into cipher text. When an operator made a keystroke on this typewriter-like device, a given code symbol was sent to the receiving party. Hitting the same keystroke sent a different code symbol; there were billions of potential combinations. To counter this encryption machine, the British at Bletchley Park used a code-breaking device named Bombe, given to them by Polish crypto-analysts at the outbreak of the war.

The Enigma code was broken by 1940. At about that time, the German High Command developed Lorenz, an even more complex coding device. To crack its codes, the British created Colossus, a forerunner to the modern computer that used 1,500 vacuum tubes and read tape at 5,000 characters per second. Cumbersome and unwieldy by today's standards, it was a major leap forward in data processing at the time. Colossus was delivered to Bletchley Park in June 1943 and was operational in time to confirm that the Germans continued to be deceived regarding the intended site of the D-Day invasion.

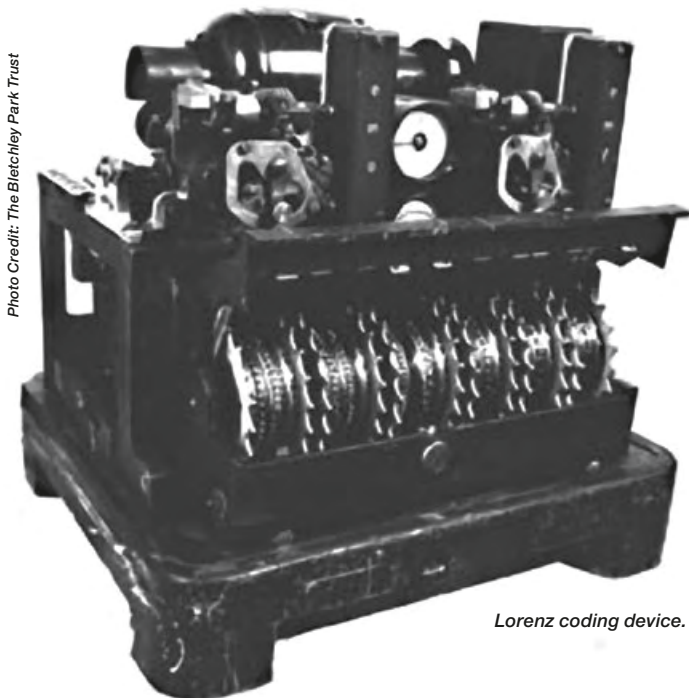
Ultimately, though, people, not machines, cracked the German codes. A community arose in Bletchley Park, 50 miles north of London, with the mix of skill and experience required to meet the daunting challenge of German intelligence.

The first group moved there on August 15, 1939, weeks before the start of hostilities. By war's end, Bletchley Park had grown to an eclectic community of 9,000 people. Among its members were not only brilliant mathematicians (including Alan Turing) and those with traditional science backgrounds, but also a poet, a schoolteacher, a novelist, chess champions, and crossword puzzle experts. (The ability to solve the *Daily Telegraph* crossword puzzle in under twelve minutes was used as a recruitment test; the winner succeeded in under eight.)

Code Makers: Navajo Code Talkers

While the English were engaged in cracking German military codes, American intelligence officials were trying to devise ways to protect their military communications. Japanese code breakers, some of whom had lived in the United States and were fluent in American English—even American slang—intercepted or sabotaged communications sent to and from units in the field. In 1942, Philip Johnston read a newspaper account of the Louisiana code military communications staff trying to use Native American personnel for code development. That inspired Johnston, the son of a Protestant missionary who grew up with Navajo children, to raise the idea of Navajo “code talkers” in a briefing to Lt. Colonel James Jones.

Until just prior to World War II, the Navajo language had never been written or translated. A new dictionary was devised to convert military terms into Navajo words, often with an origin in Navajo culture, which raised the complexity of the message and increased the difficulty of uncovering its true meaning. For instance, the Navajo word “lo-tso,” meaning “whale,” was used for “battleship;” “chay-da-gahi” (tortoise) was “tank;” and “da-he-tih-hi” (hummingbird) was “fighter plane.”



Lorenz coding device.

American Indians were trained to encode, transmit, and decode a three-line English message in twenty seconds; machine translation took thirty minutes. Though many code talkers had never been away from their reservation, all but two who were left behind to train others were deployed to Guadalcanal in U.S. Marine units. There was little time to test this code system before it was deployed under combat conditions. In addition to the Guadalcanal campaign, Navajo code talkers took part in every Pacific combat operation from 1942 to 1945, including Saipan, Iwo Jima, and Okinawa. By war's end there were 450 Navajo code talkers.

High Risk, High Rewards

Project managers at Bletchley Park and those who supervised the Navajo code talkers engaged in high-risk, high-reward situations. The English could have failed to crack the German codes, or German intelligence could have discovered that their codes were compromised. Either eventuality could have greatly increased casualties and possibly changed the war's outcome.

Success in breaking the German code produced dramatic benefits. The Royal Air Force (RAF) had advance knowledge of the German air force's (Luftwaffe) overall strategy regarding English targets during the 1940 Battle of Britain. The heroic defense of the RAF, and the strategic error of the Luftwaffe in shifting from air attacks on RAF bases to bombing London and other cities, meant English victory in the air war. Without air supremacy, the Germans could not invade England. Knowing the German codes and communication messages meant Allied naval forces could interdict the German U-boat “wolf packs” that had preyed on Allied merchant shipping in the Atlantic, particularly U.S. supply convoys that kept England alive.

Deciphering the German codes also enabled Allied forces to attack German Afrika Korps ship convoys, denying crucial supplies to General Erwin Rommel's desert army and helping to weaken his army's offensive and defensive operations. Perhaps the most dramatic success was the verification at Bletchley Park that the Germans remained deceived about the D-Day invasion site. If the Germans had known in advance the “spear tip” of the invasion target was to be Normandy, plans could have been put in motion to make the Allies pay a much steeper price.

As for the Navajo code talkers, failure to protect the code would have exposed many thousands of U.S. Marines in their Pacific island-hopping campaigns to even greater danger than they actually faced. The full impact of the code talkers may never be known. According to Major Howard Connor, Fifth Marine Division signal officer, were it not for the Navajos, the Marines would never have taken Iwo Jima. Six Navajo code talkers worked around the clock during the first forty-eight hours of battle on Iwo Jima; 800 messages were sent without error.

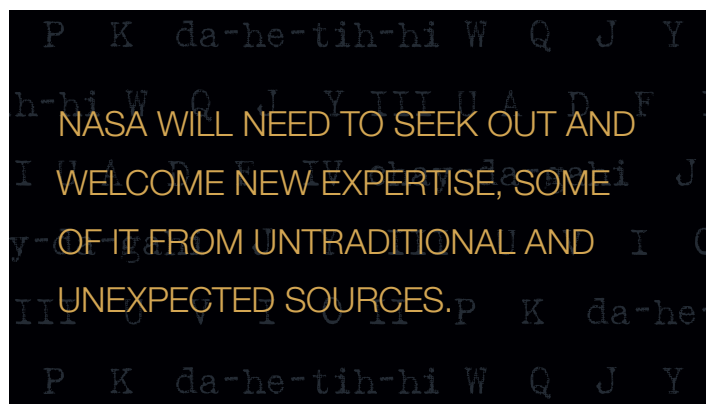
Given the damage that would have resulted if the enemy had learned the secrets of Bletchley Park and the Navajo code

talkers, project management placed a great deal of trust in these communities. That trust proved to be well-founded. The secrets of these operations went undisclosed for decades after the war's end.

Meeting NASA's Challenges

NASA does not face such wartime crises, but the Agency has encountered and met many challenges and faces many more in its evolution from “sortie” missions in low-Earth orbit to a crewed International Space Station and long-duration exploration missions. Some of the big ones include the following:

- In life sciences, adaptation to variable gravity environments, radiation mitigation, telemedicine/remote medical diagnosis and treatment, sustainable sources of food and water, and psychological and physical stress factors of long durations in closed environment systems at great distances from the Earth
- Advance material development for habitat construction and performance in extreme environments
- Energy conservation and production
- New technologies in spacecraft propulsion
- In-situ resource utilization



Successfully meeting these and other challenges will require many minds from many disciplines to share ideas and generate a varied range of approaches to common goals and objectives. As in the cases of Bletchley Park and the Navajo code talkers, NASA will need to seek out and welcome new expertise, some of it from untraditional and unexpected sources.

This process has already begun. In recent years, the Agency has begun to explore new approaches and creative, nontraditional paths to technology advancement. Today, NASA missions involving modeling and simulation draw on men and women who grew up on computer simulation games, much as Bletchley project management used crossword puzzle experts to explore cipher code formation. NASA's Centennial Challenge looks for expertise wherever it exists by using prize money to



German Enigma machine.

Photo Credit: The Bletchley Park Trust

encourage individuals or teams outside NASA to meet specific technology objectives. In 2007 the Astronaut Glove Challenge to improve glove design was won by Peter Homer, a former aerospace engineer whose background included—in addition to aerospace—sailing and sail making. Working at his dining room table, Homer crafted an improved glove design that included off-the-shelf materials.

One might argue that the English code breaker and American code maker project managers took major risks in their nontraditional approaches and personnel. They did so because the stakes were so high that new approaches, even apparently risky ones, were justified by the need to confront the deadly challenges faced. NASA's future missions will not be part of dire combat for survival, and NASA's mandate is to broadly disseminate, not hide, its goals and its knowledge. But the magnitude and difficulty of NASA's future missions, and the challenges known and still to be encountered, require the same level of commitment, determination, and trust. Only by broadly using and trusting the creative, dedicated talents of a wide and diverse community will we succeed. ●

Note: Special thanks to The Bletchley Park Trust (www.bletchleypark.org.uk) for their cooperation.

JOHN EMOND became a Presidential Management Intern at Goddard Space Flight Center in 1982 and was a contract specialist at Goddard from 1984 to 1987 before joining the NASA Headquarters Office of Commercial Programs as a policy analyst. He is presently a collaboration program coordinator within the Innovative Partnerships Program office at NASA Headquarters, with a primary focus on fostering interagency collaboration in technology development and application.



Leave No Stone Unturned, Ideas for Innovations Are Everywhere

BY COLIN ANGLE

Maybe an employee comes up with that brilliant idea for a new product. Or perhaps a customer makes a great improvement suggestion. Plus, you always have to keep your eye on the groundbreaking research and development (R&D) taking place in business, academia, and elsewhere. That's the perpetual challenge: you never know where innovation is going to come from, which is why you need to look for it everywhere.



Industry partnerships are one way to develop new products that may never make it to market otherwise. One example is the iRobot PackBot with ICx Fido Explosives Detection Kit. iRobot teamed with ICx Technologies to integrate its explosive detection technology into the combat proven PackBot platform.

Photo courtesy iRobot Corporation

At iRobot, talented people and continuous communication are the cornerstones of innovation. We hire the smartest, most curious, and very creative people and communicate with them continually. Regularly scheduled all-hands meetings, just one of the companywide communication efforts, help everyone know about the projects being tackled by the other divisions. This basic but critical communication has stimulated a tremendous amount of cross-fertilization. For example, mine-hunting algorithms developed for our government and industrial robots were used by our home robots division for the iRobot Roomba floor-vacuuming robot. Likewise, the tracks on the iRobot Looj gutter-cleaning robot were derived from the iRobot PackBot tactical mobile robot.

But building a solid foundation is only the beginning. Whether it's improving existing products or developing new ones, you have to dig, poke, and prod to keep innovation alive.

Looking Within

iRobot's willingness to listen to employees and support their ideas has been central to our innovative core. In fact, the idea for Roomba came from a group of employees who thought an easy-to-use home-cleaning robot could make a difference in people's lives. The employees were given two weeks and \$10,000 to develop the concept—a collaboration that ultimately led to the world's first affordable vacuum-cleaning robot. Roomba made practical home robots a reality for the first time and showed the world that robots are here to stay.

Today, more than three million Roomba robots have been sold worldwide. Roomba is not the only example. Looj is the result of an internal design contest; an employee made a crude prototype using a couple of motors and a bottle brush. Looj was rolling off the assembly line and into homes just nine months later. Similar internal exercises in our government and industrial robots division have led to a host of interesting capabilities for our military robots, including autonomous stair climbing.

Internal R&D has also been essential to the innovations brought to market by iRobot throughout the company's nearly twenty-year history. With internal R&D, you always have to keep an open mind, and you always have to keep pushing yourself. One way to do this is by responding to government requests for proposals (RFPs) with requirements that are so far off the beaten track that they seem almost silly. But they force you to think in different ways and often result in success.

One of those recent successes for iRobot started with a government RFP to develop a soft, flexible, mobile robot that can maneuver through openings smaller than its actual structural dimensions—not something that comes across your desk every day. We responded and have been awarded a multiyear, multimillion dollar R&D project from the Defense Advanced Research Projects Agency (DARPA) and the U.S. Army Research

Office to develop chemical robots called ChemBots. iRobot will lead a team of leading technical experts from Harvard University and the Massachusetts Institute of Technology to incorporate advances in chemistry, materials science, actuator technologies, electronics, sensors, and fabrication techniques.

Through this ChemBots program, robots that reconstitute size, shape, and functionality after traversal through complex environments—the stuff of science fiction—will become real tools for soldiers. ChemBots will be a revolutionary new robot platform that will expand the capabilities of robots in urban search and rescue, as well as on reconnaissance missions.

iRobot's internal R&D efforts have helped establish our position as the leader in innovative robotics research and development. This didn't happen overnight; in the past decade, iRobot has won a series of DARPA project awards. DARPA initially approached iRobot in 1998 to develop a robot for its Tactical Mobile Robot (TMR) program. DARPA requested white papers for a rugged, reliable, man-packable robot that could climb stairs and be used in urban combat. The other invited organizations wrote papers outlining their thoughts on urban combat robots. Only iRobot submitted a working robot prototype—along with the white paper.

WHETHER IT'S IMPROVING EXISTING
PRODUCTS OR DEVELOPING NEW ONES,
YOU HAVE TO DIG, POKE, AND PROD TO
KEEP INNOVATION ALIVE.

The first TMR units were deployed on September 11, 2001, at the World Trade Center for on-scene investigation. Units were also deployed in Afghanistan, where they assisted frontline troops in searching enemy tunnel and cave complexes. TMR demonstrated iRobot's technical capabilities to design, develop, prototype, test, and field a combat-deployable robot with advanced sensor capabilities. The follow-on contract from DARPA led to the creation of the iRobot PackBot, one of the most successful battle-tested robots in the world. PackBot is a multimission tactical mobile robot that performs search, reconnaissance, bomb disposal, and other dangerous missions while keeping warfighters and first responders out of harm's



Photo courtesy iRobot Corporation

The idea for the iRobot Roomba vacuum-cleaning robot came from company employees who thought an easy-to-use home-cleaning robot could make a difference in people's lives. The employees were given two weeks and \$10,000 to develop the concept—a collaboration that ultimately led to the world's first affordable vacuum-cleaning robot.

way. More than 2,000 PackBot robots have been delivered to military and civil defense forces worldwide.

Over the years, iRobot has earned a reputation for rapid response and innovative thinking, deploying high-quality solutions in the shortest time possible. One of the most valuable lessons we've learned is to listen to our customers, both military and consumer.

Many of the very best ideas come from talking with users; we are constantly finding ways to reach out to them, to understand what they like about our robots, and to find out what other kinds of tasks they'd like our robots to perform.

The iRobot PackBot 510 with EOD Kit, a second-generation explosive ordnance disposal (EOD) robot, was designed specifically to address evolving end-user requirements for a stronger, faster, and easier-to-use robot. The robot is 30 percent faster, drags larger objects, lifts twice the weight, and has a grip that is three times stronger than its predecessor. One of the coolest innovations is the robot's new hand controller. Modeled after video-game controllers, it makes PackBot 510 much easier and more intuitive to operate, resulting in less training time and more rapid and effective operations in the field.

In our home robots division, listening to customer suggestions has resulted in numerous product improvements and even the

development of a successful new product. After the introduction of Roomba, we were often asked for a robot that could take on the laborious task of mopping floors. In response, we developed the iRobot Scooba floor-washing robot. Scooba built on our knowledge of floor cleaning but required new expertise in wet cleaning, fluid dynamics, locomotion in a wet environment, and other challenges. Scooba's innovative technology uses a four-stage cleaning system to prep, wash, scrub, and squeegee a variety of hard floor surfaces. Unlike mop and bucket cleaning, Scooba uses only clean solution to wash floors, never dirty water. Scooba is a perfect example of an innovation that resulted from customer suggestions and by building upon a previous breakthrough.

Going Outside

While internal efforts are crucial, you also can't lose sight of the important R&D going on in other companies, research labs, and universities around the world. You also need to know when to seize opportunities; iRobot recently leveraged external R&D done by two separate entities to advance the use and capabilities of Unmanned Underwater Vehicles (UUVs) and bring those innovations to the marketplace.

iRobot recently entered into a sole licensing agreement with the University of Washington to commercialize Seaglider,

a long-range, high-endurance UUV that makes oceanographic measurements. A deep-diving UUV, Seaglider is designed for missions lasting many months and covering thousands of miles. With Seaglider, civilian, academic, and military personnel can make oceanographic measurements, perform persistent surveillance, and accomplish other missions for a fraction of the costs of traditional vessels and instruments.

To accelerate iRobot's strategic push into maritime systems, we also recently acquired Nekton Research, an unmanned underwater robot and technology company based in Raleigh-Durham, N.C. Over the next year, iRobot's newly established Maritime Programs office in that location will focus on turning the Ranger prototype, a general development platform for small UUV capabilities, into a product. The licensing agreement and the acquisition have both helped to ensure that these innovations will eventually be available to customers who need robots to conquer new underwater frontiers.

Forming industry partnerships is another effective way to develop new products that may never make it to market otherwise. One example is the iRobot PackBot with ICx Fido Explosives Detection Kit. iRobot teamed with ICx Technologies to integrate its explosive-detection technology into the combat-proven PackBot platform. The robot can detect explosive vapors emanating from Improvised Explosive Devices (IEDs). PackBot's dexterous, seven-foot arm allows the robot to place the explosive sensor close to suspicious packages and other objects, as well as reach through car windows and under vehicles. PackBot can also use its onboard capabilities to destroy IEDs, while warfighters remain out of harm's way.

We were able to quickly test the robot in the field, get feedback, and refine the product based on user requirements. PackBot's digital, modular architecture enabled fast and easy integration of the Fido payload, enabling our troops on the battlefield to quickly benefit from the new innovation.

The partnership has resulted in the first major deployment of explosives-detection robots and also demonstrates a new market

application for robots. More than 100 of the explosive-detection robots have been ordered for use by the U.S. military in Iraq. This is an exciting and promising development; there is a growing need for robots that can safely detect and disrupt explosives, not just for warfighters deployed in the Middle East, but also for first responders around the world.

Ideas Everywhere

There is no one source for new ideas and no single way to innovate. Our experience at iRobot shows the range of approaches innovative organizations can and must embrace: listen to your employees and customers; tap into R&D; forge partnerships. Ideas for innovations are lurking everywhere; it's your job to look high and low for them. ●



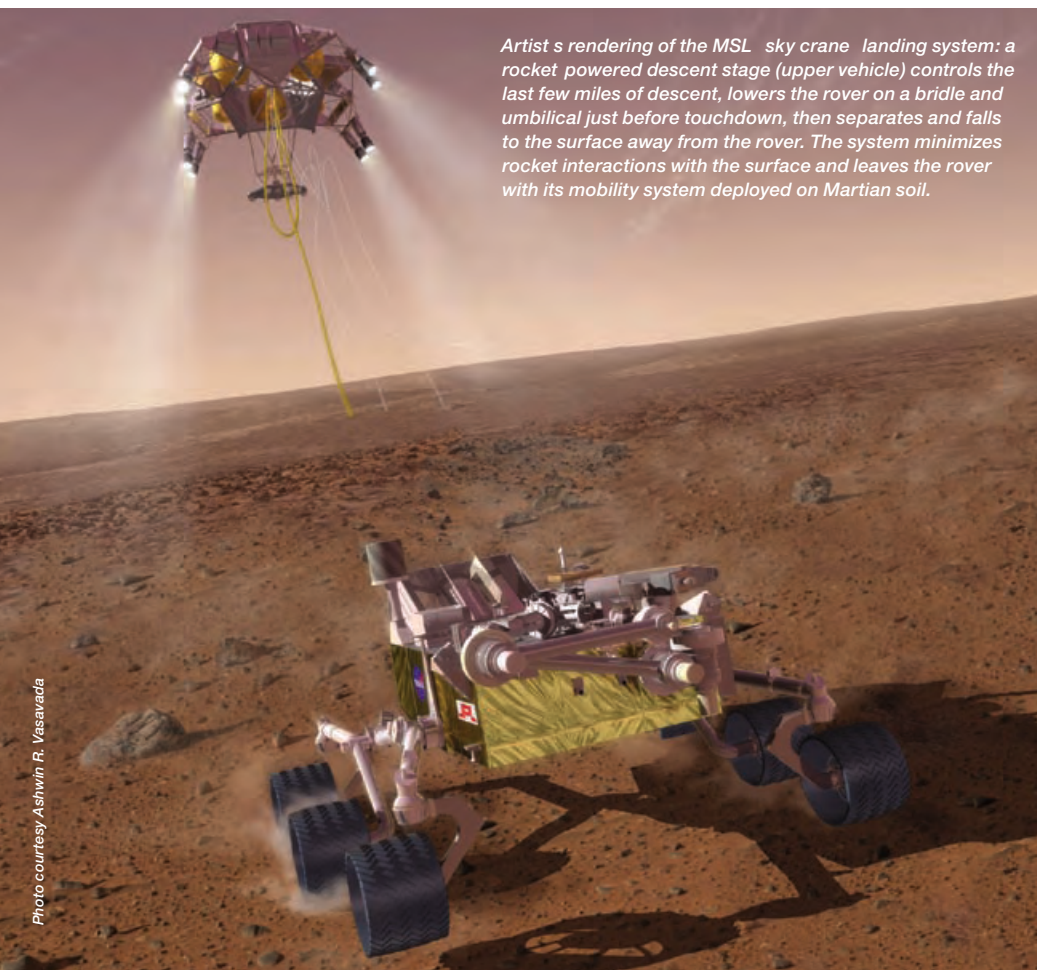
COLIN ANGLE is the chairman, CEO, and co-founder of iRobot.

MARS SCIENCE LABORATORY:

Integrating Science and Engineering Teams

BY ASHWIN R. VASAVADA

NASA's robotic exploration of Mars represents, perhaps more than any other human endeavor, both a scientific and an engineering achievement. A Mars rover must survive a violent launch and entry into the Martian atmosphere, descend and land safely, navigate over rocky terrain, manage power and data across many subsystems, communicate with orbiters and ground stations on Earth, and operate for multiple years in a dusty environment of extreme thermal contrasts. Yet these profound technical challenges are only half the story. After four decades of Mars missions, the scientific goals are equally ambitious.



Artist's rendering of the MSL sky crane landing system: a rocket-powered descent stage (upper vehicle) controls the last few miles of descent, lowers the rover on a bridle and umbilical just before touchdown, then separates and falls to the surface away from the rover. The system minimizes rocket interactions with the surface and leaves the rover with its mobility system deployed on Martian soil.

Each of the twin Mars Exploration Rovers (MER) that reached Mars in 2004 has found evidence that liquid water interacted with surface rocks and soils for a sustained period early in Martian history. Launching in 2011, the Mars Science Laboratory (MSL) rover mission will further assess Mars's habitability by exploring a new landing site with a car-sized rover carrying more than 80 kg of scientific instruments. In addition to the imaging and "place and hold" sensors from MER, MSL is equipped with a system to drill into rocks and deliver sieved powder to chemical and mineralogical laboratories inside the rover body. It will make measurements that are difficult even by terrestrial standards: high-resolution imaging devices need stable platforms and precise articulation; chemical detectors require controls on contamination and sample handling; and calibration, cleaning, adjusting, and troubleshooting must all be accomplished through the looking glass of the spacecraft's telemetry.



Photo courtesy Ashwin R. Vasavada

A family history of JPL Mars rovers: a full-scale model of the Sojourner rover (center) that accompanied the Mars Pathfinder lander in 1997; a model of Spirit and Opportunity, the Mars Exploration Rovers (left); and a model of the Mars Science Laboratory (right), with its seven-foot-high remote-sensing mast.

But what makes Mars exploration unique is not just the challenge it presents to the scientists and engineers involved, but how deeply these communities must be integrated in order to succeed. Mars exploration is a program driven by scientific goals but enabled by technical achievements. Technical design choices are motivated by the mission's science objectives and each carries the potential to either limit or enhance science return. Furthermore, upon reaching Mars, the robotic vehicle becomes a virtual scientist: the investigations of scientists on Earth are accomplished via the spacecraft and the technical teams that operate it.

The need to bring scientists and engineers together is obvious, but that essential close collaboration is not easy to achieve in practice. For example, it's tempting to view the two communities as playing distinct roles over the life cycle of a project. A mission might be defined by a team of external (for example, university) scientists, handing off their requirements to a NASA center comprising predominantly engineers who will design and develop the spacecraft, which will then return data to be analyzed by scientists scattered around the world. But having an integrated team during the critical middle stages is what ensures the mission will accomplish what was originally intended. Our strategy within the MSL project at the Jet Propulsion Laboratory (JPL) has been to integrate a small number of in-house project scientists into the engineering teams during design and development. Here are a few thoughts on how we've done that and what we've learned. While born out of a Mars mission, these strategies will be useful for any NASA mission with a scientific component, research or applied.

Who: Choose the Right People

Integrating a scientific and technical team starts by recruiting team members who are not only excellent in their disciplines but also have the skills and temperament conducive to working together. Project scientists are the liaisons between the engineers and the external science community. Strong research credentials

will ensure that they are respected by their scientific peers and trusted to convey and defend the mission's scientific objectives. Project scientists are likely to get a wide range of questions from project engineers, requiring broad knowledge and the ability to make timely judgment calls. JPL hires scientists with these criteria in mind and encourages them to stay current in their fields by dedicating a fraction of their project-supported time to related research. Sometimes it may be necessary to bring in additional expertise from the external community, perhaps by holding a teleconference or forming a working group. In such cases the project scientist can help find the experts, frame the questions, and translate the responses.

On the engineering side, JPL has found that project system engineers with research experience have a good framework for understanding how the scientific and technical aspects of a project interact. Also, the importance of temperament should never be underestimated. Scientists and engineers come from different cultures and bring their own presuppositions about each other's motivations and capabilities. Rarely does a team regret integrating, but the path may initially be rocky.

How: Understand Each Other's Work

The negative presuppositions held by each community sound something like this: "Scientists have unrealistic ideas about what can be done and always want more," or "the engineers are killing the science." It may take time and plenty of dialogue, but the goal is to go from a view of working at cross purposes ("look, if we don't land successfully, we're not going to get any science anyway!"), to one of thoughtful negotiation ("how much can we cut your energy before the science quality really suffers?"), to eventually one of working together toward a shared goal.

One of the most successful ways of moving along this path is to jointly discuss the larger scientific and technical context, not just the issues at the interface. For example, scientists need to develop an inherent appreciation for the cost, mass, schedule, and risk pressures that drive the engineers to certain choices.



The integrated MSL spacecraft in the Spacecraft Assembly Facility at JPL along with the Assembly, Test, and Launch Operations team. Visible are the cruise stage (lower segmented ring), the backshell, and heat shield. Inside the capsule are the rover and descent stage. The capsule is 15 ft. in diameter, larger than the Apollo capsule.

Photo courtesy Ashwin R. Vasavada

Meanwhile, scientists can educate the engineers on the objectives of the mission, improving their judgment about where to push constraints and where to accept limitations. These approaches have worked especially well within our team that is designing and testing the MSL entry, descent, and landing (EDL) system. This team commissions external scientists to run state-of-the-art models of Martian weather, statistically analyzes the results against the capabilities of the EDL system, and simulates the spacecraft's flight through the modeled atmosphere. Although there is enormous complexity on both sides of the science–engineering interface, we have a simple goal within the team: “No black boxes on either side.” Only by sharing our expertise with each other can we ensure accuracy and eventually a successful landing on Mars.

Where: Be in the Right Place

A large project like MSL, with more than one thousand JPL staff at the peak, is organized into systems, subsystems, and even smaller groupings. Where in this organizational structure should we integrate our dozen or so in-house scientists?

Our method is to attach an individual scientist (most of them part time) to each scientific instrument and the rover's sampling system—in fact, to every project element that has a major scientific component. We also have two scientists attached part time to the team in charge of EDL: an expert on Mars's surface and one on Mars's atmosphere. The project scientist and his two full-time deputies interact primarily with the project core staff and address project-level requirements, design trades, and other management issues. Scientists pick up additional assignments as the need and our skills dictate, on topics including the expected environmental conditions on Mars, contamination issues, or mission operations strategies and tools.

By maintaining a presence across both dimensions of the project organizational chart, we can provide scientific guidance where needed and can stay current with engineering progress. JPL has taken steps to promote this integration by giving the project scientist shared authority with the project manager on decisions that affect science. The MSL project has helped by requiring scientific representation on all major internal reviews.

But being there is what counts; a single missed meeting or review might be regretted.

When: Integrate Over All Phases of the Project

The sample acquisition and processing system on MSL has been a proving ground for our ability to integrate our scientists and engineers. A six-foot robotic arm must accurately place a jackhammer drill on a rock, then drill, collect, and sieve the rock powder and deliver measured amounts to instruments inside the rover body, all while minimizing fractionation (that is, separation by particle mass or size) and cross-contamination. Early in the project, we attempted to create a set of science requirements that would bound the capabilities of the subsystem but leave implementation choices to the engineers. This serial approach proved naïve in several ways.

For one thing, it was difficult to formulate a comprehensive set of requirements given the discovery-driven nature of the mission, with uncertainties about what might be encountered and what problems might occur. We just don't know how many hard versus soft rocks we will encounter, even though such a prediction would provide a quantitative way of estimating the life of drill bits. Furthermore, implementation choices would affect the scientific quality and integrity of samples in different ways. At one point the design team worked hard to preserve our ability to study various depths within a rock by retrieving an intact core sample. Our science team grew increasingly concerned about the fractionation that would occur when crushing the core into powder and began to prefer a less complex powdering drill, but those relative judgments weren't efficiently passed along to the engineers. We determined that it would be better if the end-user scientists were involved in the design process, weighing in on these decisions and sharing experiences from their own laboratories. Together the team could iteratively find a set of requirements and design choices.

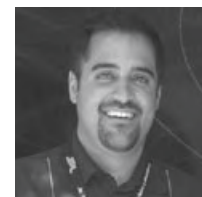
We've now settled on a process that provides scientific input by both embedding a scientist full time into the team and

by holding occasional science reviews of the subsystem with the external scientists who will eventually use the sampling capabilities. As the prototypes and flight-like models go through testing, our in-house scientist will provide samples of Mars's analog rocks and soils and analyze performance. This is an important point for all the scientific equipment on the rover: the utility of the data returned from Mars depends on the characterization and calibration performed prior to launch, not just the verification against a set of functional requirements.

In the end, integration is successful only if internalized by individual team members. Our scientists are passionate about what they hope to achieve and know their success is tied to that of the engineering teams. One turning point occurred when our sampling system team revealed a new version of their design, much more clever and capable than anticipated by the scientists who will get to use it one day on Mars. In the other direction, the engineering teams have been greatly motivated by the enthusiastic and detailed discussions of landing sites and the potential for discovery. In the end it all comes down to two things: building trust and recognizing each other's unique contributions. ●

Note: This work was carried out at the Jet Propulsion Laboratory, California Institute of Technology, under a contract with the National Aeronautics and Space Administration. © California Institute of Technology. Government sponsorship acknowledged.

ASHWIN R. VASAVADA is a planetary scientist at the Jet Propulsion Laboratory and has been the deputy project scientist on the Mars Science Laboratory mission since 2004. His research interests include the climate history of Mars, polar ices on the moon and Mercury, and the meteorology of Jupiter and Saturn.



ESA, NASA, AND THE INTERNATIONAL SPACE STATION

BY ALAN THIRKETTLE

From the 1970s until the end of the twentieth century, the European Space Agency (ESA) and NASA cooperated on a range of human space flight programs. At the start of that relationship, ESA had no experience in human space flight whatsoever, while NASA had been through several programs culminating in Apollo and Skylab, and the post-Apollo Space Transportation System (STS, the Space Shuttle) was already approved. ESA committed to develop the Spacelab modular system, the scientific laboratory of the shuttle. In this very successful program, ESA and the European industry learned how to develop and qualify human spacecraft, started an astronaut program, and built many experiments in numerous scientific disciplines.



A view of the European Columbus laboratory installed in its new home on the International Space Station. Columbus was launched with Space Shuttle Atlantis in February 2008.

Photo Credit: ESA/NASA

As a result of these ventures, ESA became a qualified human space flight partner. Cooperation with NASA was smooth, controlled, and of mutual benefit: ESA for the learning curve and the end product, NASA for the expansion of the shuttle system into an operational scientific platform.

The International Space Station (ISS) is so far the largest example of international cooperation, certainly in the aerospace world and arguably in the engineering world as a whole. The fact that decisions have to be taken among partners has sometimes been a source of difficulties, but the benefit is clear to see, in terms of both the ISS itself and the understanding developed among the players. Lessons learned from past and present ventures need to be taken into account in setting up the framework for cooperation on future exploration.

The ISS Invitation

In the mid-eighties the ESA member states initiated a series of new programs: the Ariane 5 heavy lift launcher, the Hermes manned space plane, and the Columbus program. The latter consisted of

several infrastructure elements: a pressurized module, a manned free flyer (to be serviced by Hermes), a service vehicle, and a polar platform. This infrastructure, associated with the NASA-led space station program ("Freedom"), was conceived as a combination of cooperation and autonomy for Europe. U.S. President Ronald Reagan had invited partners to join NASA in working on such a station, and on the basis of the Spacelab experience, ESA accepted the invitation. While Ariane 5 was fully implemented and remains a world leader in carrying cargo to space, Hermes was eventually canceled and the infrastructure program converged to the Columbus laboratory module in the face of economic realities.

As the station program developed, the basis of cooperation—which also embraced Japan and Canada—became apparent. It is a multilateral venture controlled by a single intergovernment arrangement, beneath which are bilateral implementing agreements between NASA and each of the partners. The overall system architecture and design are under NASA's leadership with agreed-upon standard documentation, verification procedures,



Astronaut Daniel W. Bursch (left) and cosmonaut Yuri I. Onufrienko, Expedition Four flight engineer and mission commander respectively, wearing Russian Sokol suits in the Soyuz 3 spacecraft that is docked to the International Space Station.

and common hardware elements (such as berthing mechanisms, hatches, module diameters, voltages, and interchangeable payload racks). The partners provide “independent” elements: a robotic system in the case of Canada and module/platform elements in the case of Japan and Europe. These elements interface with the centralized resource systems (power, communications, heat rejection, atmospheric conditioning, habitation, etc.) provided by NASA.

Station resources, including crew time, are shared among the partners in accordance with negotiated relative contribution evaluations. In exchange, the partners pay a corresponding share of the common system operating costs. Partners have their own operational control centers (under NASA oversight) and their own utilization programs. In exchange for provision of all the resource systems, NASA has usage rights to 50 percent of the partner elements.

Cooperation Evolution: Russia Enters into Partnership

After the end of the Cold War, discussions between the United States and Russia led to the introduction of Russia as a new partner on the station. This coincided with a major redesign of the architecture. The existing partners were invited to participate in the changes and given considerable technical visibility, although decisions were largely a bilateral matter between the two major players.

This expansion had a number of consequences, some causing concern and others clearly beneficial. The station architecture was no longer seamless but consisted of two distinct halves—the so-called United States Orbital Segment (USOS), which included the elements of the original partners, and the Russian

segment. Different voltages, different life-support systems, separate logistics, and nonstandard hardware (including hatches and berthing mechanisms) added to the complexity. On the positive side, the station gained the experience that the Russians brought—they had operated space stations including Mir for many years—the robustness of alternative transport systems (Progress and Soyuz), and more balanced power sharing between them and NASA. (Eventually, of course, the station was kept alive by the Russian transport systems following the *Columbia* tragedy. Without Soyuz and Progress, it would certainly have had to be de-manned and potentially abandoned completely.)

Aspects of ESA's Cooperative Participation

The main ESA contribution, the Columbus laboratory, is a sixteen-rack pressurized module with four external payload attachments. It was launched in February 2008 aboard STS-122. Rather than paying cash for the flight, ESA provided two of the interconnecting nodes of the station, plus other goods and services equal to the value of the launch.

Having been one of the last partners to fully commit to the station, the ESA Columbus module (the full development of which was only finally confirmed at the end of 1995) was to be the last flight in the assembly sequence. This meant a long wait for the European science community, so negotiations for early European use of the station, including the associated launch and retrieval of payloads and short-duration astronaut flights, were conducted. In exchange, ESA undertook to provide the station's Cupola, a Microgravity Science Glovebox, and a -80 degree freezer system, the so-called Melfi. Another agreement with the Japanese agency led to ESA receiving payload rack structures from Japan in exchange for delivering another Melfi to our Japanese partner. ESA also negotiated a progressively earlier slot in the assembly sequence. The February 2008 launch was some two-and-a-half years ahead of the last assembly sequence launch.

The European share of the common systems operating costs, owed to NASA, is also being “paid in kind” by the provision of

LESSONS LEARNED FROM PAST AND PRESENT VENTURES NEED TO BE TAKEN INTO ACCOUNT IN SETTING UP THE FRAMEWORK FOR COOPERATION ON FUTURE EXPLORATION.

flights of the Automated Transfer Vehicle, or ATV, a twenty-ton spacecraft launched from the spaceport in Kourou, French Guiana. The ATV can transport up to eight tons of useful cargo to the ISS and can remove several tons of trash in its destructive reentry. The first of these missions took place from April to September 2008.

There have been several European astronaut flights to ISS to date, both short-duration and longer “increment” flights. 2009 will see a further increment flight, during which Frank de Winne will assume command of the station for the last two months of his stay on board, the first non-American/Russian commander of the ISS.

The control centers of Columbus and the ATV, in Oberpfaffenhofen (Southern Germany) and Toulouse (France) respectively, are both fully commissioned and operated by European ground crews. A utilization program is under way in domains including life science, fluid physics, materials science, solar physics, space technology, and industrial exploitation, and will grow as time passes. Many experiments are conducted jointly among the partners in the different areas of the station, continuing a long tradition of scientific cooperation.

Reflections

For ESA, the Spacelab cooperation was straightforward inasmuch as we were clearly not an equal partner, either in magnitude of task or in experience. While the magnitude of our overall participation is that of a small contributor (measured as 8.3 percent quantitatively), there is closer equality in development competence, as evidenced by the amount of European hardware on board. Operationally, we are comparative newcomers in manned space, but there is consensus that Europe has produced some of the “best ever” astronauts. In order to acknowledge this change in circumstances, an increased degree of consensus management has been appropriate and necessary.

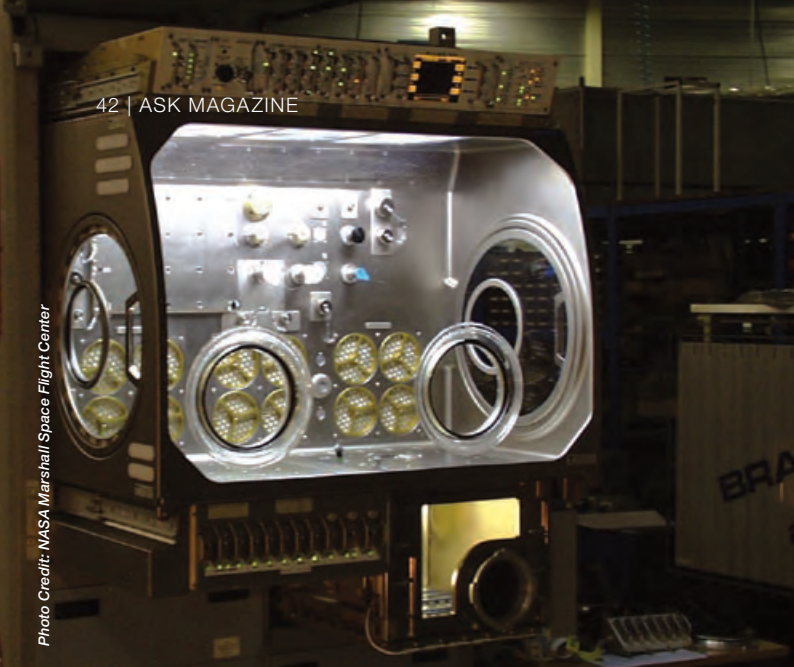
There have been periods, however, when cooperation was more difficult than it should have been. At the turn of the

century, NASA, under pressure from the U.S. government following large increases in the projected cost of the station, took a critical internal look at its expenditure and ambitions and, as a result, canceled a number of elements of the station, including the habitation module and the crew rescue vehicle. While the partners were offered “visibility” into the review, which lasted for more than two years, they had no real opportunity to influence its outcome. The decisions were made unilaterally, to the consternation of the other partners.

Europe had spent more than \$100 million on cooperation with NASA on the crew rescue vehicle and its predecessor, the X-38. There was no compensation for this investment when NASA canceled that part of the station program. This was a time when the “partnership” was disregarded in favor of the wishes of the United States. Relationships were so difficult for a time that some European governments called the European participation into question. It only survived because the top governing document, the intergovernment arrangement (IGA), has the status of a formal treaty for the European participants. The situation was aggravated by the subsequent loss of *Columbia*, which added three years of delay to the launch of the European module and hence extra (unforeseen) costs that stretched our ability to continue the program.

Ironically, the loss of *Columbia* probably led to the beginning of reparation of relationships, because all partners in the human space community feel a common cause in the recovery from such a disaster. By 2004 it was clear that ESA would suffer long delays in the launch of Columbus, and we were desperate to advance the launch of our module to an earlier slot in the assembly sequence. Working groups involving all partners were established to look at the feasibility of this.

NASA engineering and operations communities expressed reluctance, but NASA management decided that the benefit to the partnership outweighed the technical challenges; an advance of the Columbus launch in the sequence was agreed upon by all participants. This was an example of true



Interior lights illuminate the Microgravity Science Glovebox, developed by the European Space Agency and NASA for use aboard the International Space Station.

cooperation and partnership spirit. Today the relationships are as good as they have ever been. This is a function of the state of advancement of the program—it always helps when all partners are on orbit—but also of the people running the program in the various agencies and industries. Personal relationships have been forged that will form the bedrock of future cooperative ventures, notably in the exploration of the moon, Mars, and beyond.

One recent example of cooperation had to do with the computer system in the Russian service module “crashing” in 2007. Russian segment electrical power was lost, and there was a threat that the crew would have to leave the station before the batteries of the Soyuz rescue vehicle depleted. Europe had designed and developed the computers for the Russian partner, and so the European, American, and Russian engineers had to work very quickly and cooperatively, overcoming time differences and physical distances, to retrieve the situation. Although there were strong sensitivities (Europe was—rightly—convinced that “their” computers had not caused the problem, Russia did not want their service module to be seen as the culprit, and NASA did not want their overall system leadership to be called into question), everyone worked extremely closely and found the resolution in a very short time. NASA re-routed power from the USOS to the Russian segment and took over station attitude control. ESA sent spare parts and engineers to Russia to help troubleshoot, Russian engineers scrutinized their designs to establish root cause, and the onboard astronauts and cosmonauts helped with real-time hardware evaluation. This was an example of cooperation happening thanks to a common set of clearly defined, obvious objectives.

Cooperation is a way to achieve mutual objectives in an affordable, non-competitive manner. For huge programs it is the only realistic route to achievement. But there are good and bad paradigms for cooperation. We should strive for balance in the partnership. When partners need each other’s contributions—that is, when the elements are interdependent—then the whole

is greater than the sum of the parts, and the relationship is based on a necessarily equal footing. On the other hand, when one partner contributes something that is “nice to have,” rather than essential, the role of that partner in the decision-making process will suffer when times are hard.

The mutual dependence of America and Russia on the ISS, and between ESA and NASA on STS/Spacelab, are examples of mutual need enabling serious problems to be solved cooperatively. Now that all ISS partners have hardware elements in orbit, successful operations depend on overall cooperation, so the bedrock of mutual dependency is there. The strong institutional and personal relationships established during the good and bad years between the various stakeholders are flourishing in that climate of partnership. ●



ALAN THIRKETTLE is the former ESA ISS program manager.

Viewpoint: Building a National Capability for On-Orbit Servicing

BY FRANK J. CEPOLLINA AND JILL MCGUIRE

NASA's Goddard Space Flight Center pioneered satellite servicing more than thirty years ago with the historic Solar Maximum Repair Mission. That mission laid the groundwork for four successful servicing missions to the Hubble Space Telescope (HST) and the development of the HST Robotic Servicing and De-orbit Mission (HRSMD). Separately and independently, the Department of Defense (DoD) has been refining its capabilities to refuel and reconfigure on-orbit satellites robotically and autonomously. By collaborating, NASA and DoD can validate their technologies and prepare more efficiently for future space robotics missions. A national system for on-orbit servicing would provide an important capability for a broad range of missions.

NASA and DoD have been collaborating since 1958. NASA's Space Shuttle program has flown eleven dedicated missions for the DoD, providing access to more than 250 secondary DoD payloads, as well as transportation to Mir, the Russian space station. Other successful collaborations include the National Polar-Orbiting Operational Environment System, a tri-agency program that merged the DoD and Department of Commerce—National Oceanic and Atmospheric Administration polar-orbiting weather satellite programs, and Clementine, a DoD–NASA technology-demonstration mission to survey the moon. Another collaboration, one often forgotten, is the sharing of vehicles in the Evolved Expendable Launch Vehicle program; in fact, NASA developed the Atlas and Delta launch vehicles that the military has relied on for the past twenty years.

Although NASA and the DoD may use some technologies differently, and in some cases DoD work is classified, the two organizations deal with many of the same issues, technical challenges, and requirements. Propulsion, materials, avionics, and launch technologies are ideal areas of joint technology development. Even the NASA Mars exploration missions have benefited from technical collaboration with the U.S. Air Force, using the air force–developed RAD6000 32-bit microprocessors and lithium ion batteries for both planetary rovers. For the most part, NASA and DoD personnel can team together on technology

Orbital Express during on orbit testing.

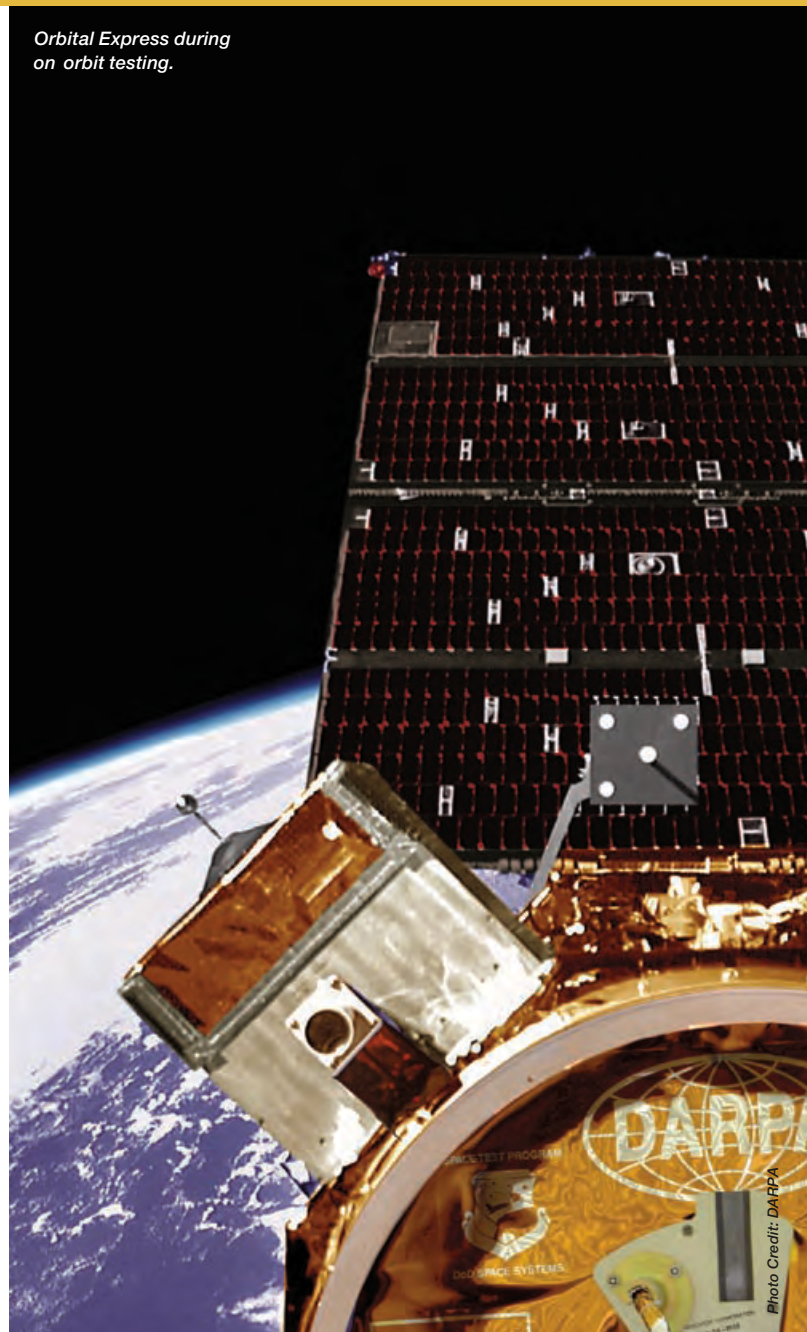
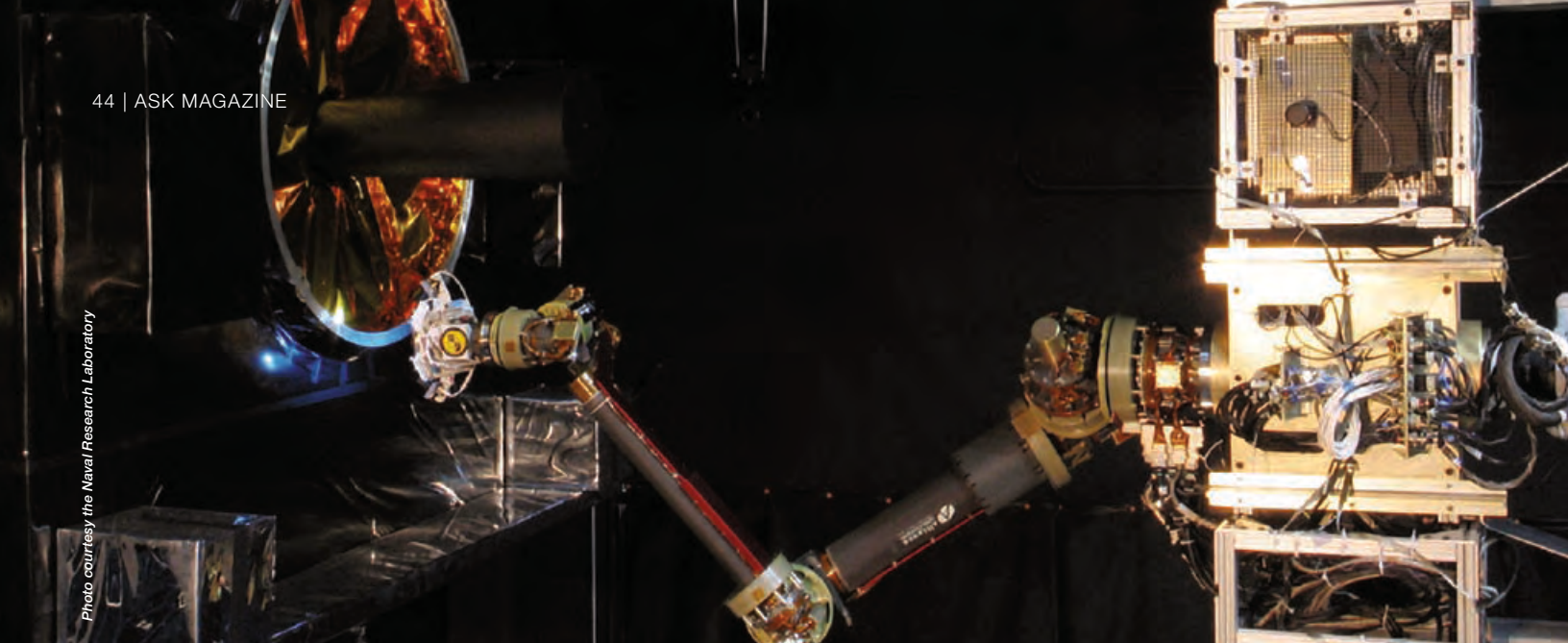


Photo Credit: DARPA



development or basic research. After the development phase, the DoD will transfer many of the technologies to be used on classified programs. Whenever the organizations' goals coincide—for instance, on ways to maneuver, communicate, experiment, or work in space—collaboration leads to the optimal access to equipment, processes, and procedures for both organizations. The collaborations can maximize resources in many areas, making optimal use of the American taxpayers' investment.

Learning from Hubble and Orbital Express

Goddard's year-and-a-half effort on HRSDM involved rapidly developing a spacecraft with a robotic grapple arm, a two-armed dexterous robot, a vision system, twenty-four robotic tools, robot-compatible Orbital Replacement Units, and ground stations to support the robotic operations. Considerable effort was spent on detailed kinematic and dynamic work-site analysis, operational scenarios, and ground and neutral buoyancy evaluations using protoflight hardware.

As a result, Goddard has developed broad capabilities, technology, and space robotics expertise in autonomous rendezvous and capture, dexterous robotics, end effectors (the working devices at the end of robot arms) and tools, teleoperation, vision/situational awareness, assembly and servicing automated tools, and robotic-controlled de-orbit. Much of that technology will be used for the Hubble servicing mission.

While this work was going on at NASA, Defense Advanced Research Projects Agency (DARPA)—the DoD's technology research center—was creating the infrastructure to support the Orbital Express mission. Orbital Express demonstrated the technical feasibility of robotic, autonomous, on-orbit refueling and reconfiguration of satellites to support a broad range of future U.S. national security and commercial space programs. DARPA led the development of two new spacecraft, a robotic grapple arm, robot-compatible Orbital Replacement Units, refueling tanks, support hardware, and ground stations.

Building on that highly successful mission, DARPA developed

the Front-End Robotic Enabling Near-Term Demonstration (FREND), which includes the robotic arm, sensor suite, and algorithms that have successfully demonstrated advanced, autonomous, unaided grappling of a simulated spacecraft with no a-priori knowledge of the spacecraft and with no standard, robotic-friendly targets or interfaces on the spacecraft.

These separate but complementary NASA and DoD activities validated technologies and sowed the seeds for collaboration during DARPA's Orbital Express mission and FREND activities, when Goddard experts were brought in to provide technical expertise and oversight. During Orbital Express, DARPA management entrusted NASA personnel with technical oversight in areas including relative navigation and advanced video guidance sensors, robotic systems, and mission operations. When necessary, NASA personnel were able to get the appropriate security clearances. Because the robotic system used on Orbital Express had many commonalities with the system developed for Goddard's HRSDM, the contractor providing the system applied several lessons learned from HRSDM to develop a simpler architecture and optimized control system for Orbital Express. It became apparent that because NASA and DoD had many common capabilities, the Orbital Express mission could have been performed sooner if the two agencies had collaborated earlier. Also, lessons learned from Orbital Express can be applied to future collaborations on in-space servicing.

Future Collaboration

Because the work accomplished at Goddard in the areas of robotic servicing and relative navigation is synergistic with much of the work done at DoD, it would be beneficial to combine efforts and funding resources for future in-space NASA and DoD space robotics missions. One potential symbiotic relationship is on the FREND program, as DARPA is looking for a mission partner who can build the spacecraft onto which the FREND robotic arms and sensors can be mounted.

Little work has yet been done to optimize FREND's end

The Front-End Robotic Enabling Near-Term Demonstration underwent full-scale rendezvous and autonomous robotics grapple testing at the Naval Research Laboratory's Spacecraft Proximity Operations Test Bed.

Artist's concept of Hubble Robotic Servicing and De-orbit Mission architecture.



Image courtesy Frank J. Cepollina

effectors, something Goddard spent considerable time doing during HRSMD. In addition to being able to integrate and test the FREND mission in Goddard's facilities, the center has considerable expertise planning on-orbit servicing operations, dating from the Solar Maximum Repair Mission in 1984 to current Hubble repair and upgrade missions. A collaborative effort incorporating this on-orbit experience and the lessons learned from a NASA–DoD FREND mission into the design and development of an Orion-type service module that can operate either in an astronaut environment or an autonomous/robotic environment would create a new national capability for on-orbit servicing of space assets that could even be applicable to commercial communications satellites.

No one should downplay the challenges of collaboration seen in the past. In the Evolved Expendable Launch Vehicle program, for example, both the Boeing Delta IV and the Lockheed Atlas V had to undergo separate DoD and NASA certification processes. This is currently the case with the SpaceX Falcon launch vehicle.

SpaceX underwent a review by Aerospace Corporation for the DoD and a separate review by NASA. Programs funded by multiple agencies present challenges when design trade-offs are required. These trade-offs lead to the accumulation of excessive requirements, one of the biggest risks of NASA–DoD collaborations.

Furthermore, although the Clementine mission was a successful low-cost technology demonstration that surveyed the moon and flew past an asteroid, a recent lessons learned study of the mission by the National Research Council found different “cultures” operating within the DoD and NASA. Differences included greater resources available to DoD than to NASA, a sense of urgency for military projects as compared with a more leisurely pace of civilian programs, less involvement by Congress and reduced micromanagement by the DoD, and a more focused, task-force-like management style at the DoD that contrasted with the broad, participatory approach

associated with NASA missions. But all these challenges have been overcome on successful collaborations in the past and can be again as long as teams of highly motivated and dedicated professionals remain aware of the them.

The ultimate goal of a NASA–DoD collaboration is to enhance the United States' ability to close the gap between concepts and implementation while reducing cost. A combined effort would maximize the United States' limited robotics research funding at a time when it trails countries such as Japan, Germany, and Russia in this area. This collaboration would provide benefits including servicing, upgrade, and refueling opportunities for expensive NASA, DoD, and other strategic assets in low- and high-Earth orbits, as well as existing commercial communications satellites in geosynchronous orbits. It would also help facilitate the needed merger between human and robotic space flight, establish the United States' preeminence in robotic space flight, and maintain U.S. dominance in human on-orbit servicing. ●

Note: While Frank J. Cepollina is the project manager in charge of Hubble Space Telescope servicing, this article reflects his personal (not official) ideas about the potential for establishing a post-2014 national U.S. capability for repair upgrade and refueling of spacecraft in high-usage orbits.

FRANK J. CEPOLLINA serves as deputy associate director for the Hubble Space Telescope Development Project at Goddard Space Flight Center. He is known as the “Father of On-Orbit Servicing” for his decades of leadership in repairing and upgrading satellites in orbit. In 2003 he was inducted into the National Inventors Hall of Fame.



JILL MCGUIRE is the Robotics and Crew Aids and Tools manager for the Hubble Space Telescope program. She has been at Goddard Space Flight Center since January 1992 and with the Hubble program since 1998.



BUILDING THE TEAM:

The Ares I-X Upper-Stage Simulator

BY MATTHEW KOHUT

The opportunity to build a new launch vehicle that can loft humans into space does not come along often. The Ares family of launch vehicles, conceived in response to the Vision for Space Exploration, presented the first chance for NASA engineers to get hands-on experience developing human space flight hardware since the development of the Space Shuttle thirty years ago.

In 2005, NASA Headquarters solicited proposals from integrated product teams, or IPTs, for different segments of the Ares I X test flight vehicle. The objectives focused on first stage flight dynamics, controllability, and separation of the first and upper stages. The launch vehicle would comprise a functional booster stage and an upper stage mass simulator, which would have the same mass as an actual upper stage but none of the functionality.

A team at Glenn Research Center prepared to bid for the job of building the Ares I X Upper Stage Simulator (USS). The

first challenge in bringing the project to Glenn was assembling a core team with the right skills to develop a winning proposal.

Even before we had gotten authority to proceed but were doing concept studies, and cost and schedule estimations, I needed a good systems engineer to look across this conceptual simulator that we were coming up with and help us identify if we were missing any functions, said Vince Bilardo, who headed the proposal team and would eventually become the project manager. We needed a good systems engineer to help us create a draft functional allocation.

Workers unwrap segments of the upper stage simulator at Kennedy Space Center. The segments were built at Glenn Research Center and shipped to Kennedy for assembly.

Photo Credit: NASA/Jim Grossman

As the proposal development period for the upper-stage work progressed, Bilardo drafted Bill Foster to serve as his lead systems engineer. Foster began attending systems engineering technical interchange meetings while Bilardo ran concept teams that drew up a series of designs ranging from high fidelity and expensive to low fidelity and inexpensive. “[Vince] had different teams laying out concepts, and that’s where the first ‘tuna can’—the design we actually ended up with—came up. He ran those concept teams over a three-day period, and that kind of kicked everything off,” Foster said.

The Glenn team continued to define its concepts and cost estimates as the Constellation program developed the requirements for the test vehicle. “The requirements were pointing us toward a higher-fidelity simulator. So some of our concepts started to fall to the wayside while the higher-fidelity one—the expensive one—was really the only one that was going to pay off,” said Foster. “When we rolled the cost all up and Constellation was figuring out their budget, they said, ‘We’re not doing high fidelity because it’s way too expensive.’”

A few weeks later, Glenn came back with a trimmed-down version of its low-fidelity proposal. “This low-fidelity launch did a few things. One, it gave us good flight data about whether we could launch this long, skinny rocket. Second, it was fairly inexpensive. Third, it was going to be an early launch [2009] to get this early data, whereas the high-fidelity version pushed out into 2011,” said Foster. “That’s what got us turned on [approved]. At that point we started ramping up people.”

An In-House Development

In May 2006, the Glenn team received provisional authority to proceed with the USS as an in-house project, meaning the Glenn team would design, develop, and build the hardware in its own facilities using its own technical workforce, rather than contracting the job out to private industry. After a probationary period, the project got full authority to proceed in August.

The selected design required manufacturing eleven segments of half-inch-thick steel that stretched 18 ft. in diameter and 9.5 ft. tall—the tuna-can shapes that gave the simulator its nickname. The job would incorporate all the basic hardware development functions: cutting, rolling, welding, inspecting, sandblasting, painting, drilling, and tapping for instrumentation.

Since the project team was beginning with no in-house expertise in large-scale fabrication or manufacturing, it required an entirely new set of procedures that documented each step of the building and assembly process in detail. Bilardo called on Dan Kocka, a recently minted engineer who had spent most of his career as a technician, to serve as the production-planning lead. Kocka had never assumed these duties before. “I said to Vince, ‘Are you sure I’m the right guy for this?’” Kocka recalled. Bilardo was confident that Kocka’s unique background would be an asset that would outweigh his relative lack of experience. Once the project got under way, Kocka’s doubts dissipated; he found that his ability to think like an engineer and a technician served him well. “I really was in a good position to have both of these things going on at the same time in my mind,” he said.

An additional challenge that fell heavily on Kocka concerned demonstrating compliance with AS 9100, an aerospace manufacturing quality standard. Glenn’s management team was making a centerwide effort to achieve AS 9100 certification. For the USS, this meant putting in place rigorously documented procedures that met with the approval of both the Safety and Mission Assurance organization and the technicians doing the work. The AS 9100 standard added another level of rigor to the process of designing and building space flight hardware.

Preparing for a fabrication job of this size and scope also demanded a wholesale renovation of a facility: new cranes, new assembly platforms, and a new sheet-metal roller. This meant retrofitting an older manufacturing shop floor that was large enough to accommodate the hardware. The facility modification had to be done quickly—in about three or four months—so the project could begin work as scheduled.



Bob Peters, of Kendel Welding and Fabrication, welds part of an internal access support for the Ares I X upper stage simulator at Glenn Research Center in Cleveland.

Photo Credit: NASA/Quentin Schwinn (RS Information Systems, Inc.)

“From a project management perspective, I needed somebody who could handle all that facilities work and go work with the facilities organizations and the facilities directorate at Glenn to start planning, designing, and implementing the overhauls that we needed to accomplish in a very short period of time,” said Bilardo. He found Jack Lekan, an experienced project manager who was finishing up another job at the time. “Jack was a long-time Glenn guy who had excellent contacts across the center and a skill set that was very much oriented to team building and cooperation and working well with the various performing organizations. [He was] perfect for the job.”

Ramping Up

Managing the USS project required constant interaction with the other three NASA field centers responsible for Ares I-X: Langley Research Center (systems engineering and integration office), Marshall Space Flight Center (first stage, avionics, and roll control system IPTs), and Kennedy Space Center (integration

and test functions as well as the launch itself). Bilardo spent a significant amount of time traveling or otherwise coordinating with his counterparts at these centers, so he needed a deputy project manager who could handle the “down-and-in” details of running the project on a daily basis.

He turned to Foster, his lead systems engineer, who had project management experience from his years on microgravity science projects where he’d served as both the project manager and systems engineer. With Foster moving over to project management, the team needed a new lead systems engineer. They brought in Tom Doehne, who was just finishing up a trade study for the upper-stage thrust vector control system of the Ares I vehicle.

Doehne’s primary focus was on managing the design integration of the simulator hardware, documenting the design in the Design Definition Memorandum, and developing the project requirements. The design evolved and the requirements database kept expanding as the larger Ares I-X management team kept adding requirements for the USS. Doehne realized he needed more systems engineers to support the project. “Initially, I was the only systems engineer as we were developing this task, and we had a lot of work up front that we were trying to do in the early July–November [2006] time frame,” he said.

As the systems engineering workload increased leading up to a Systems Requirement Review, Doehne had trouble finding qualified systems engineers since the new Orion and Ares I projects at Glenn had been ramping up during the past year. Eventually, he was able to transition two civil servants who were in the Space Mission Excellence Program as well as some experienced contractor support. “We took qualified engineers from other areas of the center who were in training as systems engineers. They received real project experience, and we were able to complete the large volume of work that was in front of us,” Doehne said.

In addition to knowledge and experience, Doehne valued team members who could remain engaged and flexible on a project with an aggressive schedule and a rapidly changing context. “Team dynamics is also a very important key to building a successful project team and shouldn’t be mistaken for something that isn’t needed,” he said. “In today’s projects with limited budgets and aggressive schedules, we need to work as a cohesive team unit and have the ability to adapt to a dynamic work environment to achieve our common goal.”

“Welding Is Not Easy”

The scale of the USS demanded a manufacturing capability that didn’t exist at Glenn. The recent focus of the center’s manufacturing efforts had been on microgravity payloads that called for highly intricate machining of sophisticated instruments, not on the rough fabrication skills needed to roll, weld, and attach large segments of a launch vehicle. This

reorientation toward heavy manufacturing posed challenges both in terms of the workforce and the organization.

Glenn had several highly skilled machinists among its civil service workforce, but it had few fabricators and a critical shortage of welders. Since the center no longer had enough work to fully utilize the majority of its machinists, the project management team, in consultation with Glenn's upper management, set out to retrain a cadre of about twenty machinists as welders.

The effort was well intentioned, but it did not work as planned. It turns out welding is not easy, said Foster, who had advocated for hiring outside welders. Even with training, it took years of practice as a welder to achieve the level of proficiency that the job demanded. Flight quality welds must pass a litany of tests, including radiographic and ultrasonic inspections by certified weld inspectors. Unless you are welding day in and day out for a living, it's really difficult to maintain the level of skills required to execute flawless welds that are going to fly on a flight test for NASA," said Bilardo.

The next step was to hire welders on contract. The project reached out to some local non union shops, which began sending over welders for qualifying tests. Again, the necessary skill level proved to be a formidable barrier. They were probably washing out at a 60 percent rate, said Foster. The project retained the services of only one of these shops, and they still needed more welders. A call went out for union welders. Even then, with top notch welders, we were getting about a 25 percent washout rate," said Foster.

The drawn out hiring process cost the project time that it hadn't built into its schedule. Having found enough qualified welders, the project then had to align the number of welders on the shop floor with the work flow. We needed a lot of welders at the beginning, but then we cut back because we were not able to keep them busy, said Foster. Then the pendulum swung too far. We cut back too far because things were going well and we hadn't gotten into the complicated segments. When we got into the complicated ones and the welding went back up, all of a sudden we needed welders again.

The juxtaposition of union and non union welders in the manufacturing facility created other issues. Some union welders did not want to work alongside their non union counterparts; in one instance, union welders walked off the job. We ended up deploying them [union and non union personnel] on different segments so they didn't have to rub elbows on the same build stand, said Bilardo.

Getting Smarter

There was probably no way to foresee the welding problems; the team learned by doing. As the USS entered the home stretch of fabrication before preparing to ship segments to Kennedy Space Center for integration and testing, Bilardo reflected on

the difficulties the project had encountered: About two thirds of our [cost] growth is due to requirements and scope growth and events outside our IPT, and about one third has been within our control and really attributable to what I would call maturing estimates. There are things that you now put in your budget that you couldn't have guessed that you needed. Or you did guess at it, and you guessed low, because you weren't smart enough. You just get smarter over the course of the project. ●

Ares I X construction work area at Glenn Research Center.



Photo Credit: NASA Glenn Research Center

Things I Learned on My Way to Mars

BY ANDREW CHAIKIN

When I was five, in 1961, one of my favorite books was called *You Will Go to the Moon*. It tells of a young boy who accompanies a team of astronauts on a lunar voyage. (*How* he gets to do this is never explained.) After landing at a lunar base, they put on space suits and go bounding across a bright, cratered moonscape. At the end of the story, they climb to the top of a tall hill and look into the starry sky, where they spot Mars, beckoning them still farther into space. “Mars is a long, long way from the moon,” says the narrator. “What would you find on Mars? No one knows yet. But someday you may go there, too! Then YOU will see.”

This Mars Global Surveyor Mars Orbiter Camera image shows a suite of south mid latitude gullies on a crater wall that may have been formed by runoff of liquid water.

Photo Credit: NASA/JPL/Malin Space Science Systems

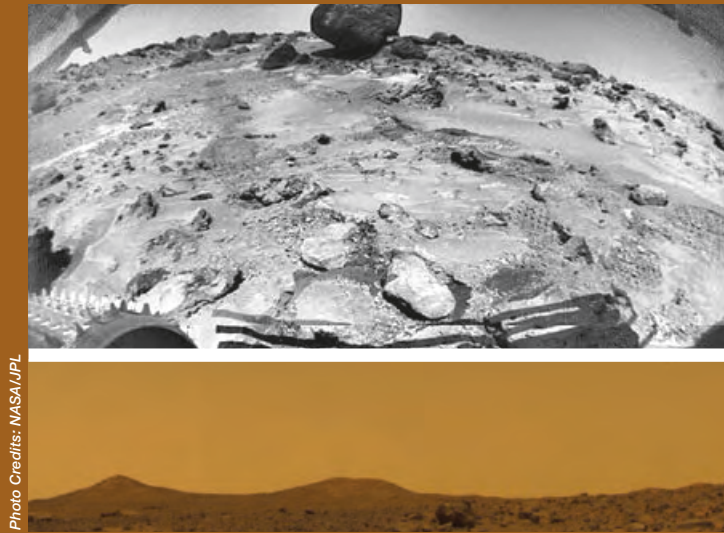


Photo Credits: NASA/JPL

This Sojourner image, taken on Sol 70, shows rocks and rover disturbed soil. The large rock in the distance is "Yogi. Much of Yogi visible in this image cannot be seen from the perspective of the Pathfinder lander.

The true color of Mars as seen by Pathfinder.

As much as I wanted to believe the promise, "you will go to the moon," I'm still waiting for my chance. And no one has been to Mars, which shines like a beacon on a dark and distant shore, awaiting the first footsteps of human beings. But I'm not complaining. I feel incredibly fortunate to have witnessed one of the most compelling sagas of the Space Age: the transformation of the red planet from a vague apparition in astronomers' telescopes into a world brimming with wonders and mysteries. I even got to take part in that adventure as an undergraduate intern at the Jet Propulsion Laboratory (JPL) on the first Viking landing in 1976. Since then, I've covered Mars exploration as a science journalist, and the red planet has lost none of its allure. For my book, *A Passion for Mars*, I talked to some of the people who have been the driving force behind this epic quest. What they told me includes some fundamental lessons—about exploration and about life.

Check your hubris at the door. From the surprisingly moonlike surface seen in Mariner 4's first crude close-ups of the planet in 1965 to the geologic wonderland unveiled by the Mariner 9 orbiter in 1972, Mars has surprised scientists every time they've been able to see it in new detail. That was especially true for Caltech geologist Bruce Murray, who served as a junior member of the Mariner 4 science team and went on to become one of the leading experts on Martian geology. After three flyby missions showed a cratered, apparently unevolved landscape, Murray was not prepared for the world revealed by Mariner 9's orbital survey, brimming with giant volcanoes, vast canyons, and ancient dry river valleys. And the surprises didn't end there; they continued with every new mission. Three decades later, a new orbiter called Mars Global Surveyor, outfitted with a high-powered camera designed by Murray's former student Mike Malin, was sending back spectacularly detailed views that revealed a host of baffling new surprises. By 2001 Murray was calling Mars "the land of broken paradigms." If there is one thing Mars teaches us, it is that when we think we have it all figured out, we'd better think again.

Learn to be comfortable with ambiguity. One of the reasons Mars has held such a powerful grip on the human imagination is the enduring question of whether it harbors some form of life. But even after more than forty years of exploration by spacecraft, the answer remains elusive. In 1976 the twin Viking landers arrived on Mars carrying a trio of experiments designed to search for microbial life, but the data they sent back suggested not life but a surprisingly exotic surface chemistry laden with highly reactive compounds called peroxides. Meanwhile, an onboard gas chromatograph/mass spectrometer failed to detect organic molecules at a parts-per-billion level, suggesting that the very dust of Mars is hostile to life. In the end, the Viking results failed to settle the life-on-Mars question one way or the other, much to the chagrin of many journalists covering the mission.

And yet, to the media and the public at large, Viking's failure to detect life was translated into a certainty that Mars is a dead planet. And so things stood for twenty years, until a team of NASA scientists claimed to have found evidence for fossil bacteria inside a Martian meteorite. That claim has remained controversial, but the ensuing discussions had the effect of dragging the entire subject of exobiology from the fringe to the mainstream. And more recently, evidence that water may flow beneath the surface today in subterranean aquifers has raised anew the possibility of Martian biology. True resolution of the question may not be had until astronauts go to Mars to conduct detailed on-site explorations. Until then, we can all take our cue from what Carl Sagan told frustrated members of the press late in the Viking mission: "What I would urge on you is an increased tolerance for ambiguity."

Be prepared to go against the flow. In the mid-1980s, geologist Malin faced resistance from his colleagues when he proposed an extremely powerful camera for an upcoming Mars orbiter; the other scientists claimed they already had all the pictures of Mars they would ever need. Malin was convinced that an unknown Mars lurked below the limits of resolution, and after winning approval from NASA, he led the development

This false color image taken by the Mars Exploration Rover Opportunity's panoramic camera shows that dune crests have accumulated more dust than the flanks of the dunes and the flat surfaces between them.



Photo Credit: NASA/JPL/Cornell

of a lightweight, low-cost instrument called the Mars Observer Camera (MOC) that went to Mars aboard the ill-fated Mars Observer, which was lost shortly before entering Mars orbit in 1993. Malin had to wait another four years before a second MOC finally reached Mars aboard Mars Global Surveyor in 1997. And, just as he'd hoped, his camera discovered a Mars no one expected, where entire landscapes have been buried and exhumed over billions of years; where strange formations defy explanation; where climate change is now apparently under way. It is a world of ancient lake deposits and fossil river deltas, where signs exist that water has flowed across the surface in the recent geologic past—and, perhaps, even today. Malin is just one of countless Mars explorers who challenged conventional wisdom and were rewarded with spectacular discoveries.

Do the right thing. Exploring Mars is at the very edge of what we humans know how to do, and each new mission has tested our ingenuity. From 1968 to 1976 engineer Gentry Lee helped lead the engineering effort behind the billion-dollar Viking mission. Engineers were charged with developing an onboard computer to steer the landers to a soft touchdown, a biology package with three separate experiments crammed into a one-foot cube, and countless other cutting-edge components. To make everything more difficult, NASA required that the spacecraft undergo a grueling heat sterilization process before launch, in order to avoid bringing terrestrial bacteria to Mars. Seeing these challenges solved—and seeing the first pictures from the surface of Mars while helping to direct the mission at JPL—made Viking a high point in Lee's life. He came out of that experience with an understanding of what it takes to do cutting-edge exploration, as exemplified by the hard-as-nails Viking project manager, Jim Martin. In Lee's words, Martin practiced "an absolutely ruthless technical Darwinism. If you were not capable of doing your job, you were gone."

A quarter-century later, in 2001, Lee returned to JPL to help oversee the young engineers working to create the Mars Exploration Rovers Spirit and Opportunity. The landing system

for the rovers was nothing like Viking's; they would "bump down" onto Mars cushioned by a set of giant airbags, just as the Mars Pathfinder lander had in 1997. But the rover teams struggled as they tried to adapt the Pathfinder system to the larger and more massive rovers. In testing, there was one technical showstopper after another, from failed parachutes to shredded airbags. Meanwhile, Lee and other Mars veterans subjected the young engineers to harsh reviews, threatening the project with cancellation. The team responded by digging in to redesign the parachute and the airbags; then, just months before launch, they created a system of small, computer-controlled rockets to counteract the Martian winds during the final descent. Even after the rovers were safely on Mars, the crises didn't end; most harrowing was Spirit's temporary loss of sanity due to computer problems just days before Opportunity was due to land. But each time, the teams stepped up with the same commitment, walking the high wire between success and failure, and saving the mission. In the end, Lee says, "The whole story [of the Mars Exploration Rovers] is one of heroism, where people would not accept the risks of not changing something and would accept the risk of the hours and hours of new work that would have to be done in order to make it right."

Be in it for the long haul. Mars explorers must be driven to accomplish their goals, but they must also be willing to endure delays and disappointments. Cornell planetary scientist Steve Squyres spent ten years having his proposals for Mars missions rejected by NASA before his instruments were finally chosen for the mission that eventually became the Mars Exploration Rovers. That decade-long ordeal, Squyres told me, was far worse than the stress of actually creating the rovers and getting them ready in time for their 2003 launch window.

But no one epitomized this essential mix of patience and determination more than NASA Administrator Tom Paine, who was at the Agency's helm during the first Apollo lunar landings. Even as Neil Armstrong and Buzz Aldrin took the first footsteps on the moon in July 1969, Paine hoped to convince

President Richard Nixon to support an onward-and-upward space program, culminating in human expeditions to the red planet. It wasn't just NASA that Paine wanted to advance; it was the human race. He believed space exploration had the power to transform civilization, not just with new technologies to improve life on Earth, but by giving humans the ability to settle other worlds. He wrote passionately about our future in space, of new branches of humanity taking root on alien shores in a renewal of the pioneer spirit of the American West. But Nixon and his people weren't interested, and the hope of human Mars missions was squelched indefinitely. In the summer of 1970, as the White House continued to slash the space budget, Paine left NASA and returned to private life.

Paine's Martian dreams never left him, and in 1984 he met up with a group of young people determined to put Mars on NASA's agenda, including several former grad students from the University of Colorado at Boulder. At a time when Mars wasn't even on NASA's radar screen, these young visionaries were trying to tackle the daunting obstacles to a Mars voyage, from keeping Mars-bound astronauts safe from solar flare radiation, to growing crops in a greenhouse on Mars, to obtaining rocket fuel from Martian resources. For the next decade Paine worked to help them, speaking at their conferences, writing papers, and adding his voice of experience to their efforts. Paine understood that the quest for Mars requires a commitment that spans generations, one in which you may not live to see your efforts bear fruit. He died in 1992; a chapter for the "Case for Mars" proceedings volume *Strategies for Mars* was the last thing he ever wrote.

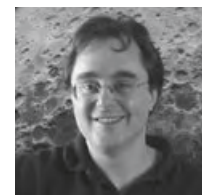
Passion: Don't leave home without it. It's what Squyres talked about when I interviewed him in the spring of 2004, a few months after Spirit and Opportunity reached Mars. "Every one of us who participates in a mission like this feels an incredible passion for what we do," Squyres told me. "That's what got us to the launchpad. It was that passion."

When I think of the passion for Mars, I think of the great science fiction writer Ray Bradbury, who has called Mars a

way station on humanity's journey toward immortality as an interstellar species. I think of Apollo 11 astronaut Mike Collins, who says the desire to journey to the red planet is part of "a primal urge to go outward bound A desire for what is next, what's farther away, what we have not been to. Outward bound. And I don't mean to get the photographs back; I mean to actually be there. To see, to touch, to smell, to die, to live there."

And I think of my college geology professor, Tim Mutch, who led the Viking lander imaging team that obtained the first pictures from the surface of Mars. Mutch, who was also a mountaineer, and who was killed during a Himalayan climb in 1980, once likened Viking to climbing Mount Everest: "There's a slim chance you'll make it," he told one interviewer, "but if you do, success is unequivocal!" Going to Mars ourselves, following in the footsteps of our robotic avatars, will be an Everest for the entire human species. And when we do, the reward will come not just from getting there, but who we become along the way: True to our nature as explorers and discoverers, harnessing our ingenuity to move all of humanity forward. No matter what we find on Mars, as Mutch once said, "The quest is the fundamental thing." ●

Science journalist **ANDREW CHAIKIN** is best known as the author of *A Man on the Moon*. His newest books, written with his wife, Victoria Kohl, are *Voices from the Moon* and *Mission Control, This Is Apollo*. www.andrewchaikin.com





Apollo

A Detective Story

BY GENE MEIERAN

Almost forty years ago, when I worked for Fairchild Semiconductor, I received an unusual telephone call from Andy Procassini, head of Fairchild Quality Assurance. Andy asked me, first, could I keep confidentiality about the topic he was calling about and, second, could I immediately come to the Mountain View facility and meet with him and a few other Fairchild folk? Of course, the only possible answer to such a request is, “Yes, sir. Be there in thirty minutes.”

I drove to the Mountain View plant and was ushered into a conference room with six or seven others. Andy immediately came to the point. “You know,” he said, “that Apollo 11 successfully landed on the moon on July 20, and all three astronauts are now safely home. Well, they [NASA] are planning on fueling up Apollo 12 for an early November launch, but there’s a problem. That’s why you people are here.”

While all systems tested “go” on the Apollo 12 Saturn rocket, Command Module (CM), and Lunar Module (LM), a problem had been detected in a later version of the LM radar transponder being put together at the Grumman facility in New Mexico. The failure was in a Fairchild linear amplifier identical to the one already installed and deeply buried in the electronics of the Apollo 12 LM. If that one failed, docking would be impossible. For obvious reasons, this was not an acceptable risk.

The questions put to the head of quality and reliability at Fairchild were, considering these devices had been tested umpteen times before installation and were operating properly, what went wrong with the failed amplifier? And, given that fueling of Apollo 12 was scheduled for the next ten days or so, what was the likelihood that the current properly operating device would fail during the mission?

Andy asked us for an answer within days and told us to be ready to go to Johnson Space Center in Houston to discuss the results of our investigation and analysis.

Building a Team

The five of us hardly knew each other; Mike was from marketing, Frank and Charlie from manufacturing, and I was from research and development. We had little in common, but here we were with a major problem that had to be resolved in days. We immediately got together and analyzed the data Fairchild had received. We had no physical evidence yet; the device that had failed was part of a disassembled LM in New Mexico. Other linear amps from the same batch were also being

retested, including those in the lunar modules for Apollo 13 and 14, but the devices were not available.

On the other hand, we had all the test data; the devices made available for NASA met the highest test standards available, Mil-Standard-883, and all devices in this batch had been tested and retested. So the first and most obvious question was how had our highest test standards passed devices that so quickly failed? Either the test procedures were at fault and we had passed bad devices, or the device failed because of something that occurred after the tests. Since the people who assembled the radar unit were not part of Fairchild and had obviously tested the device and unit subsequent to our selling the devices, we immediately suspected some sort of failure that occurred after device assembly into the radar module. We could rule out examination of our test procedures (even though we did look into these) and recognize that the device somehow failed after assembly.

We were indeed fortunate that our hastily assembled team got along; we did so because we all recognized we were becoming part of history, and we did not want history to record that Fairchild caused a delay of the second Apollo lunar landing. Furthermore, we all knew of each other, at least by reputation, and respected each other’s technical abilities, so there were no serious ego problems. Finally, we had a deadline to meet. There is nothing like a hard deadline to promote cooperation among dedicated technologists.

The Investigation

The next set of data to reach us was disheartening; other devices in the same batch, including another Apollo LM device, had also failed in NASA tests as they concurrently tried to trace the nature of the problem. The NASA problem escalated into a major field problem, as this device had been sold to a number of other customers, including the Department of Defense. If Fairchild had a batch of faulty devices incorporated into many sensitive applications, there could be enormous consequences beyond delaying a scheduled Apollo liftoff.

Charlie Gray and Frank Durand were responsible for manufacturing quality control, so they immediately got to work looking at the manufacturing records of these devices. In anticipation of exactly this kind of situation, Mil Standard devices had extensive traceability back to the sources of all parts used. My role was to analyze failed devices and come up with a plausible story of how and why they failed and, furthermore, make some sort of recommendation about the future of the specific device that was still functioning properly and installed on the Apollo 12 LM, already on the launchpad and being readied for fueling.

Our first act was to gather a number of other high reliability devices manufactured at the same time and retest them. Ordinarily, this should be unnecessary, since the high reliability testing was extensive and redundant, and 100 percent of the devices should pass rescreening. Imagine our surprise when a number of our stored devices failed this test, for test characteristics similar to those that failed the NASA tests. At about the same time, we received and retested some of the failed NASA devices. (Of course, they failed, too!)

The next step was obvious: open the hermetically sealed devices and see if we could identify the cause of failure. This part of the failure analysis was trivial; the cause was as obvious as it was astonishing. Basically, no bond wires connected the chip to the outside world. None! So solving why the devices failed was indeed trivial, but how had the wires disappeared?

The assembled devices were encased in a ceramic package sealed with a high temperature sealing glass in a special furnace, a process called hot cap sealing, prior to final testing. First, the completed chip was attached to a cavity in one part of the package, using conventional die attach processes. A metal lead frame was embedded in a thin layer of a high melting temperature sealing glass; this lead frame was the conduit of current and voltage to the external world from the embedded chip. The chip was connected to the lead frame through aluminum wires wire bonded to the aluminum coated lead frame, again using conventional semiconductor assembly processes.

The wire bonded bottom half of the ceramic package was then sealed to an upper cavity. In the hot cap sealing process this upper part of the package, which contained a layer of high temperature sealing glass, was heated to a temperature sufficient to melt the layer of glass, and the top part of the package was pressed onto the bottom part of the package, also heated to melt its glass layer. The two layers of molten glass would join and weld the parts of the package together. The chip was hermetically embedded in the sealed cavity, and the electrical signals would



Photo Credit: NASA

After the transporter carried the 363 foot high Apollo 12 Saturn V space vehicle to Launch Complex 39A and before fueling began, later productions of a small part embedded in the lunar module began to fail rigorous NASA testing.

pass through the glass seal by way of the embedded metal lead frames. It turned out that the temperature at which the ceramic parts were heated needed to be controlled to within a few degrees centigrade. This process had failed. The devices were sealed at too high a temperature; this excessive temperature was the most important cause of subsequent device failure.

Talking to NASA

Fortunately, the failure analysis took only a few days, so we had time to go to Houston to discuss the issue before a forced delay in fuel loading of Apollo 12 was to begin. Andy Procassini suggested we as a team go to Houston to tell them of our findings.

We arrived the day before our review with a host of NASA decision makers. We spent the night at our motel rehearsing our message. We discussed our strategy for the meeting and decided on the answer we knew we must be prepared to give and defend at the end of our presentation. Mike was chosen to talk about the devices, the architect to talk about its characteristics, and I would talk about the nature of the failure and its implications for Apollo 12.

Photo Credit: NASA



Lunar Module 6 for the Apollo 12 lunar landing mission is moved to an integration work stand in the Kennedy Space Center's Manned Spacecraft Operations Building.

The Meeting

We were ushered into the meeting at 8:30 a.m. We listened to NASA and Grumman engineers define the exact nature of their problem: a potential for a failed radar transponder after the LM left the surface of the moon, resulting in an impossibility of docking with the CM. The Grumman engineer gave an impressive talk about the device, stating what would happen if this or that particular pin failed for just about every possible combination of pin failures. This guy knew his radar system!

As he was talking, I looked around the room. Ten or twelve NASA officials, including Jim McDivitt and George Low, the ultimate decision maker, sat at a long table along with engineers responsible for the CM, the LM, the radar system, the fueling operation, and other elements. The hanging lights illuminating the table left the rest of the room in gloom; in this gloom were the attendees from Fairchild, from Grumman, from other Apollo spacecraft manufacturers, scientists, and engineers—perhaps another dozen people.

My presentation was quite simple. The hot glass sealer had exceeded its temperature for a brief time, heating the glass beyond

its normal sealing temperature. As a result, the glass seal was porous and allowed moisture to diffuse into this otherwise hermetic package. High levels of moisture combined with contaminants infused at the same time corroded the aluminum bond wires, leaving the appearance of no bond wires. Jim McDivitt asked if that meant aluminum would always dissolve in the presence of moisture, implying that, if so, the devices on Apollo 12 and 13 were time bombs ready to fail at any time. I said no, people always boiled water in aluminum containers. It took more than moisture or even contaminants; only specific contaminants attacked the thin aluminum oxide layer that protected all aluminum from instant corrosion. George Low suggested that since the contaminants present in the failed devices were likely present in the unfailed devices, they were still time bombs. In my view, the failed devices had failed months or years ago when the non-hermetic packages had been exposed to sufficient moisture and contaminants; currently operating devices were not likely to fail in the future, especially devices embedded in protective plastic as part of the lunar module assemblies.

George Low then asked the question I will always remember: “Dr. Meieran, would you fly this bird?” This was at 11:25 a.m., according to the clock that looked like Big Ben to me, on the wall in back of the long conference table. My response was, “Yes, I think it is safe to fuel Apollo 12, as the probability of this device failing is very, very small.” I knew that moisture diffused into even a badly sealed package and aluminum dissolved at a measurable rate quite fast compared to the time between assembly of the device and its encapsulation in the LM radar system. It seemed reasonable to believe that any corrosion that would occur had occurred already. This hypothesis was confirmed by examination of a large number of devices with different date codes.

For the next half hour, the NASA engineers discussed the implications of our findings. As the minute hand on the clock approached twelve, George Low announced, “It’s a go.” Looking back, I think my comment about being able to boil water in aluminum containers made the difference. While all these people were highly intelligent engineers, they were not corrosion scientists. Using a practical example they could relate to helped them understand my recommendation and trust it. ●



GENE MEIERAN is an Intel Senior Fellow.

The Knowledge Notebook

What We Owe the Past

BY LAURENCE PRUSAK



Not long ago a few of us who work on this magazine were talking about creating some sort of knowledge map of a NASA program—perhaps Kepler or even Apollo. We discussed trying to draw an easy-to-comprehend illustration of who contributed knowledge to the project and maybe even what and when they contributed. You get the picture. Well, you can't get the actual picture. After some reflection, we realized just how large a task this would be and how difficult even to figure out where to draw the line, because one could quite easily make the case for including Isaac Newton, or Leibniz, or Einstein.

Even if we limited our map to the actual time that the project was funded, we would face an uphill battle. So many do so much and *their* work depends on many others. Where do you start and stop?

Responding to someone who noted how much more we know than previous generations, the poet T.S. Eliot said, "Yes, but they are what we know." This is as true for science and engineering as it is for more humanistic endeavors. So very much of what we know in 2009—some estimates go to more than 90 percent—is handed to us on a plate. Economic historian Joel Mokyr calls it a real "free lunch."

As the present presses down on us with its constant demands, we all wish we knew more of this and that—that our lives and work would go so much more smoothly if only we knew more about chemistry or nuclear physics or some other subject. But stop and think for a minute about how much we do know that we didn't have to figure out or research or travel for or spend years in a lab to acquire. All that knowledge is a bequest to us from all those in the near and distant past who

worked on seemingly intractable topics in science and technology.

It is also salutary to think just how many unsung people actually contributed to inventions that we often attribute to lone geniuses. A new book by Gar Alperovitz and Lew Daly, *Unjust Deserts*, has some interesting things to say on this subject. For example, while almost everyone thinks of Alexander Graham Bell as the sole inventor of the telephone, there was another very viable claimant to that title. The work of an Italian immigrant named Antonio Meucci very likely preceded Bell's, but Bell patented his invention first because Meucci had trouble coming up with the \$10 patent fee. This "unknown" fact hasn't been entirely unknown.

Many long years ago when I was in college, I sat waiting in a small park for a girlfriend who lived in an Italian neighborhood in Brooklyn, New York. There was a small statue in the park dedicated to Meucci stating that he was the true inventor of the telephone. I knew it wasn't a joke, as the city had even named the park after him, but I was astounded that this very obscure man had beaten Bell to the phone and nobody but some New York Italians seemed to know it or teach it or even mention it. I couldn't even find anyone then, in the pre-Google age, who could tell me much about this man and his invention.

I mention Meucci here to make the same point Alperovitz and Daly do. So many people do so much science and engineering that all achievements and inventions depend on, as Newton famously put it, "standing on the shoulders of giants." It is only our very strong need in the United States to believe in individualism that makes some doubt

the truth of this. It isn't surprising that Bell and Meucci and yet another contender came up with the telephone at roughly the same time. That new technology was in the air because of all the other inventions and theories leading up to it that were known to these and no doubt other technicians. This is why Sputnik wasn't such a big surprise to those working in the field. Or why Newton and Leibniz invented calculus at the same time, and Darwin and Wallace both came up with theories of evolution. Only our need to reward individuals constrains our understanding of how deeply social all major inventions and intellectual developments really are.

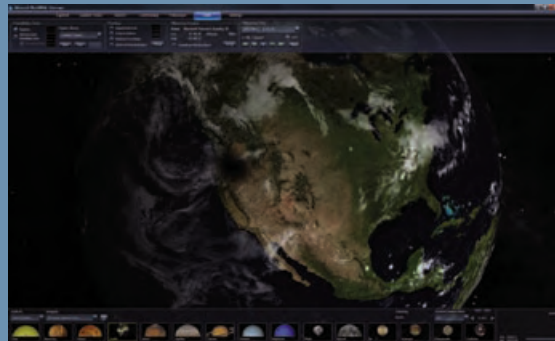
My goal here is not to take away any glory or diminish anyone's achievements. It is only, in a time of transition and baton-passing at NASA, to ask us to take some time to think about how many brilliant men and women have contributed to our own achievements and how we all stand on the shoulders of all those who came before us. ●

... STOP AND THINK FOR A MINUTE ABOUT HOW MUCH WE DO KNOW THAT WE DIDN'T HAVE TO FIGURE OUT OR RESEARCH OR TRAVEL FOR OR SPEND YEARS IN A LAB TO ACQUIRE. ALL THAT KNOWLEDGE IS A BEQUEST TO US FROM ALL THOSE IN THE NEAR AND DISTANT PAST ...

ASK interactive

NASA in the News

NASA announced plans to partner with Microsoft to develop technology that will make planetary images and data more readily available to the public. “Making NASA’s scientific and astronomical data more accessible to the public is a high priority for NASA, especially given the new administration’s recent emphasis on open government and transparency,” said Ed Weiler, associate administrator for NASA’s Science Mission Directorate in Washington. The project, WorldWide Telescope, is an online virtual telescope that allows users to explore NASA content, such as high resolution scientific images and data from Mars and the moon. Additional information and a free download of WorldWide Telescope can be found at www.worldwidetelescope.org.



Learning and Exploration

NASA has released *FROZEN*, its second major production for the Science on a Sphere platform. Developed by the National Oceanic and Atmospheric Administration, the Science on a Sphere technology gives viewers the sense that a globe is suspended in space before them as clouds and images of other climate features swirl over its surface. Goddard Space Flight Center’s Science Visualization Studio helped collect data from various satellites and turn them into images for the spherical film. To learn more about *FROZEN* and find cities where this unique cinema-in-the-round will be installed, visit www.nasa.gov/externalflash/frozen.

Web of Knowledge

Want to keep up with NASA’s activities? The Agency has jumped into social networking sites to give the public more insight into its activities, including live reporting during spacewalks and answers to questions in real time. Connect with NASA through social networks such as Twitter, Facebook, YouTube, and more at www.nasa.gov/collaborate/index.html.

For More on Our Stories

Additional information pertaining to articles featured in this issue can be found by visiting the following Web sites:

- **Constellation:** www.nasa.gov/mission_pages/constellation/main/index.html
- **Space Shuttle:** www.nasa.gov/mission_pages/shuttle/main
- **Mars Science Laboratory:** mars.jpl.nasa.gov/msl

feedback

We welcome your comments on what you’ve read in this issue of *ASK* and your suggestions for articles you would like to see in future issues. Share your thoughts with us at askmagazine.nasa.gov/about/write.php.

Not yet receiving your own copy of *ASK*?

To subscribe, send your full name and preferred mailing address (including mail stop, if applicable) to **ASKmagazine@asrcms.com**.

If you like *ASK Magazine*, check out *ASK the Academy*

ASK the Academy is an e-newsletter that offers timely news, updates, and features about best practices, lessons learned, and professional development. Learn more at **askacademy.nasa.gov/**.
