

①
Bulletin 54
(Part 1 of 3 Parts)

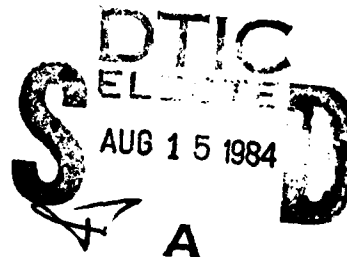
AD-A148 131

THE SHOCK AND VIBRATION BULLETIN

Part 1
Welcome, Keynote Address
Invited Papers MIL-STD-810D
MIL-STD-810D Panel Session

JUNE 1984

A Publication of
THE SHOCK AND VIBRATION
INFORMATION CENTER
Naval Research Laboratory, Washington, D.C.



Office of
The Under Secretary of Defense
for Research and Engineering

Approved for public release; distribution unlimited.

84 08 09 046

DTIC FILE COPY

SYMPOSIUM MANAGEMENT

THE SHOCK AND VIBRATION INFORMATION CENTER

J. Gordan Showalter, Acting Director

Rudolph H. Volin

Jessica Hileman

Elizabeth A. McLaughlin

Mary K. Gobbett

Bulletin Production

**Publications Branch, Technical Information Division,
Naval Research Laboratory**

Bulletin 54
(Part 1 of 3 Parts)

THE SHOCK AND VIBRATION BULLETIN

JUNE 1984

**A Publication of
THE SHOCK AND VIBRATION
INFORMATION CENTER
Naval Research Laboratory, Washington, D.C.**

The 54th Symposium on Shock and Vibration was held in Pasadena, California, October 18-20, 1983. The Jet Propulsion Laboratory in Pasadena was the host.

**Office of
The Under Secretary of Defense
for Research and Engineering**



A/

CONTENTS

PAPERS APPEARING IN PART I

Welcome

WELCOME	1
Robert J. Parks, Associate Director, Space Science and Exploration, Jet Propulsion Laboratory, Pasadena, CA	

Keynote Address

KEYNOTE ADDRESS	3
Robert S. Ryan, George C. Marshall Space Flight Center, Huntsville, AL	

Invited Papers

DNA ICBM TECHNICAL R&D PROGRAM	23
Colonel Maxim I. Kovel, Director, Shock Physics Directorate, Defense Nuclear Agency, Washington, DC	
SOME DYNAMICAL ASPECTS OF ARMY MISSILE SYSTEMS	43
Dr. James J. Richardson, Chief, Structures and Mechanics, U.S. Army Missile Command, Redstone Arsenal, AL	
AIR FORCE SPACE TECHNOLOGY CENTER SPACE TECHNOLOGY - EMPHASIS 84	55
Colonel Frank J. Redd, Vice Commander, Air Force Space Technology Center, Kirtland AFB, NM	
REFLECTIONS ON TRENDS IN DYNAMICS - THE NAVY'S PERSPECTIVE	59
Henry C. Pusey, Consultant, NKF Engineering Associates, Inc., Vienna, VA	
ELIAS KLEIN MEMORIAL LECTURE - MODAL TESTING - A CRITICAL REVIEW	65
Strether Smith, Lockheed Palo Alto Research Laboratory, Palo Alto, CA	
SOLUTIONS TO STRUCTURAL DYNAMICS PROBLEMS	77
Dr. George Morosow, Martin Marietta Corporation, Denver, CO	
WHERE IS THE REAL LITERATURE ON AIRBLAST AND GROUND SHOCK?	83
Dr. Wilfred E. Baker, Southwest Research Institute, San Antonio, TX	

MIL-STD-810D

TAILORING INITIATIVES FOR MIL-STD-810D ENVIRONMENTAL TEST METHODS AND ENGINEERING GUIDELINES	87
David L. Earls, Air Force Wright Aeronautical Laboratories, Wright-Patterson AFB, OH	
ACCELERATION RESPONSES OF TYPICAL LRU'S SUBJECTED TO BENCH HANDLING AND INSTALLATION SHOCK	91
H. Caruso and E. Szymkowiak, Westinghouse Electric Corporation, Baltimore, MD	
IMPACT OF 810D ON DYNAMIC TEST LABORATORIES	101
Dr. Allen J. Curtis, Hughes Aircraft Company, El Segundo, CA	
THE CHANGING VIBRATION SIMULATION FOR MILITARY GROUND VEHICLES	113
Jack Robinson, Materials Testing Directorate, Aberdeen Proving Ground, MD	
PANEL DISCUSSION - MIL-STD-810D	125

PAPERS APPEARING IN PART 2

Ship Shock

TWO-DIMENSIONAL SHOCK RESPONSE OF A MASS ON ENERGY-ABSORBING SHOCK MOUNTS

R. E. Fortuna and V. H. Neubert, The Pennsylvania State University, University Park, PA

OPTIMUM DESIGN FOR NONLINEAR SHOCK MOUNTS FOR TRANSIENT INPUTS

K. Kasraie, Firestone Tire & Rubber Company, Central Research Laboratories, Akron, OH and V. H. Neubert, The Pennsylvania State University, University Park, PA

THE DEVELOPMENT OF A METHOD FOR THE SHOCK-RESISTANT SECURING OF LARGE BATTERIES IN SUBMARINES

A. Jansen, Royal Netherlands Navy, The Hague

SHIPBOARD SHOCK RESPONSE OF THE MODEL STRUCTURE DSM; EXPERIMENTAL RESULTS VERSUS RESPONSES PREDICTED BY EIGHT PARTICIPANTS

R. Regoord, TNO-IWECO, Delft, the Netherlands

DIRECT ENERGY MINIMIZATION APPROACH TO WHIPPING ANALYSIS OF PRESSURE HULLS

K. A. Bannister, Naval Surface Weapons Center, White Oak, Silver Spring, MD

Shock

WATER IMPACT LABORATORY AND FLIGHT TEST RESULTS FOR THE SPACE SHUTTLE SOLID ROCKET BOOSTER AFT SKIRT

D. A. Kross, NASA/Marshall Space Flight Center, Marshall Space Flight Center, AL, N. C. Murphy, United Space Boosters, Inc., Huntsville, AL, and E. A. Rawls, Chrysler Corporation, New Orleans, LA

AN OBJECTIVE ERROR MEASURE FOR THE COMPARISON OF CALCULATED AND MEASURED TRANSIENT RESPONSE HISTORIES

T. L. Geers, Lockheed Palo Alto Research Laboratory, Palo Alto, CA

ALTERNATIVE SHOCK CHARACTERIZATIONS FOR CONSISTENT SHOCK TEST SPECIFICATION

T. J. Baca, Sandia National Laboratories, Albuquerque, NM

SHOCK RESPONSE ANALYSIS BY PERSONAL COMPUTER USING THE EXTENDED IFT ALGORITHM

C. T. Morrow, Consultant, Encinitas, CA

LEAST FAVORABLE RESPONSE OF INELASTIC STRUCTURES

F. C. Chang, T. L. Paez, and F. Ju, The University of New Mexico, Albuquerque, NM

LOW VELOCITY, EXPLOSIVELY DRIVEN FLYER PLATE DESIGN FOR IMPACT FUZE DEVELOPMENT TESTING

R. A. Benham, Sandia National Laboratories, Albuquerque, NM

EXPERIMENTAL INVESTIGATION OF VIBROIMPACT OF TWO OSCILLATORS

C. N. Bapat and S. Sankar, Concordia University, Montreal, Quebec, Canada

MODELS FOR SHOCK DAMAGE TO MARINE STRUCTURAL MATERIALS

D. W. Nicholson, Naval Surface Weapons Center, White Oak, Silver Spring, MD

A STUDY OF THE EFFECT OF MASS LOADING ON THE SHOCK ENVIRONMENT

Q. Z. Wang, Beijing Institute of Strength and Environment Engineering, Beijing, China and H. B. Lin, Chinese Academy of Space Technology, Beijing, China

Blast and Ground Shock

ASSESSMENT OF SEISMIC SURVIVABILITY

R. E. McClellan, The Aerospace Corporation, El Segundo, CA

GROUND SHOCK EFFECT ON SOIL FIELD INCLUSIONS

R. E. McClellan, The Aerospace Corporation, El Segundo, CA

PENETRATION OF SHORT DURATION AIRBLAST INTO PROTECTIVE STRUCTURES

J. R. Britt and J. L. Drake, Applied Research Associates, Southern Division, Vicksburg, MS

A COMPUTATIONAL PROCEDURE FOR PEAK INSTRUCTURE MOTIONS AND SHOCK SPECTRA FOR CONVENTIONAL WEAPONS

S. A. Kiger, J. P. Balsara, and J. T. Baylot, USAE Waterways Experiment Station, Vicksburg, MS

PRELIMINARY DESIGN CRITERIA AND CERTIFICATION TEST SPECIFICATIONS FOR BLAST RESISTANT WINDOWS

G. E. Meyers, W. A. Keenan, and N. F. Shoemaker, Naval Civil Engineering Laboratory, Port Hueneeme, CA

PAPERS APPEARING IN PART 3

Structural Dynamics

STRUCTURAL MODIFICATIONS BY VISCOELASTIC ELEMENTS

P. J. Riehle, Anatrol Corporation, Cincinnati, OH

STOCHASTIC DYNAMIC ANALYSIS OF A STRUCTURE WITH FRICTIONAL JOINTS

Q. L. Tian and Y. B. Liu, Institute of Mechanics, Chinese Academy of Sciences and D. K. Liu, Space Science & Technology Centre, Chinese Academy of Sciences, Beijing, China

MODAL ANALYSIS OF STRUCTURAL SYSTEMS INVOLVING NONLINEAR COUPLING

R. A. Ibrahim, T. D. Woodall, and H. Heo, Department of Mechanical Engineering, Texas Tech University, Lubbock, TX

DISCRETE MODIFICATIONS TO CONTINUOUS DYNAMIC STRUCTURAL SYSTEMS

Y. Okada, Ibaraki University, Hitachi, Japan, B. P. Wang, University of Texas at Arlington, Arlington, TX, and W. D. Pilkey, University of Virginia, Charlottesville, VA

REANALYSIS OF CONTINUOUS DYNAMIC SYSTEMS WITH CONTINUOUS MODIFICATIONS

B. P. Wang, University of Texas at Arlington, Arlington, TX, Y. Okada, Ibaraki University, Hitachi, Japan, and W. D. Pilkey, University of Virginia, Charlottesville, VA

A POLE-FREE REDUCED-ORDER CHARACTERISTIC DETERMINANT METHOD FOR LINEAR VIBRATION ANALYSIS BASED ON SUB-STRUCTURING

B. Dawson, Polytechnic of Central London, London, England, and M. Davies, University of Surrey, Guildford, Surrey, England

DETERMINATION OF SHEAR COEFFICIENT OF A GENERAL BEAM CROSS SECTION BY FINITE ELEMENT METHOD

C. M. Friedrich and S. C. Lin, Westinghouse Electric Corporation, Bettis Atomic Power Laboratory, West Mifflin, PA

Machinery Dynamics

GEAR CASE VIBRATION ISOLATION IN A GEARED TURBINE GENERATOR

R. P. Andrews, Westinghouse Electric Corporation, Marine Division, Sunnyvale, CA

EFFECT OF COUPLED TORSIONAL-FLEXURAL VIBRATION OF A GEARED SHAFT SYSTEM ON THE DYNAMIC TOOTH LOAD

S. V. Neriya, R. B. Bhat, and T. S. Sankar, Concordia University, Montreal, Quebec, Canada

PERCISION MEASUREMENT OF TORSIONAL OSCILLATIONS INDUCED BY GEAR ERRORS

S. L. Shmutter, Ford Motor Company, Manufacturing Processes Laboratory, Dearborn, MI

THE ANALYSIS BY THE LUMPED PARAMETER METHOD OF BLADE PLATFORM FRICTION DAMPERS USED IN THE HIGH PRESSURE FUEL TURBOPUMP OF THE SPACE SHUTTLE MAIN ENGINE

R. J. Dominic, University of Dayton Research Institute, Dayton, OH

Vibration Problems

TRANSIENT VIBRATION TEST CRITERIA FOR SPACECRAFT HARDWARE

D. L. Kern and C. D. Hayes, Jet Propulsion Laboratory, California Institute of Technology, Pasadena, CA

**VIBRATIONAL LOADING MECHANISM OF UNITIZED CORRUGATED CONTAINERS WITH CUSHIONS
AND NON-LOAD-BEARING CONTENTS**

T. J. Urbanik, Forest Products Laboratory, USDA Forest Service, Madison, WI

LEAKAGE-FLOW INDUCED VIBRATIONS OF A CHIMNEY STRUCTURE SUSPENDED IN A LIQUID FLOW

H. Chung, Components Technology Division, Argonne National Laboratory, Argonne, IL

**THE EXPERIMENTAL PERFORMANCE OF AN OFF-ROAD VEHICLE UTILIZING A
SEMI-ACTIVE SUSPENSION**

E. J. Krasnicki, Lord Corporation, Erie, PA

EFFECT OF AIR CAVITY ON THE VIBRATION ANALYSIS OF LOADED DRUMS

S. De, National Research Institute, W. Bengal, India

SESSION CHAIRMEN AND COCHAIRMEN

<u>Date</u>	<u>Session Title</u>	<u>Chairmen</u>	<u>Cochairmen</u>
Tuesday, 18 Oct. A.M.	Opening Session	Dr. Ben Wada, The Jet Propulsion Laboratory, Pasadena, CA	Dr. J. Gordan Showalter, The Shock & Vibration Information Center, Naval Research Laboratory, Washington, DC
Tuesday, 18 Oct. P.M.	Elias Klein Memorial Lecture Plenary A	Dr. J. Gordan Showalter, Shock and Vibration Information Center, Naval Research Laboratory Washington, DC	
Tuesday, 18 Oct. P.M.	Ship Shock	Mr. Gene Remmers, David Taylor Naval Ship Research and Development Center, Bethesda, MD	Dr. Michael Paksysa, NKF Engineering Associates, Vienna, VA
Tuesday, 18 Oct. P.M.	Space Vibration	Mr. Jerome Pearson, Air Force Wright Aeronautical Laboratories, Wright-Patterson AFB, OH	Mr. John Garba, Jet Propulsion Laboratory, Pasadena, CA
Wednesday, 19 Oct. A.M.	Plenary B	Mr. William J. Walker, Boeing Aerospace Company, Seattle, WA	Dr. George Morosow, Martin Marietta Corporation, Denver, CO
Wednesday, 19 Oct. A.M.	Structural Dynamics	Mr. Edward Fleming, The Aerospace Corporation, Los Angeles, CA	Dr. John Gubser, McDonnell Douglas Astronautics Company, St. Louis, MO
Wednesday, 19 Oct. A.M.	MIL-STD-810D Session I, Rationale	Mr. John Wafford, Aeronautical Systems Division, Wright Patterson AFB, OH	Mr. Robert Hancock, Vought Corporation, Dallas, TX
Wednesday, 19 Oct. P.M.	Shock	Mr. Ami Frydman, Harry Diamond Laboratories, Adelphi, MD	Mr. Martin Walchak, Harry Diamond Laboratories, Adelphi, MD
Wednesday, 19 Oct. P.M.	MIL-STD-810D Session II, Implementation and Use	Mr. Rudolph H. Volin, Shock and Vibration Information Center, Washington, DC	Mr. W. W. Parmenter, Naval Weapons Center, China Lake, CA
Thursday, 20 Oct. A.M.	Blast/Ground Shock	Mr. William Flatheu, U.S. Army Engineer Waterways Experiment Station, Vicksburg, MS	Mr. George Coulter, U.S. Army Ballistic Research Laboratory, Aberdeen Proving Ground, MD
Thursday, 20 Oct. A.M.	Machinery Dynamics	Dr. David Fleming, NASA Lewis Research Center, Cleveland, OH	Dr. Hanson Huang, Naval Surface Weapons Center, Silver Spring, MD
Thursday, 20 Oct. P.M.	Vibration Problems	Dr. Robert S. Reed, Jr., Naval Surface Weapons Center, Silver Spring, MD	Dr. Larry Pinson, NASA Langley Research Center, Hampton, VA
Thursday, 20 Oct. P.M.	Short Discussion Topics	Mr. Howard Camp, Jr., U.S. Army Electron Research and Development Command, Ft. Monmouth, NJ	Mr. E. Kenneth Stewart, U.S. Army Armament Research and Development Command, Picatinny Arsenal Dover, NJ

WELCOME

Mr. Robert J. Parks
Associate Director for Space Science and Exploration
Jet Propulsion Laboratory
Pasadena, CA

I'd like to welcome all of you to Pasadena and to the Jet Propulsion Laboratory. I do this on behalf of Dr. Allen who is the Director of the Jet Propulsion Laboratory (JPL) and who could not be here this morning. In fact he is in the Soviet Union, and it would have been a little bit difficult to commute to the meeting this morning. On his behalf and on behalf of all the rest of us at JPL, we certainly do welcome you here to what I'm sure will be a very useful and productive session.

We at JPL certainly are able to fully appreciate the importance of the activities that you are undertaking, and we endorse these efforts completely. We want to do everything we can to help out and support these activities. I am sure that arrangements have been well made, and I don't anticipate any problems, but arrangements can be made to help with whatever turns out to be needed. Probably the biggest contribution is that Ben Wada has been able to play a role in putting all of this program together. We are very pleased about that.

Ben mentioned how long I've been at JPL. Over all those years, the kind of activity you are discussing here this morning has been a key part of our space science and exploration activities at JPL. The latest example of this is the vibration, shock, and environmental testing of a structural model of the GALILEO spacecraft.

The GALILEO spacecraft, as you may be aware, is scheduled to be launched in 1986, and will carry a combined Orbiter and Probe to Jupiter. It will send the Probe into the upper atmosphere of Jupiter down to about 10-20 bars, and it will make the first direct measurement of that atmosphere. Then the Orbiter will stay around for another 20 months or so and observe the planet, its many satellites and its unusual environment.

The design of the GALILEO spacecraft has turned out to be quite challenging. In many respects it is the most complex, or capable, dual spin spacecraft that has been put together so far. So we found quite a few engineering challenges in putting it together and testing it to make sure it is all right. But it is in that phase right now, and as far as I'm aware, it's been going very well.

Although I understand most of the unclassified sessions will be held here, there is a series of classified sessions which will be held at JPL. I also understand that there is a planned visit to JPL on Friday for any of you who are able to make it. We certainly welcome all of you and encourage all of you to visit if it fits in with your plans.

Again, I would just like to say "Welcome" to all of you and give you our best wishes and good luck in your further endeavors here. Thank you.

KEYNOTE ADDRESS

Robert S. Ryan
George C. Marshall Space Flight Center
Huntsville, Alabama

Good morning, ladies and gentlemen. It is an honor and privilege for me to be here. I am pleased to have the opportunity to address the 54th Shock and Vibration Symposium. I bring you greetings from NASA and the Marshall Space Flight Center. We at NASA have supported this group for many years, being involved in many aspects of what you do. What an impressive record you have! It is indeed my pleasure to be a part of this meeting. With all the talents assembled, I seriously question the merits of what I have to say, and yet for some reason, I do have some things I really want to say.

There are many definitions of what a keynote address should be and what it should accomplish, so I went directly to the expert who invited me here, Ben Wada, for the answer. Just as I expected, he gave me an impossible task delineating four objectives. Let me use Figure 1 to illustrate my guidelines. First, I must wake you up. Secondly, I should shake you up. Thirdly, I am to entertain you. Fourthly and finally, I am to soar you to new heights. All to be accomplished in 30 to 45 minutes. Really, the task is achievable, but not by me. The informal discussions associated with being together in conjunction with the formal papers are the way these objectives are met. I am a firm believer in the value of this yearly convention and what it accomplishes. The theme you have chosen is excellent, "Old Problems - New Solutions."

The ancient Greeks had a legend that every five hundred years, the Phoenix, a mythical bird, burst into flames and was reduced to ashes. From these ashes, the Phoenix bird rose again, renewed in youthful vigor. Although this is only a legend, in actual life we have found it necessary to begin again with little more than the ashes of the past, rising to new heights with great vigor. I do not believe we are in the position of having only ashes left from the past; however, the principles of a new start anchored in the past is very sound. Meetings like this serve this purpose well. We are away from the job, home, etc., which puts us in a good psychological state for new visions of solutions to old problems. By training, we are trapped in the very effective method of "linear

thinking," a logical way to solve problems. What we need is a good case of lateral thinking, a jolt to the side, that provides a new starting point. From this new starting point, the old standby, "linear thinking," again serves us in good stead. What is so hard for us to accomplish is this lateral or side leap which provides the ideas for a new solution. I predict you will get many high voltage, lateral jolts while you are here. Specifically, I want to talk to you today from the vantage point of NASA, where we have been, where we are, and where we are going. First, from NASA's viewpoint and then from the disciplines associated with the Shock and Vibration Information Group. This approach seems compatible with your overall theme, "Old Problems - New Solutions." We in space exploration need new solutions to old problems, as well as new visions for future problems and their solutions. I will not address the aeronautics side of NASA, since it is not a part of my experience base.

I. NASA - PAST, PRESENT and FUTURE

NASA is presently celebrating its 25th anniversary year, evolving from the NACA organization which had its beginning in 1915. We at NASA have our roots firmly anchored in aeronautics, which is still part of our charter. To fly was the first step toward space and in a real sense could be the final step. The agency was signed into law as a civilian space agency by President Eisenhower on July 29, 1958. Most of us remember with vividness the shock of Sputnik that resulted in NASA's birth. My first sighting of Sputnik came near sunset with such an impressive glow that it is still etched clearly in my mind. We have accomplished much as a government/industry team in those 25 years, resulting in an impressive technology resource for our nation. The thrill of space exploration has not diminished. Each new achievement brings renewed emotional peaks, not only for us involved but for the public in general. MSFC received more letters, etc., after STS-8 flew than from any previous flight. Time will permit looking at only a few snapshots of these accomplishments. Figure 2 summarizes to some extent where we have been.

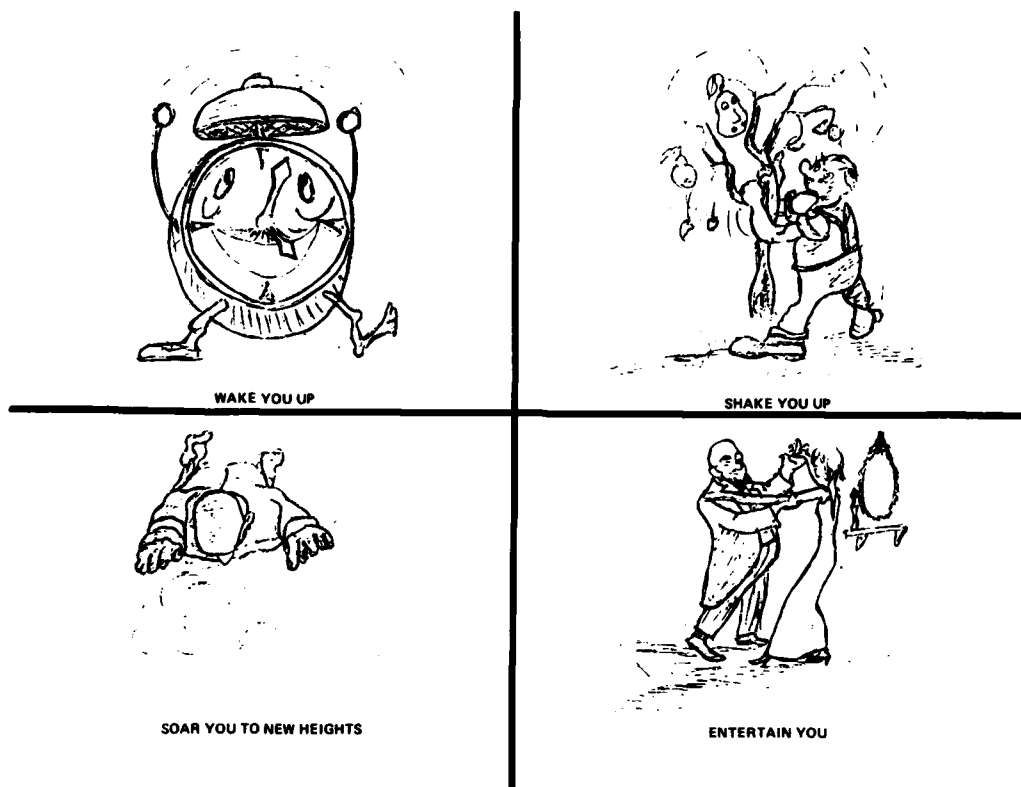


Fig. 1 - Objectives of a Keynote Address

WHERE HAVE WE BEEN ?

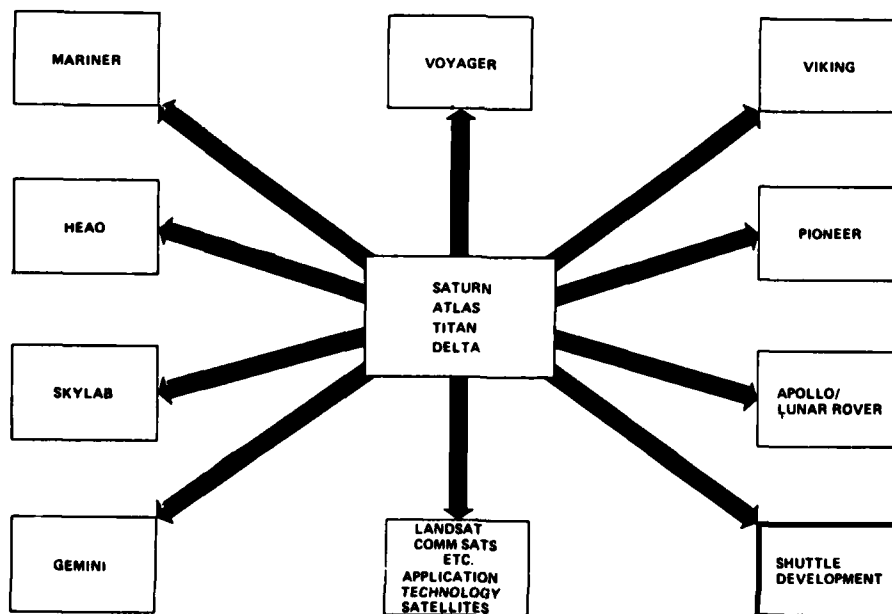


Fig. 2 - Previous NASA Launch Vehicles and Payloads

In the center is listed most of the launch vehicles used with some of the key programs or payloads on the perimeter. In the right-hand corner is the Shuttle development looming at us as an extension of the earlier vehicles, obviously more recently but still in our past as far as development goes. Let us not dwell on these accomplishments, nor how we overcame many difficulties, setbacks, and problems. Briefly, we should recall the grandeur, excitement, national prestige, and even more importantly, the fresh look at our planet Earth that accompanied them. Just to refresh your memory, Figure 3 was made when we were just into early Shuttle development and in the middle of Skylab. Emphasized are Marshall Space Flight Center's management roles. Snapshots of the Lunar Rover, Saturn, HEAO, Skylab and Shuttle are shown.

Apollo, with the lunar landing and lunar exploration, was indeed one small step for man - one giant leap for shock and vibration. Apollo was composed of 11 manned flights involving 29 astronauts, 12 of whom placed their footprints on the Moon. In addition, there were two manned Earth orbit preliminaries, three circumlunar flights, and six lunar landing missions. The results of Mariner (1978), Pioneer (1978), Viking (1976), Voyager (1977), again raised us to new heights, providing new insights into our universe and our origins. The people here at JPL can better tell these stories, although emotionally and in some special cases technically, we shared together. Skylab (1973) was our first orbiting space station (Figure 4). Out of near disaster came the highly successful exploration of Earth and beyond, telescope, materials processing, earth resources, containing three missions of 28, 59, and 84 days, proving the resiliency and necessity of man in space.

Space Shuttle Columbia lifted off the pad in 1981 (Figure 5) followed by 7 more flights, 5 of which were dubbed developmental, while the last 3 were operational. With this step, man has the capability to routinely and efficiently enter space.

Squeezed in between, from Marshall's viewpoint, were the three HEAO missions (Figure 6) launched using Air Force vehicles adding greatly to our scientific knowledge. Briefly, this is where we have been -- to the Moon, the planets, and beyond, using efficiently both manned and unmanned space exploration.

The next question peeking over the horizon is, "Where are we today with the Shuttle development behind us?" A record of eight successful Space Shuttle launches and number 9 (Spacelab 1) ready for launch later this month is in the books. We jointly are in full swing with the next phase, "utilization of space," with the Space Shuttle the work horse. Figure 7 illustrates some of the activities we are into now or will be very shortly.

The first Spacelab mission flies this month with two other Spacelab missions to follow shortly. Spacelab is one of our laboratories for utilization of space. Many options are available for various experiments and space exploration. Work is in progress on the Western Test Range facilities with the first Space Shuttle launch scheduled there within two years. Space Telescope is progressing towards a 1986 launch, providing scientists with their greatest opportunity yet to explore our universe. The technical challenges associated with designing, verifying, and operating a telescope of this nature is mind boggling. Pointing accuracy, length of operations time, etc., are unprecedented. The Long Duration Exposure Facility is moving steadily towards launch. A solar wing (SAFE), forerunner of space power, will be launched within 18 months, including for the first time "on-orbit" dynamic testing using remote sensors. Tracking Data Relay Satellites are in orbit and are going in orbit to serve space as well as mankind. Special missions are moving ahead rapidly. Also, many get-away specials can be flown on Shuttle on space availability status. These are experiments that can be quickly installed or substituted as space becomes available. The facilities and procedures are developed and working at KSC for Shuttle payload processing and integration with the Orbiter. Routine space operations using the Shuttle are here.

With that brief snapshot of where we are, let's look at where we are going. The President last year delineated a space policy, and NASA has formulated a set of goals and objectives to carry out this space policy. Figure 8 lists these goals. Specific objectives have been developed for each of these goals.

Clearly, the keys are man's presence in space, low-cost Shuttle operations, space science, space technology, and aeronautics. Coming out of the pack as a focus for some of these goals is the Space Station (Figure 9). You will be hearing much about this in the future. Obviously, a major focus is making the Shuttle operations routine and cost effective with all this implied.

If one looks at the goals and what is in the works, many exciting possibilities are in various stages of ideas, plans, or development. Figure 10 lists some of the various areas of the goals we are committed to. We are working the upper stages as a means of higher orbits and planetary missions. Utilization of the Shuttle for experimentation in all forms from get-away specials to Space Telescope are key areas. Many orbiting observatories from IRAS to AXAF are moving forward and offer exciting potentials. Galileo will be with us shortly as will the Tethered Satellite. Spacelab will be a major activity for many years. Shuttle performance improvements are a continuing goal. Planning boards already have many concepts for large lift

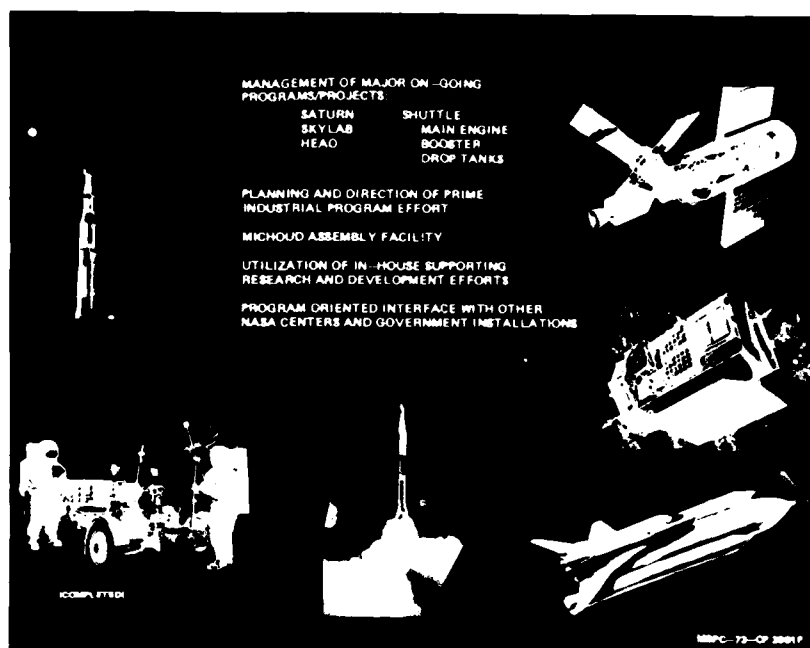


Fig. 3 — Marshall Space Flight Center Program Management Missions

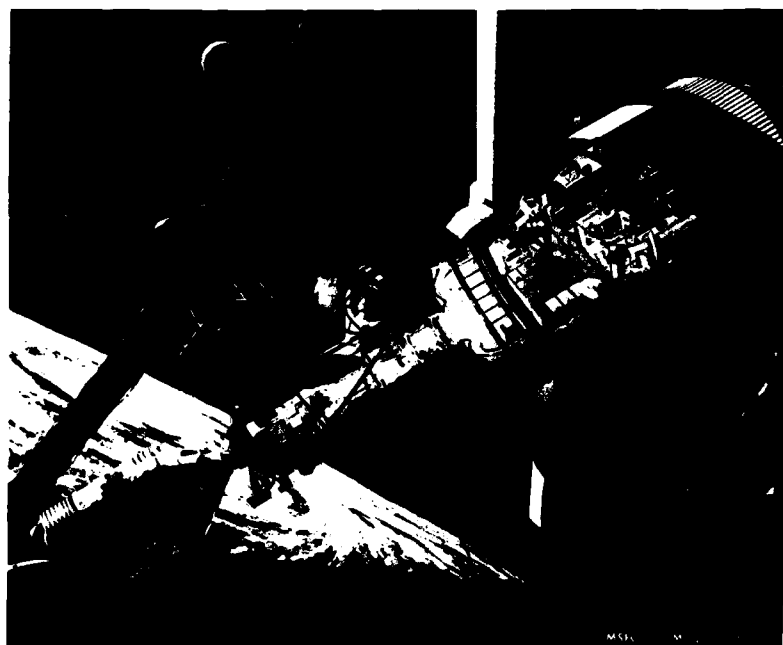


Fig. 4 — Skylab

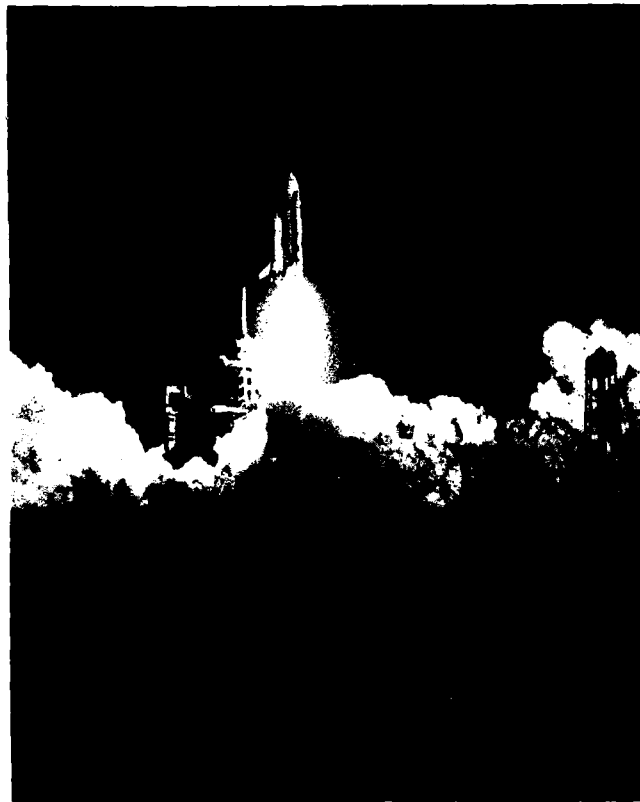


Fig. 5 — Space Shuttle Columbia Lift Off in 1981

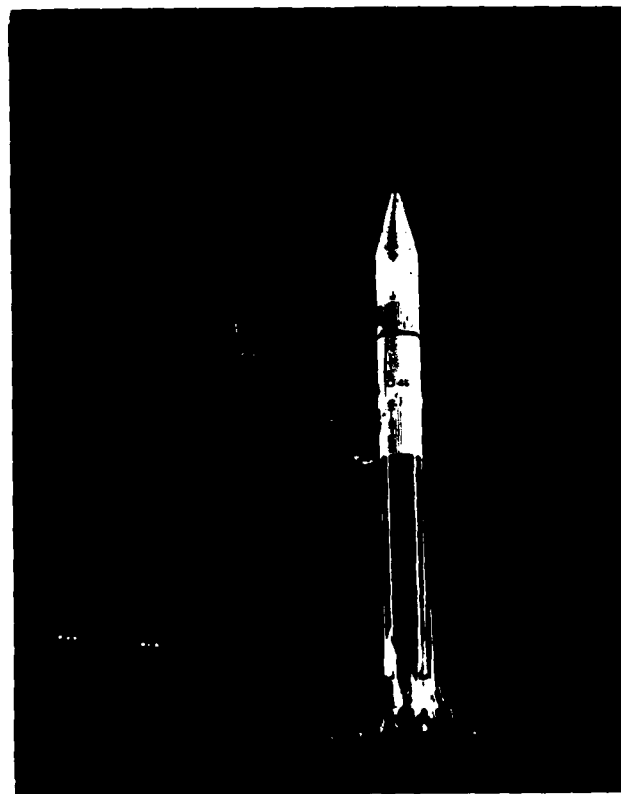


Fig. 6 — Atlas-Centaur with HEAO-1 Spacecraft

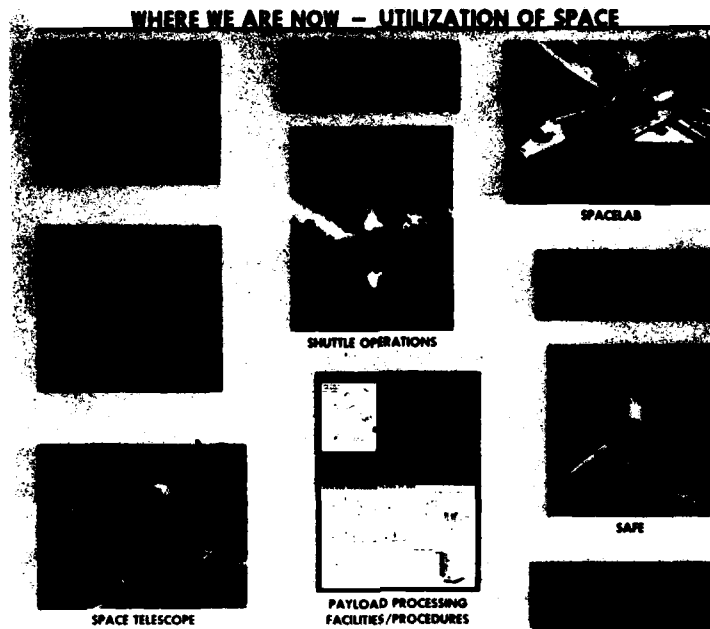


Fig. 7 — Utilization of Space

- PROVIDE FOR OUR PEOPLE A CREATIVE ENVIRONMENT AND THE BEST OF FACILITIES, SUPPORT SERVICES, AND MANAGEMENT SUPPORT SO THEY CAN PERFORM WITH EXCELLENCE NASA'S RESEARCH, DEVELOPMENT, MISSION, AND OPERATIONAL RESPONSIBILITIES.
- MAKE THE SPACE TRANSPORTATION SYSTEM FULLY OPERATIONAL AND COST EFFECTIVE IN PROVIDING ROUTINE ACCESS TO SPACE FOR DOMESTIC AND FOREIGN, COMMERCIAL AND GOVERNMENTAL USERS.
- ESTABLISH A PERMANENT MANNED PRESENCE IN SPACE TO EXPAND THE EXPLORATION AND USE OF SPACE FOR ACTIVITIES WHICH ENHANCE THE SECURITY AND WELFARE OF MANKIND.
- CONDUCT AN EFFECTIVE AND PRODUCTIVE AERONAUTICS PROGRAM WHICH CONTRIBUTES MATERIALLY TO THE ENDURING PREEMINENCE OF U. S. CIVIL AND MILITARY AVIATION.
- CONDUCT AN EFFECTIVE AND PRODUCTIVE SPACE SCIENCE PROGRAM WHICH EXPANDS HUMAN KNOWLEDGE OF THE EARTH, ITS ENVIRONMENT, THE SOLAR SYSTEM, AND THE UNIVERSE.
- CONDUCT EFFECTIVE AND PRODUCTIVE SPACE APPLICATIONS AND TECHNOLOGY PROGRAMS WHICH CONTRIBUTE MATERIALLY TOWARD U. S. LEADERSHIP AND SECURITY.
- EXPAND OPPORTUNITIES FOR U. S. PRIVATE SECTOR INVESTMENT AND INVOLVEMENT IN CIVIL SPACE AND SPACE-RELATED ACTIVITIES.
- ESTABLISH NASA AS A LEADER IN THE DEVELOPMENT AND APPLICATION OF ADVANCED TECHNOLOGY AND MANAGEMENT PRACTICES WHICH CONTRIBUTE TO SIGNIFICANT INCREASES IN BOTH AGENCY AND NATIONAL PRODUCTIVITY.

Fig. 8 — NASA Goals and Objectives

WHERE WE ARE GOING

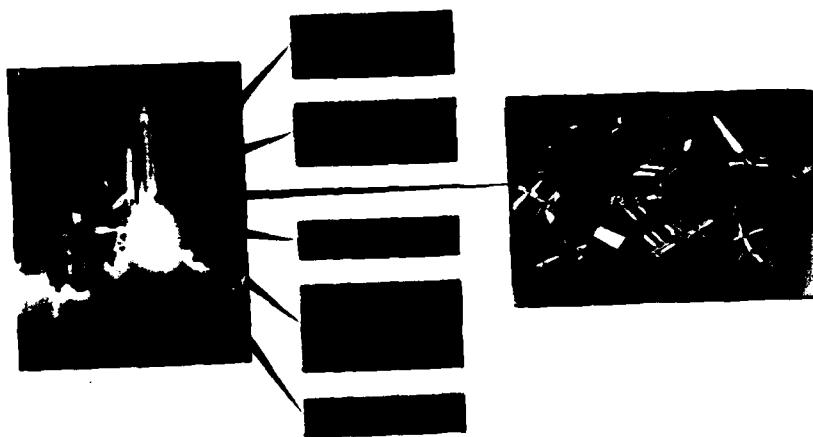


Fig. 9 — Future Space Programs — Space Station

- UPPER STAGES FOR HIGH ORBITS AND PLANETARY EXPLORATION
 - IUS
 - CENTAUR
 - OTV
 - PAM
- GETAWAY SPECIALS
- STUDENT EXPERIMENTS
- SPACELAB
- SPACE POWER SYSTEM
- TETHERED SATELLITE
- SPACE STATION
- PLANETARY
 - GALILEO
- COMET STUDIES
 - INTERNATIONAL HALLEY WATCH (IHW)
 - ISEE -3
- SPACE TELESCOPE
- ORBITING OBSERVATIONS
 - INFRARED ASTRONOMICAL SATELLITE (IRAS)
 - SHUTTLE INFRARED TELESCOPE FACILITY (SIRTF)
 - COSMIC BACKGROUND EXPLORER (COBE)
 - GAMMA RAY OBSERVATORY (GRO)
 - ADVANCED X-RAY ASTROPHYSICS FACILITY (AXAF)
- GLOBAL ENVIRONMENT
 - ACTIVE MAGNETOSPHERIC TRACER EXPLORERS
 - EARTH RADIATION BUDGET EXPLORER (ERBE)
 - ADVANCED UPPER ATMOSPHERE RESEARCH SATELLITE
- LARGE LIFT VEHICLES
- SHUTTLE PERFORMANCE IMPROVEMENTS
 - FILAMENT WOUND CASE SOLID ROCKET MOTORS (FWC SRM)
 - WEIGHT REDUCTIONS
 - COMPOSITES
 - HIGHER PERFORMANCE PROPULSION SYSTEM

Fig. 10 — NASA's Future Activities and Programs

launch vehicles and Shuttle derivations. This thumbnail sketch completes our survey of NASA.

II. SHOCK AND VIBRATION - PAST, PRESENT AND FUTURE

I would like to turn to the technical areas you are most concerned with, viewing them from my vantage point as a structural dynamicist, who has spent many years working the control disciplines. Clearly, we in the technical disciplines associated with shock and vibration face our greatest challenges. Visions of the future with a firm understanding of where we are now and where we are to focus are mandatory if we are to meet the goals/objectives of our great organizations and our commitment to excellence and the future of mankind. Figure 11 attempts to answer in visual form where we have been, where we are now, and where we are going. The chart has three messages: first, it shows how some of the major technical disciplines are becoming more and more complex, pushing the state-of-the-art or beyond; second, and possibly even more important, it shows that we must solve these technical challenges with less money and schedule time, while maintaining the same basic level of reliability; and last, it says if we are to accomplish these two major challenges, more complex systems at lower cost and schedule time, with the same reliability, then we cannot just do it with bigger and bigger, faster and faster computers. We must change our focus to effective, productive, innovative engineering, which means fresh approaches to new problems. This means new training methods, new management techniques, innovative organization patterns, and simplified analysis and test techniques.

I wish we had time to talk in detail about each of these discipline areas; we do not; however, I have chosen only a few to look at and will leave the main task to the experts in the various sessions, which is really the purpose of this conclave. In Figure 12 we see a more detailed description of structural dynamics from an overall viewpoint, providing additional detail over the previous slide. Key problem areas are more accuracy, faster turn-around, operational verification, and special testing. High performance is a parallel complexity factor for these disciplines, particularly in terms of life-time which implies accurate environments, material characteristics, and fracture mechanics. The Shuttle Main Engine is an excellent example of this situation, high energy concentration, weight and volume constraints, with high temperature and pressure. This results in high static or mean stress with very small allowances for alternating stress before the endurance limit is reached. This implies that the high cycle alternating stress levels operate on the flat part of the S-N curve, creating high sensitivity to small changes in alternating stresses (Figure 13). Either one must reduce the mean stresses or increase the endurance limit (material choice) to solve the problem or accurately predict the environments and the dynamic characteristics. Figure 14

illustrates this situation for the Space Shuttle Main Engine, where the latter approach was taken when the engine performance requirements were upped nine percent to meet additional Space Shuttle performance requirements. All essential elements of the lifetime problem are shown. Key to solving the SSME fatigue issues has been threefold, (1) structural dynamic test (model verification), (2) materials selection and material properties characterization, and (3) hot firing ground test environment and response measurements, lifetime verification being accomplished through the hot firing ground certification program. Very accurate predictions and their verification were the key to getting the Shuttle at the operational stage it is at today.

Figure 15 treats the area of component criteria in more detail. Here, we have gone from a limited data bank, single-axis prototype testing approach to the future, requiring multi-axis with accelerated time testing. Data banks must be extended to three dimensions. Analysis must consider multi-3D modes instead of single-axis, single modes as well as alternate approaches, such as SEA (statistical energy analysis). Pattern recognition will be a key development area in conjunction with analytical/operational verification of many subsystems and components. This results from the large volumes of data with many parameters that we must evaluate.

The area of dynamic response has made great strides moving from rigid-body analysis to limited number of elastic modes to the present Space Shuttle system response analysis of 400 modes, including wind, thrust parameters, control, etc., in a Monte Carlo analysis (Figure 16). Approximately 30 parameters are varied in the analysis. Parallel with this is the very accurate jitter analysis of optical systems, such as Space Telescope where response of the optical system to noise and control devices (momentum wheels) must be kept to very low values (arc milliseconds (0.0087)). Modes through 120 Hz are required for this analysis. In the future, unlimited number of modes in conjunction with growing structures which are designed from stiffness will come into being. Many of these structures will be very complex, composed of many elements in an unsymmetrical manner. Localized nonlinear damping will dominate the response creating new analysis techniques as well as definition and verification techniques. Total system analysis with appropriate trades will be involved and is a major challenge.

Structural modeling has made great strides from both the analytical approaches and testing standpoint. Models have moved from equivalent beams to large finite element systems. General purpose finite element programs exist, such as NASTAN, SPAR, and industry peculiar codes. Nonlinear and equivalent macro element modeling looms on the horizon (Figure 17). Testing of very large systems on the ground is now

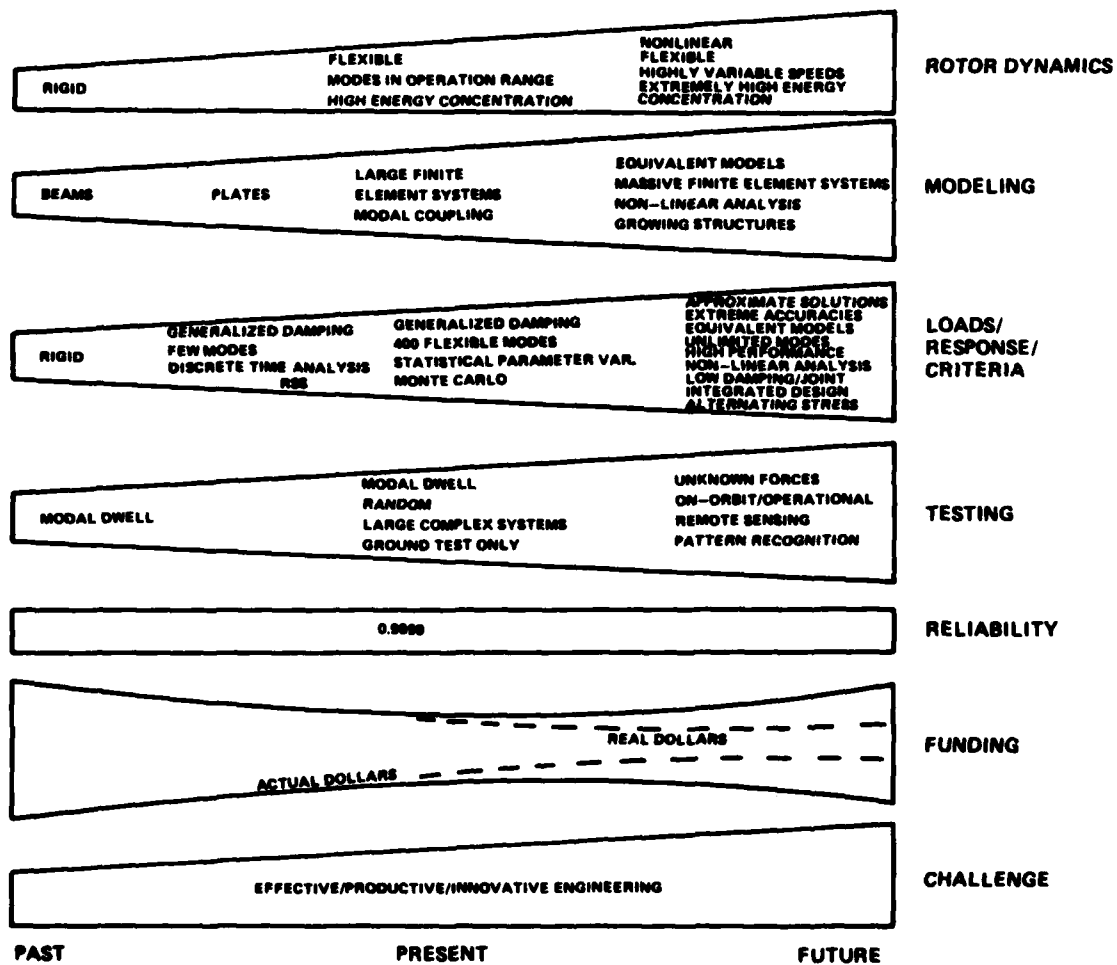
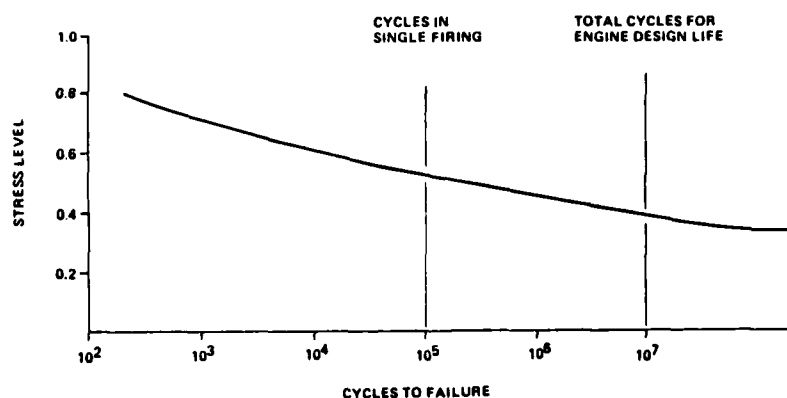


Fig. 11 — Past, Present and Future, Major Technical Disciplines and Challenges

MODELING	BEAMS	TEST VERIFIED LARGE NUMBER FINITE ELEMENTS GENERALIZED DAMPING FIXED SYSTEM	ANALYTICAL/OPERATIONS VERIFIED NON-LINEAR SUPER FINITE ELEMENT SYSTEM LOCAL/JOINT NON LINEAR DAMPING GROWING STRUCTURE
ROTOR DYNAMICS	RIGID	HIGH ENERGY CONCENTRATION FLEXIBLE MODES IN OPERATION RANGE	EXTREME ENERGY CONCENTRATION NON-LINEAR FLEXIBLE HIGHLY VARIABLE SPEEDS HIGH PERFORMANCE LOW WEIGHT
LIFETIME	MINUTES	HOURS	YEARS REFURBISHMENT
RESPONSE	RIGID	GENERALIZED DAMPING 400 FLEXIBLE MODES STATISTICAL PARAMETER ASSESSMENT	LOCAL/JOINT/NONLINEAR DAMPING UNLIMITED MODES EXTREME ACCURACIES (PERFORMANCE) GROWING STRUCTURES APPROXIMATE SOLUTIONS EQUIVALENT STRUCTURES OPTIMIZED DESIGN
TESTING	MODAL DWELL	RANDOM TIME MODAL DWELL GROUND (FULL AND SCALE)	KNOWN FORCE UNKNOWN FORCE ON ORBIT/OPERATIONAL REMOTE SENSING
ENVIRONMENTS	TEST	TEST ANALYTICAL APPROXIMATIONS NONDIMENSIONAL	COMPUTATIONAL ANALYSIS OPERATIONAL ADJUSTMENTS
VIBRATION CRITERIA/TESTING	LIMITED DATA BASE	SINGLE AXIS EXTENDED DATA BANKS LIMITED ANALYTICAL APPROACHES PROTOFLIGHT	MULTI-AXIS TIME ACCELERATION OPERATIONAL VERIFICATION
PAST		PRESENT	FUTURE

Fig. 12 — Overall View of Structural Dynamics

GENERAL PROBLEM STATEMENT ON LIFETIME PREDICTIONS



WEIGHT, VOLUME, SYSTEMS PERFORMANCE CONSTRAINTS, AND ENVIRONMENTS DO NOT ALLOW DESIGN BELOW THE ENDURANCE LIMIT; THEREFORE, SMALL CHANGES IN ALTERNATING STRESS CAUSE LARGE CHANGES IN LIFETIME

Fig. 13 — General Problem Statement of Lifetime Predictions

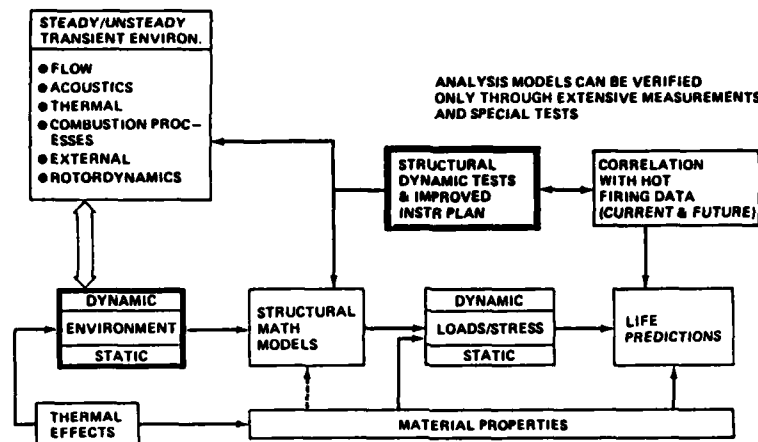


Fig. 14 — Space Shuttle Main Engine Dynamic Problems and Methods for Their Solution

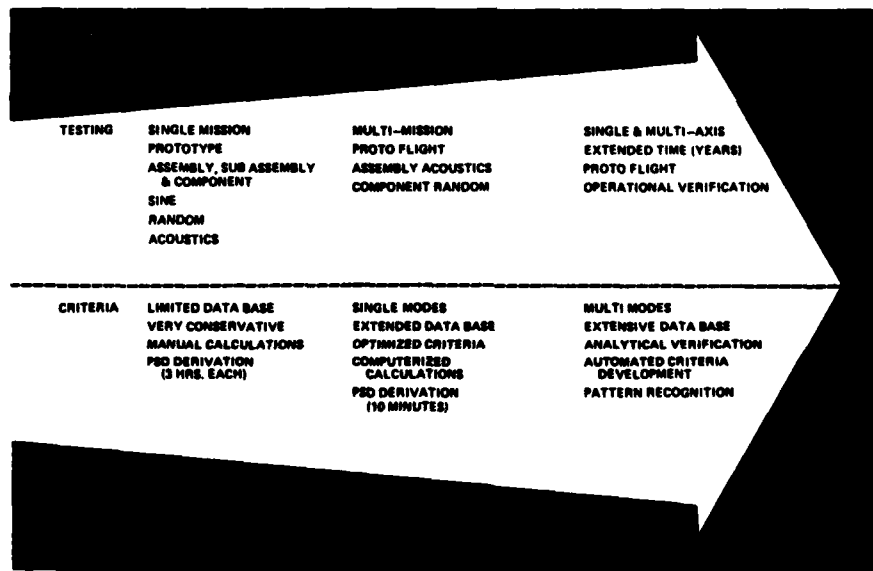


Fig. 15 — Advances in Vibro-Acoustic Component Criteria



Fig. 16 — Advances in Structural Dynamics Response Analyses

routinely accomplished. The largest to date was the full scale launch configuration of the Space Shuttle. Many space configurations on the planning boards cannot be tested on the ground, requiring analytical verification and/or on-orbit operational testing. Determination of localized, joint damping for large space structures, in particular for fine pointing, high performance systems, will be required. This will be a major challenge.

Let me move next to what are some very basic and some new challenges I believe we face that must be solved if the goals of NASA are met. I believe many of these challenges also transcend into the various industries and programs you are concerned with.

Future space missions, in particular the satellites and the Space Station, move conceptually into a more complex regime. Figure 18 illustrates this in two aspects, (1) design approaches and (2) expected lifetime. Notice in the past, space vehicle designs were strength designed with large safety factors tested to acceptable limits and, in general, the operations time was short. In the present, we are still in the strength design regime; however, safety factors are limited. NDE and fracture mechanics are used. Fatigue is a constant concern requiring much attention. Many structures are analytically verified instead of test verified, particularly at the system level. Operation time is still short with the exception of a few spacecraft and satellites. Future plans move from a strength design approach to systems that are designed for stiffness, including very accurate control on deformations and responses. Refurbishment and maintenance must constitute a prime part of the engineering tasks, integrated design approaches, in conjunction with analytical and operational verification techniques. Figure 19 shows the dilemma we have in structures/structural dynamics disciplines. The overall conflicting technical requirements of increasing cost, time, complexity, and risks versus programmatic requirements of decreasing cost and time lead to five major problems which must be solved in the near future to meet near-term goals.

1. Design loads cycle time/complexity. One-year lead cycles must be reduced to approximately three months or less if future Shuttle manifests are met.

2. Payload experiments response accuracy (loads) is a very real problem. Current indications from first Shuttle flights indicate that experiment responses are being grossly overpredicted. Analytical system models also show extreme sensitivities to small system changes which are obviously incorrect.

3. Fine pointing requirements of systems, such as Space Telescope, antennas, etc., are requiring extremely accurate models and knowledge of subsystems.

4. Qualification and verification using protoflight (flight article) testing or analytical verification moves us into a new regime.

5. High performance induced problems, such as lifetime and quality control.

Currently, we are attacking these problems with finer models, larger, faster computers, graphics, detailed statistical assessment, and detailed testing. The challenge is to move to innovative approaches that use equivalent models, integrated design, organization adjustments, and motivational and educational programs.

Figure 20 illustrates this changing approach for large space systems versus the traditional. In the past, the structure was analytically characterized and test verified. This structural model is used to design the control system which is then test simulated. As indicated with the arrows, feedback occurs between the various design and verification activities producing a finely tuned and verified systems before it flies. Large space systems cannot follow this traditional approach. These structures cannot be ground tested as a total unit. Only limited element tests can be performed. This means that the control design is accomplished using analytical models with verification accomplished in simulations. This means either the control system must be very complex, such as adaptive systems, so that it is not sensitive to unknown or unpredicted structural characteristics or the system must be changed on-orbit. To accomplish the latter requires on-orbit structural dynamic characterization with the ability of control system logic update to accommodate these changing structural characteristics. The system is further complicated by the requirement for changing and growing configurations as missions and uses evolve. This figure also illustrates some of the concepts and programs now underway, including two planned flight experiments, SAFE and SADE.

Figure 21 illustrates some challenging concepts for construction of these systems using common elements in both volume and trusses. Building block design and verification tools are a large part of concept as well as assembly techniques. Many other options exist including deployables, erectables, on-orbit manufacturing (beam machines). The book is still open on the approaches to be used.

As stated previously, the solutions to the challenges that are on the forefront of our disciplines must be solved in innovative ways. Computer graphics, computer-aided design, and manufacturing, pattern recognition tools and special software are some of the current techniques that need further development. Figure 22 shows a graphics work station with special software that allows the engineer to take computer-aided design tapes and build a

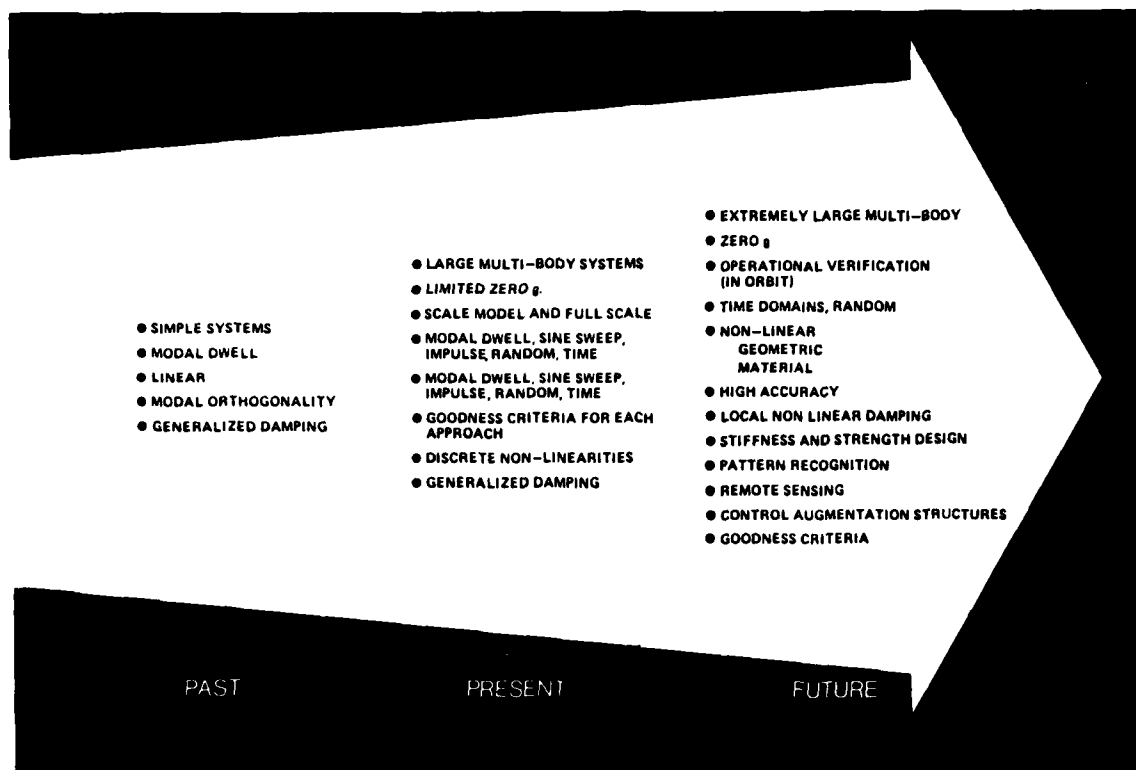


Fig. 17 — Advances in Structural Modeling and Testing

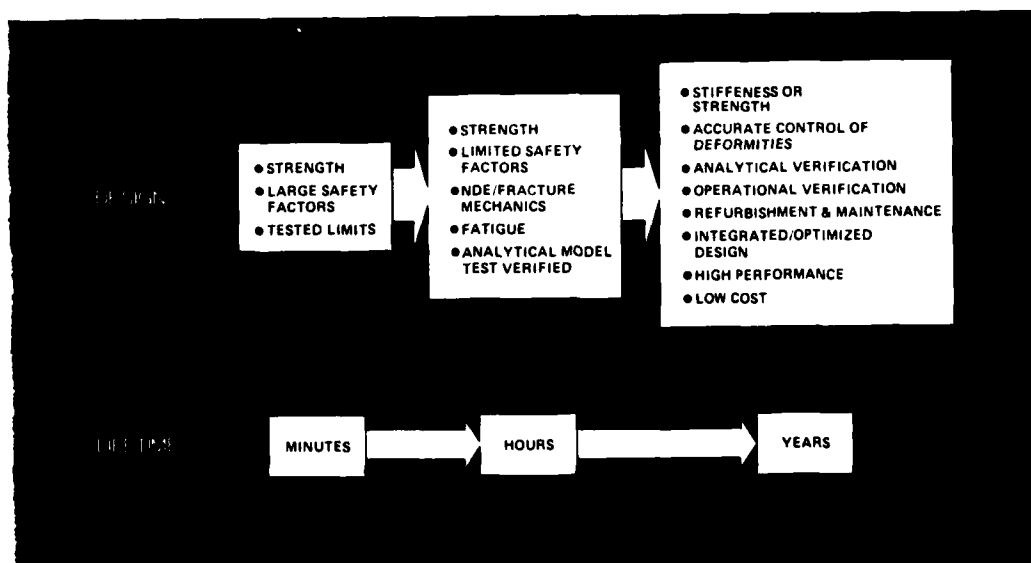


Fig. 18 — Future Design Requirements for Satellites and the Space Station

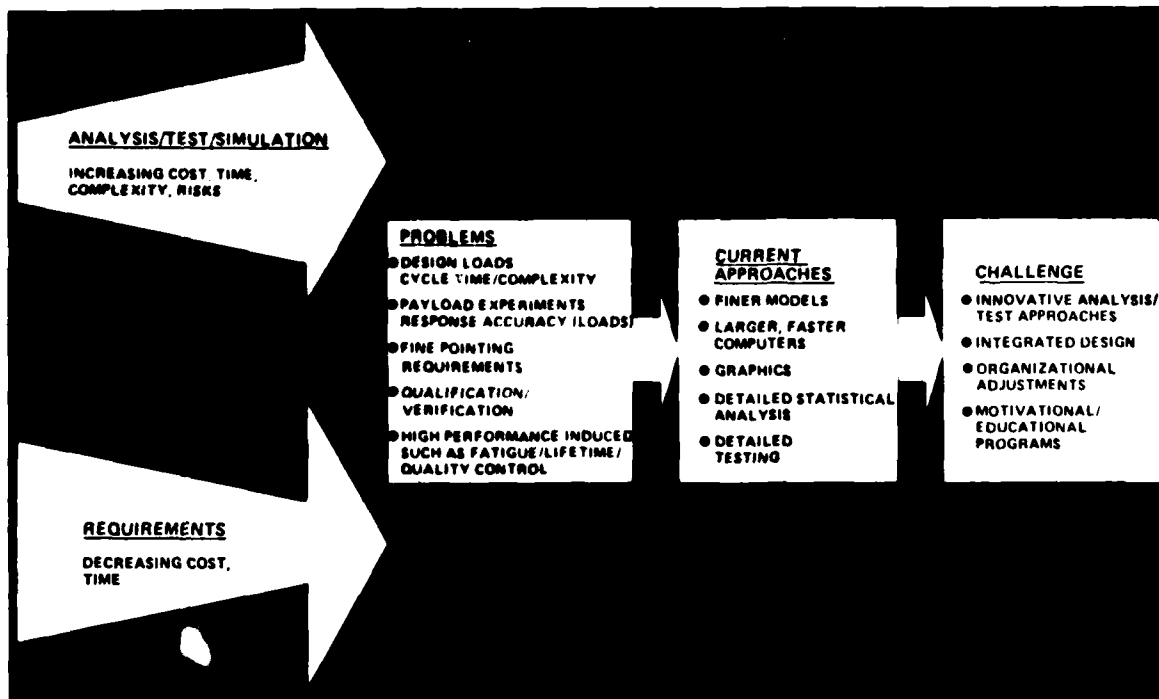


Fig. 19 — Problems, Current Approaches and Future Challenges in Structural Dynamics

LARGE SPACE SYSTEMS TECHNOLOGY (STRUCTURAL DYNAMICS/CONTROL/TEST INTERACTION)

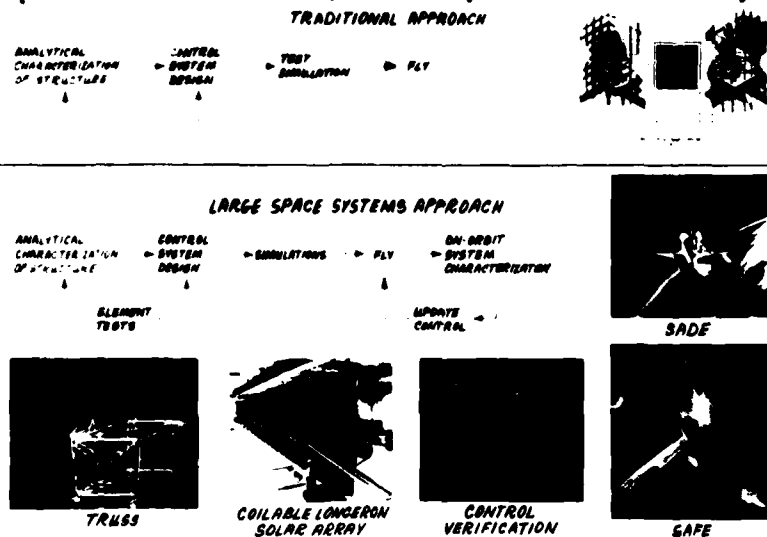


Fig. 20 — Large Space Systems Technology (Structural Dynamics/Control/Test Interaction)

COMMONALITY OF SPACE STATION STRUCTURE ELEMENTS

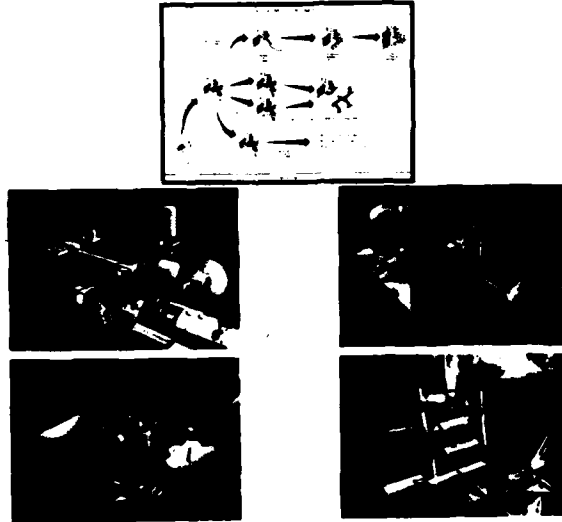


Fig. 21 — Commonality of Space Station Structure Elements

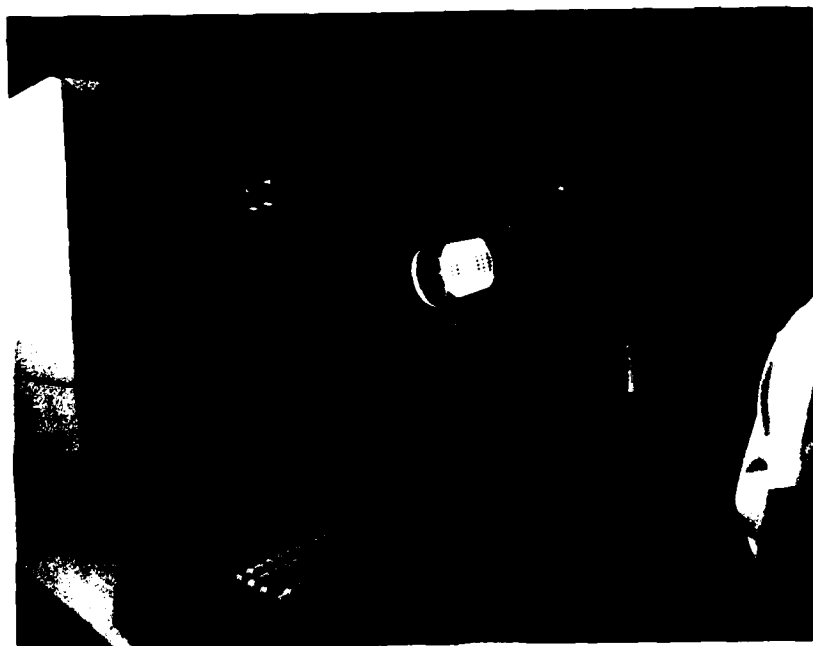


Fig. 22 — Computer Graphics Work Station

finite element structural model with the results of a mass and stiffness matrix tape. Models developed by subcontractors, etc., can be checked quickly for errors using this system. Obviously, these present systems are the doorways to exciting things in the future.

The chart in Figure 23 summarizes some of the technology gaps that exist in the rapidly evolving field of structural control interaction, a discipline that is very exciting and has a very special interest to me. Three discipline areas are used, (1) structures, (2) controls, and (3) systems. Gaps are developed in three broad areas, (1) techniques, (2) tools, and (3) test (verification parameter data). Techniques deal with approaches to solving known problems, whereas tools describe techniques required to apply the techniques. Tests deal in a generic sense with verification. Generally, the areas discussed previously are contained in this matrix. Readers can study it in detail for more insight. Clearly, we must get ready for this exciting multi-discipline challenge to the future associated with space stations and large space structures.

Figure 24 summarizes some of the challenges that we face if the NASA goals and objectives are to be met. Clearly, some of these challenges as stated are controversial. Many think we are already doing these things. To some extent, I agree; however, it is a matter of degree. We must make further, larger steps if we are to be successful. Organization structure must be under continuous evaluation if integrated system design, etc., is accomplished. Discipline-oriented organizations can be a detriment to this type analysis. Obviously, analysis time must be greatly reduced and productivity increased. Many would like to continue separate external and internal loads analyses. We must re-evaluate this to see if we do not need to remove this conservative approach by accomplishing dynamic stress analysis in lieu of using equivalent static external loads to determine internal stresses. We must greatly enhance our ability to do computational fluid analysis, particularly in the area of internal flows such as rocket engines. More emphasis must be placed on designs that are amenable to maintenance, growth, and quality control.

Again, I want to emphasize the requirement for innovative pattern recognition techniques as we deal with more and more data as a function of ever increasing numbers of parameters. Information not interpretable is useless. Payload analysis time must be cut by approximately an order of magnitude as well as improvement in overall productivity.

With the chart on Figure 24, I close what to me are some key challenges we face. Our major task is defining the approaches for meeting these challenges. What is the starting point becomes the key question. A poet, hundreds of years before Plato said, "Before the gates of excellence, the high gods have placed

sweat. Long is the road thereto and steep and rough at the first, but when the height is won, then is there ease." In all of this sweat, however, we need time to think, to meditate. We need leisure for good ideas to work their way into our consciousness. Remember, ideas can be worth ten years of hard eight-hour-a-day work. Finally, whenever we have problems, we must follow the grand old rule, "Go back to the basics," a rule every athlete knows well. You do not solve the game, lick the course, at best, you must keep playing and living with it and going back to basics. Being a wood worker by hobby has taught me repeatedly the lesson of basics, a lesson I believe applies to our engineering trade as well.

1. Tools must be sharpened properly and finely honed.
2. Alignment must be very accurate.
3. Special jigs are mandatory to accomplish many tasks correctly.
4. Materials must be of highest quality.
5. Work must be very accurately laid out. Measure and remeasure again before cutting.
6. Product must be hand rubbed and polished to produce a fine finish.

Our answers lie, then, in (1) sweat, (2) leisure time for germinating ideas, and (3) a proclivity for going back to the basics. I believe that is what this meeting is all about, so let's get on with it.

BIBLIOGRAPHY

1. Anderson, Frank, Jr.: Orders of Magnitude, NASA SP-4404, 1981
2. Haggerty, James: Spinoff 1983. National Aeronautics and Space Administration, May 1983
3. The User's Guide to Spacelab Payload Processing. NASA John F. Kennedy Space Center, March 1983
4. Levine, Arnold: Managing NASA in the Apollo Era. NASA SP-4102, 1982
5. Spitzer, Cary R. (Editor): Viking Orbiter Views of Mars. NASA SP-441, 1980
6. Newell, Homer: Beyond the Atmosphere, Early Years of Space Science. NASA SP-4211, 1980
7. Schneider, William C. and Hanes, Thomas E. (Editors): The Skylab Results, Advances in Astronautical Sciences. American Astronautical Sciences. American Astronautical Society Publication, 1975
8. Life in the Universe, Proceedings of a

APPLICATION OF MODERN DISTRIBUTION CONTROL CONCEPTS TO VERY FLEXIBLE STRUCTURE (GAPS)

	TECHNIQUES	TOOLS	TESTS (VERIFICATION) PARAMETRIC DATA
STRUCTURES	● LIGHTWEIGHT, ETC., PASSIVE DAMPERS ● JOINT DAMPING ● AUXILIARY STIFFENING ● DEFINITION OF PARAMETER VARIATION DATA BASE ● REMOTE SENSING ● REALTIME I. D. AREAS OF GAPS ● PARAMETRIC DATA ● COMPUTATIONAL AND ANALYTICAL TOOLS ● TEST TOOLS	● TRUNCATION TOOLS ● STATE IDENTIFICATION ● NONLINEAR ● GEOMETRIC ● MATERIALS ● JOINT DAMPING ● POINT LOADS ● PARAMETER VARIATION DEFINITION ● GROWTH ● TRANSIENT I. D. AND DYNAMIC CHARACTERIZATION ● MACRO-ELEMENTS ● STATISTICAL COMBINATION ● LOW G. SIMULATION TECHNIQUES ● ON-ORBIT ENVIRONMENT CODES	● ON-ORBIT ● EXCITATION ● ACQUISITION ● FORCING FUNCTION ● NONLINEAR EFFECTS ● GOODNESS CRITERIA ● VERIFICATION OF TOOLS ● SENSITIVITY DATA DEFINITION AND COMBINATION VERIFICATION ● GROUND SIMULATION OF ZERO g. ● JOB STRUCTURE ● SCALE MODEL
CONTROLS	● DISTRIBUTED CONTROL TECHNIQUE (PRACTICAL) ● SHAPE CONTROL ● PRACTICAL EVOLUTIONARY CONTROL ● DISTURBANCE ISOLATION ACTIVE / PASSIVE ● ACTUATOR/SENSOR PACKAGE FOR DISTRIBUTED CONTROL ● ROBUSTNESS (INSENSITIVE)	● DESIGN TOOL FOR DISTRIBUTED ● DESIGN TOOL FOR EVOLUTIONARY	● DISTRIBUTED CONTROL AND SHAPE CONTROL G/F/H ● EVALUATION - GROUND ● ISOLATION - GROUND
SYSTEMS	● OPTIMIZATION APPROACHES ● BLENDING BETWEEN CONTROL AND STRUCTURE ● DEFINITION OF PARAMETER VARIATION DATA ● MEANS OF ACCOMMODATING ORBITAL ENVIRONMENT DISTURBANCE EFFECTS ● ROBUSTNESS (INSENSITIVE)	● OPTIMIZATION APPROACHES (LARGE NUMBER OF MODES, COST CRITERIA, PERFORMANCE CRITERIA) ● STATISTICAL COMBINATION APPROACHES ● ON-ORBIT ENVIRONMENT CODES	● VERIFICATION OF STATISTICAL COMBINATION APPROACHES ● VERIFICATION OF STRUCTURAL CONTROL INTERACTION CONCEPTS

Fig. 23 — Technology Gaps in the Field of Structural-Control Interaction

- DESIGN FOR MAINTENANCE, REFURBISHMENT, GROWTH, QUALITY CONTROL ENHANCEMENT, AND NDE/FRACTURE MECHANICS.
- DESIGN AND VERIFICATION FROM AN INTEGRATED SYSTEM STAND POINT CONSIDERING ALL DISCIPLINES SIMULTANEOUSLY.
 - INNOVATIVE ORGANIZATIONAL APPROACHES
 - INNOVATIVE TRAINING APPROACHES/ENGINEERS KNOWLEDGEABLE IN SEVERAL DISCIPLINES
 - INNOVATIVE ANALYSIS TOOLS
 - SINGLE MULTI-DISCIPLINE SAFETY FACTOR
 - MATERIAL SELECTION/COMPOSITES
- DESIGN SYSTEM TO BE FORGIVING AND ROBUST
- DEVELOP INNOVATIVE PATTERN RECOGNITION AND EVALUATION TOOLS TO HANDLE LARGE DATA SETS ASSOCIATED WITH LARGE COMPLEX SYSTEMS
- DEVELOP APPROACHES FOR REPLACING QUASI-STATIC EQUIVALENT LOADS WITH DYNAMIC STRESS DESIGN LOADS
- REDUCE ANALYSIS TIME FOR PAYLOAD LOADS INTEGRATION BY FACTOR OF 3 TO 10
- INCREASE PRODUCTIVITY THROUGH IMPROVED ANALYSIS TECHNIQUES, TESTING TECHNIQUES, AND MANAGEMENT APPROACHES
- EXTEND THE CAPABILITY FOR COMPUTATIONAL FLUID AND STRUCTURAL ANALYSIS SIGNIFICANTLY

Fig. 24 — Challenges in Meeting NASA's Goals and Objectives

Conference held at NASA Ames Research Center, Moffet Field, CA, NASA CP-2156, June 19-20, 1979

9. Bilstein, Roger E.: Stages to Saturn, A Technological History of the Apollo/Saturn. NASA SP-4206, 1980
10. Morrison, David: Voyages to Saturn. NASA SP-451, 1982
11. Fitzgerald, Jr., Paul E. and Gabris, Edward A.: The Space Shuttle Focused Technology Program: Lessons Learned. Astronautics and Aeronautics, February 1983
12. Sackheim, Robert L. (Editor): Liquid-Rocket Propulsion Technology. Astronautics and Aeronautics, April 1983
13. Fitzgerald, Jr., Paul E. and Savage, Melvyn: Cost as a Technology Driver. Astronautics and Aeronautics, October 1976
14. Abrahamson, Lt. Gen. James A. (USAF): Shuttle Operational Expectations. Astronautics and Aeronautics, November 1982
15. Wilson, George C.: Space Program Lifting Off to New Regime. Astronautics and Aeronautics, July/August 1983
16. Olstad, Walter: Targeting Space Station Technologies. Astronautics and Aeronautics, March 1983
17. Future Space Transportation Systems Study. Astronautics and Aeronautics (4 articles), pp. 36-56, June 1983
18. Moore, Jesse W.: Effective Planetary Exploration at Low Cost. Astronautics and Aeronautics, October 1982
19. Special Section on Interactive Computer Graphics. Astronautics and Aeronautics, April 1983
20. Special Section on Space Station Technology. Astronautics and Aeronautics, March 1983
21. Diaz, Al and Casani, John R.: Galileo 1986 on Centaur. Astronautics and Aeronautics, February 1983
22. Special Section on Artificial Intelligence. Astronautics and Aeronautics (5 articles), July/August 1983
23. Keaton, Paul: Why Billions Can and Should be Spent on Space. Astronautics and Aeronautics, May 1983
24. Card, Michael F.: Trends in Aerospace Structures. Astronautics and Aeronautics, July/August 1983
25. Amos, Anthony K. and Goetz, Robert C.: Research Needs in Aerospace Structural Dynamics. 20th Structures, Structural Dynamics, and Materials Conference, AIAA, St. Louis, MO, April 4-6, 1979
26. Morosow, George; Dublin, Michael; and Kordes, Eldon E.: Needs and Trends in Structural Dynamics. Astronautics and Aeronautics, July/August 1978
27. Garrick, I. E.; and Reed, III, Wilmer H.: Historical Development of Flutter. AIAA/ASME/ASCE/AHS 22nd Structures, Structural Dynamics, and Materials Conference, Atlanta, GA, April 6-8, 1981
28. Garrick, I. E.: Aeroelasticity - Frontiers and Beyond. AIAA Journal of Aircraft, September 1976
29. Ryan, R. S.; Salter, Larry; Young, George M., III; and Munafo, Paul M.: SSME Lifetime Prediction and Verification, Integrating Environments, Structures, Materials, The Challenge. Space Shuttle from Challenge to Achievement Conference, Johnson Space Center, June 1983
30. Ryan, R. S. and Jewell, R. E.: A Preliminary Look at Control Augmented Dynamic Response of Structures. AIAA Structures/Structural Dynamics/Materials Conference, Lake Tahoe, Nevada, May 1983

INVITED PAPERS

DNA ICBM Technical R&D Program

Col. Maxim I. Kovel
Defense Nuclear Agency
Washington, D.C.

INTRODUCTION

For those of you who are unfamiliar with the Defense Nuclear Agency (DNA), we are not quite as well known as NASA, I would like to go over a brief introduction, and then I will describe the DNA program having to do with the basing of primarily small missiles. The kinds of basing we are talking about are hard silos and hard mobile launchers. I will also include a few details from the Ballistic Missile Office program since it is a parallel program. I should also mention the whole program has some funding questions associated with it. I don't know if you follow the Congressional actions; the House Appropriation Subcommittee deleted some resources from the Air Force program which covers hard silos, and the Senate Appropriation Subcommittee kept them in. So now we are all waiting for full committees and joint committees to get together to determine if what you are about to see will really take place to the extent that I will point out. It is all beyond our control.

(Fig. 1) DNA has operational roles, and it also has research and development roles. Probably two thirds of our budget is really in the research and development area; the other third is for operations. We work for the Undersecretary of Defense Research and Engineering and for the Joint Chiefs of Staff. We interact with both military and civilian agencies; and within the services, we interact with subagencies such as the Ballistic Missile Office of the Air Force (BMO).

Our Headquarters are in the Washington, DC area. The lower right hand corner of Figure 1 briefly shows our organization. DDOA stands for the Deputy Director for Operations and Administration. I report to the Deputy Director for Science and Technology (DDST), I will expand on that later. FCDNA is our Field Command which is located in Albuquerque. They field the DOD underground nuclear tests, and they also field high explosive tests such as DIRECT COURSE which is scheduled to take place on the 26th of October. AFRRI is the Armed Forces Radiobiological Research Institute which is located near the Naval Hospital in Bethesda,

Maryland, and they are concerned with the biological effects of radiation.

Figure 2 depicts our role relative to the services and to the national laboratories. We do not develop weapons; they are developed by the services. We don't build warheads; that is done by the Department of Energy. We consider the nuclear effects environment, the vulnerabilities, and the lethalties of systems. Our primary role is effects research and testing. We also look for ways to improve the hardness of existing systems against nuclear effects and possibly the directed energy effects of the future. Since we explore advanced concepts, we are mostly a technology organization.

Within the organization of the Deputy Director for Science and Technology (DDST) is the Shock Physics Directorate which is made up of three divisions (Fig. 3). The Hard Mobile Launcher program has a program manager in the Aerospace Systems Division. The Hard Silo program has a program manager in the Strategic Structures Division. We have combined program operation as closely as possible with the technology base program to maximize the resources available. The Assistant for Experimental Research is Dr. Gene Sevin, and he is overseeing the total ICBM basing program for the DDST. I might mention that the Radiation Directorate is working on small missile electronics as well as the electronics involved in the basing aspect. Within the Aerospace Systems Division we are also looking at the small missile itself in terms of its hardening and vulnerability.

Since I am the Director of the Shock Physics Directorate, I thought I'd tell you what the Directorate does (Fig. 4). We are primarily concerned with the mechanical effects of shock, thermal and nuclear radiation on mobile and fixed weapon systems, and that is why both the hard silo and the hard mobile launcher are within this Directorate. We also supervise, from the Headquarter's standpoint, the Underground Nuclear Weapons Effects Test program

at the Nevada test site as well as the high explosives site. I mentioned DIRECT COURSE, and for those of you who are not familiar with DIRECT COURSE, it is a one kiloton simulation at a scaled height of burst which will be conducted at White Sands. A sphere about 33 feet in diameter and about 167 feet above the ground will be filled with 600 tons of an ammonium nitrate and oil mixture explosive. It will have nearly 200 major experiments in the area surrounding it. This event will occur on the 26th of October assuming we have no more lightning storms to destroy all of our instrumentation. (This was reference to a lightning storm in September which destroyed a large number of gages.)

DNA ICBM BASING R&D PROGRAM

The basic objective of both the Air Force and the Defense Nuclear Agency is to provide decision-makers with as much information as possible in selecting a basing mode within a relatively short period, something like three years. The two basing modes we are considering are silos and mobile hardened launchers for small missiles. We are particularly interested in two things in regard to the silos. We are interested in how hard we can build them, and we are also interested in the range of hardening to see what is most feasible and cost-effective. Just making it very hard is not the whole solution; and of course, we don't know exactly how hard we can build them. In regard to mobile small missile launchers, we are primarily interested in the hardened launcher itself. It is important to realize that we are talking primarily about the basing of small missiles for the present. But, should anyone want to consider a larger missile, a small missile is roughly half the size of an MX. So everything is well within the scaleable range of science.

Figure 5 shows a schematic of a generic small missile. The gross weight of the missile itself is only in the 30,000 pound range, and this is Congressionally mandated. The range is about 6,000 nautical miles; it is four feet in diameter and 44 feet long. To handle this, the size of the silo is about 110 feet deep and about 12 feet wide when you put the missile and equipment into it. So the silo is considerably bigger than the missile. In regard to the size of the mobile launcher itself, you could take roughly two or maybe two and a half times this to get the total size of the vehicle with the launcher on it.

The biggest problem will be in the terminal guidance system as far as the missile itself is concerned. Basically, I think we know how to build missiles since we have been building missiles for a long time; but if it will be a small missile, you have to make it very accurate, and a lot of effort will have to go into that. How to mount all of these systems within the missile and have them survive the

environments the missile may have to fly through is a key problem within the small missile community.

Of course, the question sometimes becomes, "Why should DNA be so involved in developing some of the silos?". An extract from the report of the President's Commission states that DNA should have a major role in the basing decision. It provides the impetus by which we are jointly working with the Air Force Ballistic Missile Office. If you look back at the DOD authorization bill, you will find that they specifically identified money to be provided to DNA by the Air Force since Congress wanted to ensure that people from the technology base, who are not specifically systems oriented, would be involved in the review and consideration of how hard you can make a basing system. It is a parallel program which is very closely integrated with the Air Force; and they are looking to DNA to provide information on environments, simulation, and instrumentation development and techniques. We are talking about superhard silos, but still you must have missiles and equipment inside that can withstand associated shocks coming through. At the same time we are also considering low level pressures for a missile or a launcher that will be sitting on the ground. We must find a way to have something with a missile on it survive an appropriate overpressure environment.

We will have a program in advanced hardening technology for silos, and we are looking at what the future might hold. We are considering ways of being clever about making something hard against very high shocks or very strong ground motion. That cannot be done cheaply with the current technology; and it cannot be done completely effectively with the current technology, so we are looking for better ways to approach it.

ADVANCED SILO HARDNESS R&D PROGRAM

Figure 6 is a summary of the major objectives of this program. I think the first one, resolving uncertainties, is very important. Those uncertainties are primarily in the environment definition. One of the most important considerations is the cratering. This is because we have gone from concern for mostly an airburst to the point where we are concerned with surviving a ground burst. If a ground burst occurs, and if it produces a large crater, how big is that crater? There is little point in building a very hard silo if it will be within the crater because it will not survive if it is flipped over subjected to high accelerations or moved completely out of firing capability orientation by the ground motion. So there are some problems that have to be resolved with regard to the environment. Also, materials problems are involved; how do you put those materials together into a structure that will survive the associated effects?

Besides the shell itself, or the silo, you have to worry about the shock isolation system. Most of the work on the shock isolation system will be done by the Air Force. This initial program, by the way, is to arrive at a solution, or concepts which may resolve the problem, within about a 3-year period. That would yield information which would assist someone in deciding how to base a missile before they go into Full Scale Engineering Development (FSED). Much of the work will then continue during the FSED period. As you know, the mobile launcher concept is more popular with the Congress than this is. Everybody would like to see a mobile launcher, but it has its problems. It is a little more expensive, it requires more people, and it requires a little more space; but on the other hand, it guarantees that you will concentrate on a small missile. Those may be some of the considerations. You will not have a large missile on the launcher.

From a technical standpoint we will evaluate any advance silo designs that we can come up with. We will develop simulators, and that raises another problem. We know how to simulate one kilobar; we have done that. We can simulate one and one half kilobars, but we can only simulate a part of the environment at one time. We do some simulation with high explosives. We do some simulation with underground testing, and we do some simulation with laboratory tests. But you can't put them all together unless you can test in the atmosphere, and I don't think that is in the cards.

We have to be able to measure what we find, and that is part of simulation and instrumentation (Fig. 7). Within the simulation development and the instrumentation development, we are having great difficulty with gauges being able to measure the environment that we are actually creating. We are not sure of the environment we are creating. The very high pulse spike at the beginning of the air blast simulation tends to wipe out the air blast gauges. We can make some stress measurements, but we are not sure how far you can back those out. That is the way calculations are being done. You look at the results of many of the experiments that are going on now; and you think you know what it is, but you are not sure how it got there.

I mentioned the cratering program; and I particularly want to point out that we hope to go back to the Pacific Proving Ground to reexamine the craters, and try to understand how they were actually formed. Their formation mechanism will then go into the codes which we will use to either verify, deny, or argue with the calculations that we currently use. Remember, I said the cratering is most important in determining whether it makes sense to build a very hard silo. So we have a considerable investment in that area.

Figure 7 also summarizes some of the things that are being done in the environments area. We must simulate one to six kilobars, both height of burst and surface burst. We are concerned with the soils in the area where the silo will be located, and we are concerned with materials from which the silo is constructed. We are also interested in simulating ground shock and cratering. We have a test program to develop what we call a cratering and related effects simulator, and this program will be going on for about the next two years. We have a near source simulator in the Yuma Proving Ground Area, which is the first part of this test program, and hopefully it will go off sometime in December. I have already mentioned instrumentation, and that program is extensive in developing the proper instrumentation. Most of the silo field testing will be done by BMO on intermediate and large scale structures, but we will do some testing on a 50 inch diameter silo to develop new concepts. These are differences in the effects of air bursts and ground bursts. I mentioned the problem of the silo being in the crater earlier; however, the silo being buried under a large amount of debris presents another problem. How do you get the missile out if you are covered over with 30 or 40 feet of debris because you are still within the lip area of the crater? The Air Force is working on that problem.

Some historical trends in silo hardening are apparent (Fig. 8). These include the Minuteman and the MX Baseline. Super hard is just taking the baseline and putting a little more steel in it and making some minor changes to it. The ultra-hard silo is another possible concept and we have done some testing. An ultra-hard silo that has about 3% steel and a liner inside was designed by Weidlinger and Karagozian. It was built at the Waterways Experiment Station, and it was tested there in November, 1982. By brute force, one can get up pretty high; but it is a matter of how much steel you put into it and how big you make it. (Fig. 9) Techniques are available for building ultra hard silos. One idea is to decouple silos from the ground motion, and it shows what to do with a silo if you are really clever and you want to isolate it from the shock, or have part of the silo take the shock loading and be destroyed while the rest of the silo survives. If you have some other good ideas on how to decouple structures in the ground from the ground motion, we would love to hear about them.

I mentioned the shock isolation system. We had some shock isolation mockups in some of our tests. We used two types of shock isolators; one was a spring-hydraulic type, and the other was a straight hydraulic type mounted either on the silo wall or on the mockup of a cannister inside. I mentioned earlier that the Air Force is looking at this program. They are doing a lot of work on shock isolation systems; they are working primarily with Boeing. Figure 10 shows some alternate vertical shock isolation system

concepts. This shows the enhanced hardness concept, but we are going more toward the superhard concept now. This is what is possible with advanced technology; hydropneumatic springs, dual isolators, and computer controlled active damping. Figure 11 shows some alternate lateral shock isolation system concepts. At present we put lots of foam around the missile. They are looking at ways of installing recoverable dampers so if you have more than one shock, you will not end up with the crush-up on the first blast taking away your capability for eliminating lateral shock. The Ballistics Missile Office is considering the construction of a shock isolation testing facility which would be capable of doing a full scale test on the silo, at least at the small silo dimensions (Fig. 12). Again, I am not sure if it will be built now, but it was originally in the program for the Air Force.

Another consideration for reducing the ground shock motion is to do something before the shock reaches the silo by putting some sort of barrier around the silo to absorb the ground shock (Fig. 13). This has been done commercially to protect pumping stations, for example. We have even had the Steel Industry very interested, and the Steel Workers Union has suggested putting a lot of iron pipes in the ground and have the pipes act as dampers for the shock. Obviously, this is only effective against ground burst, (we are talking about lateral motion and this has nothing to do with the overpressures coming down). The effectiveness varies considerably with the type of soils with which you are involved. Figure 14 shows another possible concept; here the crushable material would be a low porosity concrete perhaps.

Some calculations were made using an analytical model of a ground shock isolation system which was subjected to a 27 megaton field surface burst (Fig. 15). Figure 16 shows the calculations of the effectiveness of the three different types of barrier that were studied. With no barrier, the stress was considerably higher, but all of the barriers provided some reduction of the lateral ground shock. All of the foams were effective; in the case of velocity, the foam delayed or reduced the velocity considerably; and in the other case, as soon as the foam got locked up, it just translated and delayed the time. However, the total displacement appears to be relatively the same; it just takes place over a longer period of time, and therefore, the chance of protecting against the acceleration is that much better.

HARDENED MISSILE LAUNCHER R&D PROGRAM

The Air Force envisions the hard mobile basing concept as transporting a number of small missiles randomly over base roads on hardened launchers. I do not know how many missiles would be involved, possibly 500 in the ground

and 500 on a mobile launcher; possibly all of the missiles could be on mobile launchers, or all of the missiles could be in the ground. The desired maximum speed capability of the hardened transporters limits their size, and this limits their hardness. Many suspension problems on these transporters are also foreseen. These missile transporters will be located on a number of bases within the United States, and they do take up a large area. One of the main problems will be communicating with and controlling them.

Next I would like to discuss the DNA Hardened Mobile Launcher R&D Program. I might mention that BMO will build a large blast facility. They would like to be able to test a full scale hard mobile launcher. We are working on the best way to develop such a simulator. The single most important factor in defining the environment is the nature of the non-ideal air blast (Fig. 17). We are not sure what constitutes the non-ideal air blast. What is its magnitude? What are its effects? Many calculations have been made, but we have to make some better measurements if we are to simulate it. Then we must be able to simulate it if we want to use it to test the launcher. Testing the launcher against the ideal situation is not satisfactory. So, first we are trying to understand non-ideal air blasts; and second, we are looking for ways to simulate non-ideal air blasts. EMP will also be present; and for the most part, we will examine EMP environments at very low levels. We will not develop any simulators for EMP until about the 1986 timeframe. The Hard Mobile Launcher, if it is a system to be used, will not go into the field until probably the 1992 timeframe which is roughly when the small missile is supposed to be available.

There are many factors that contribute to the non-ideal air blast; these include the type of surface that you are operating over, and how that interacts with the thermal and blast loading. It also includes the effects of thermal radiation of boundary layers. We are concerned with the synergistic effects between EMP, radiation, and air blast since this system is above the ground and in an area where there is quite a bit of radiation. With the height of burst type attack on a ground/surface target, a double Mach stem area occurs. You have the thermal precursor coming out, and you have a great deal of material picked up and entrained and then becoming part of the air blast (Fig. 18). We are the most concerned with this area of the non-ideal cycle. We think we understand this situation, but it is when you get into the Mach stem that you have an ideal type situation. Still, there is a lot of material in there that we have to be able to simulate. The reason for the importance of the non-ideal air blast is that the dynamic pressure of a non-ideal air blast seems to go up compared to the ideal air blast; while at the same time, the overpressure is down in the same area (Fig. 19). The combination can be more destructive.

With the overpressure going down you cannot count on the overpressure fixing or applying a force to hold the launcher in place. It goes down at the wrong time, just about the same time that the dynamic loading occurs.

We would like to have ways of simulating these types of weapon effects (Fig. 20). The bars represent things that are relatively new or that we don't really know how to do. Dynamic Air Blast Simulators (DABS) are not new, but the size of the Dynamic Air Blast Simulators and the pressure levels that we are talking about are new. We are even considering large cavity underground tests. We recently ran a cavity underground test of about an 11 meter radius called MINI JADE. We are considering running similar tests using a nuclear driven shock tube; that is a possibility. We are presently working to simulate a small scale air blast using a modified shock tube. This will be elevated into an intermediate scale air blast simulation before we finally figure out how to do the full scale air blast simulation, which could take a shock tube that might be 3,000 feet long. From Figure 21 you will get the idea of the dimensions for a full-scale shock tube. In regard to simulating the precursor, if we cannot put in the appropriate thermal loading on the ground to cause a precursor to be formed, we might be able to simulate that by using a high sound speed gas. None of the material from the detonation should be allowed to reach the target or to interfere with the measurements. Thus, the gas driver may have to be detonatable; it might be compressed air. A lot has to be determined; how do we build a large shock tube that can be used a number of times? Building a large disposable shock tube each time is expensive; after you have done four shocks, you have spent about 25 to 30 million dollars. For that same money, you can build a permanent shock tube, or one that is partially permanent and partially self destructive.

With respect to the hard mobile basing hardening and validation, BMO has put out, or is putting out, a large number of contracts to about 5 contractors to develop new concepts. They will go through about a 10-month period of developing new ideas. Then these ideas will be narrowed down to two, and they will go into a little more full scale research program. We will review what is going on, and we will provide them with the environments, the simulation capabilities, and the instrumentation. We will also consider the effectiveness of their hardening techniques, and we will consider ways of hardening. Two of the things we are concerned with are the rigid body response and how to overcome the dynamic loading to keep the launcher on the ground. Figure 22 shows the dynamic loading of the transporter.

Basically, we have a force which lifts up and possibly a ground motion effect. If the seal leaks, the dynamic loading comes in underneath to force the transporter up, and if

the overpressure is reduced at the same time the lifting force will be considerably greater than the force on the top. Figure 23 demonstrates the combination of the dynamic loading and the overpressure change which results in the reduced stability of the system. A great deal depends on the shape. On DIRECT COURSE, we have four contractors testing different shapes, different types of seals, and different dimensions; but it is not clear that you can accurately scale those features so the tests are not considered definitive. They are just the first attempt to do something on this height of burst test, and we expect to do more height of burst tests in the future. Figure 24 shows some of the seal candidates which are being considered and also the idea of putting something underneath the launcher which will anchor it to overcome the sliding problems.

Figure 25 is a program schedule that shows the DNA program feeding into the BMO program with our program concentrating on the first three years. We are providing this information to try to get a Full Scale Engineering Development input, but our program will be continuing at a reduced level from what it will be during the first three years.

SUMMARY

We are considering both hardened silos and mobile launchers. We are trying to incorporate existing technology but we are looking for new concepts. Major problems are to reduce the uncertainties in the environments and to develop the necessary simulation and instrumentation; the bottom line is to prove that what we have done is correct. All of this is concentrated in the first two years of the program. This is the third year. If you want to talk to anyone about these programs, the technical director for the Hard Silo is Dr. Kent Goering, and the technical director for the Hard Mobile Basing is Dr. Paul Rohr. Both of those people are in the Shock Physics Directorate of the Defense Nuclear Agency. That covers the DNA Program with a little of the BMO Program thrown in.

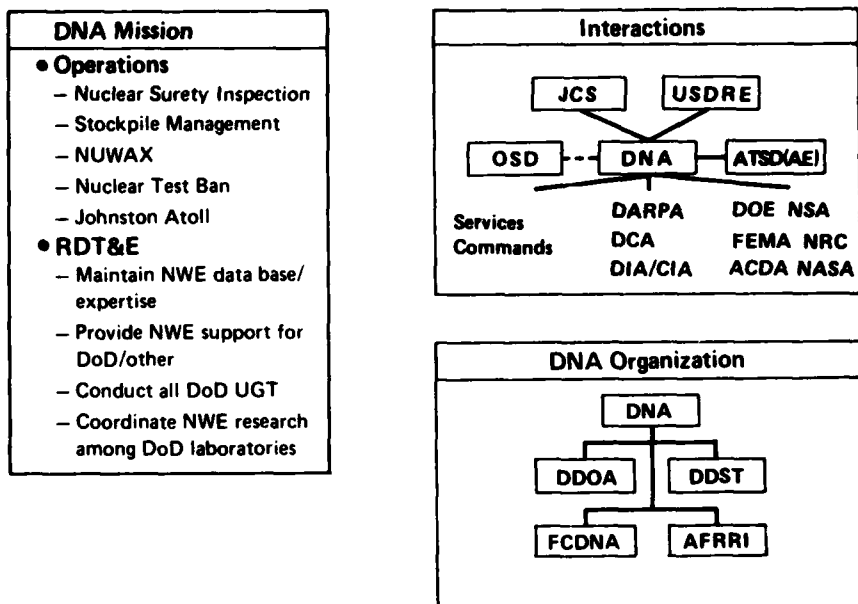


Fig. 1 — Defense Nuclear Agency Mission, Interactions and Organization

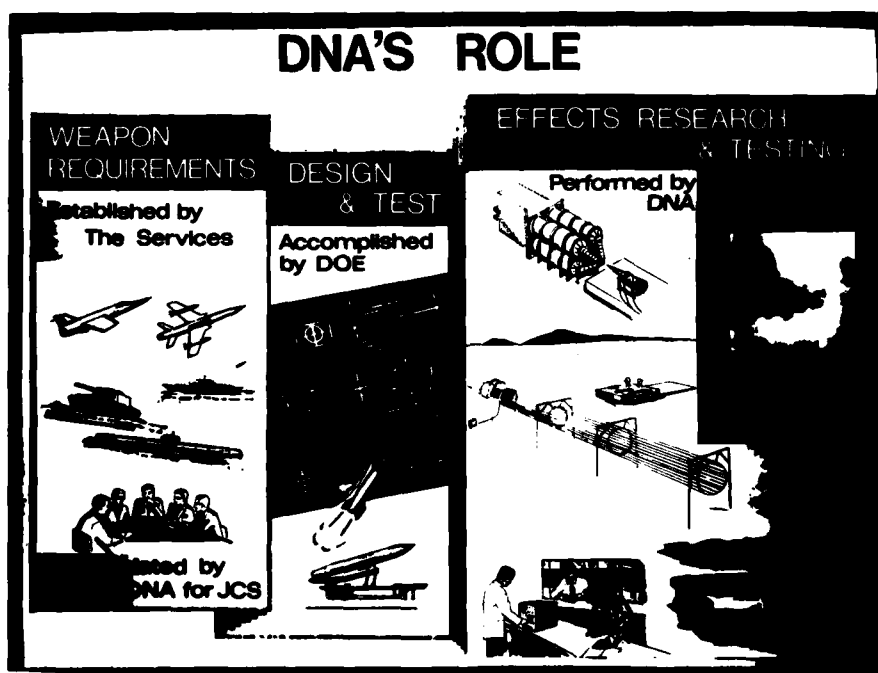


Fig. 2 — The Role of the Defense Nuclear Agency in Weapons Development

DDST Organization

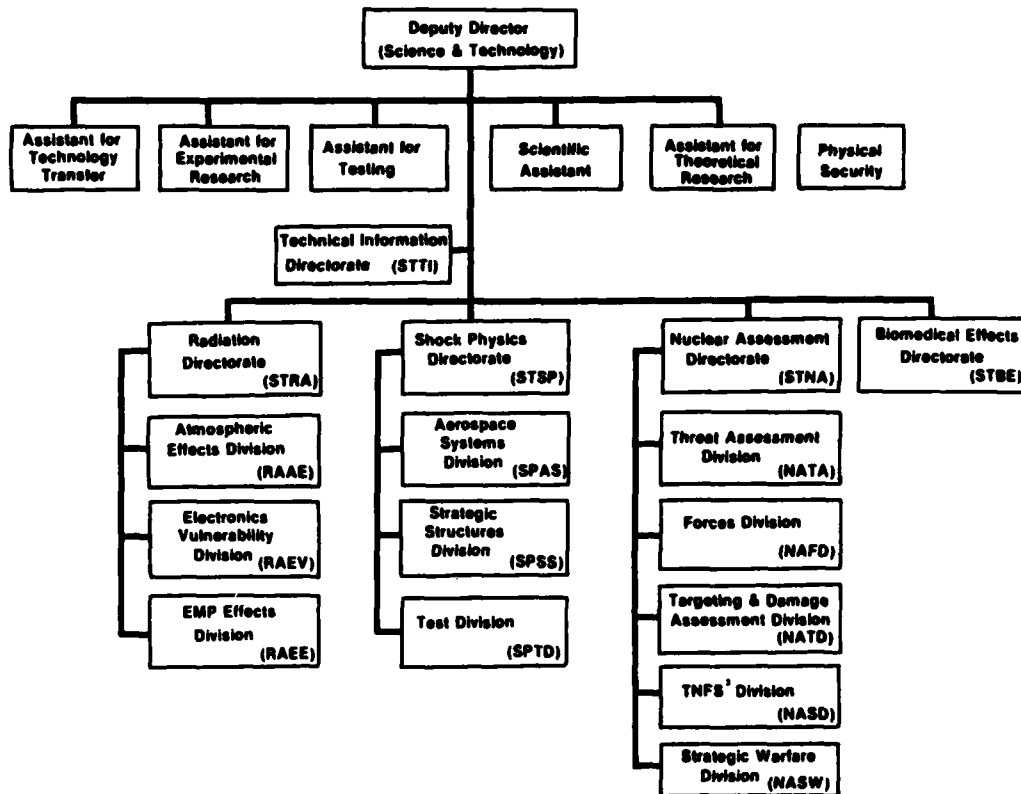


Fig. 3 — Organization Chart for the Deputy Director of Science and Technology

Shock Physics Directorate Mission

- DNA focal point for
 - Thermal/Mechanical effects of nuclear weapons
 - Directed energy
 - Test instrumentation development
- Manage integrated research and test program to satisfy DoD requirements for shock effects information on mobile and fixed weapons systems, structures, and components
- Management of DoD underground nuclear weapons effects test program (UGT)
- Planning and management of high explosive simulation (blast and thermal) of nuclear weapons effects, to include simulation development

Fig. 4 — The Mission of the Shock Physics Directorate

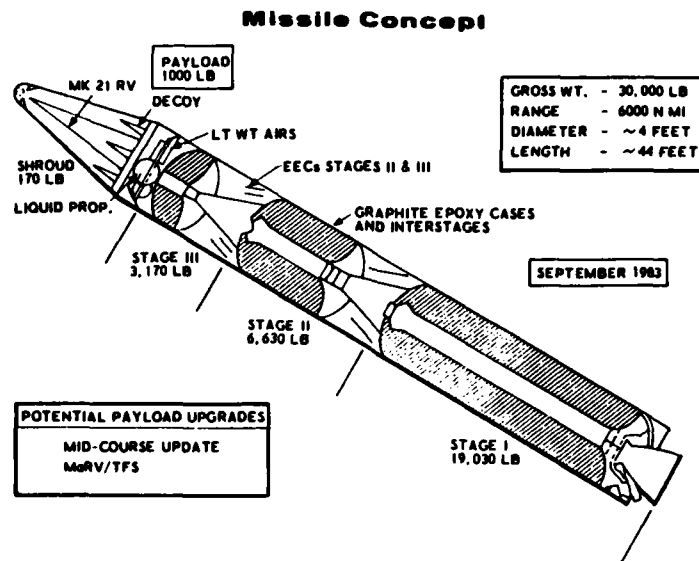


Fig. 5 — Generic Small Missile

Objectives

1. Resolve uncertainties about silo hardness to establish maximum credible hardening levels.
2. Examine workable concepts over a range of hardness levels.
3. Test most promising concepts within three years.

Technical

- *1. Evaluate - advanced silo design technology/concepts.
2. Develop NWE simulators/test beds.
- *3. Examine attainable hardness levels, which are cost-effective.
- *4. Test models of most promising concept(s).

*Areas in common with BMO

Scope

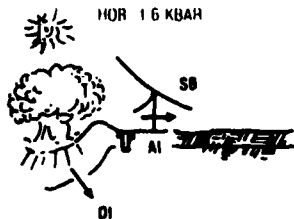
- *1. Structural/systems technology.
 2. Testing technology (simulation and instrumentation)
 3. Environment definition.
 - *4. Field testing and concepts screening.
- *Areas in common with BMO

Level of Effort

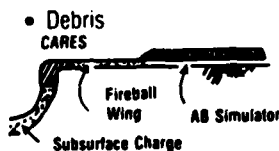
Fig. 6 — Advanced Silo Hardening Program — Summary

Environments

Fratricide Environs
Silo Basing Hardness



- Material properties
 - In-situ
 - Dynamic
- Airblast
- Cratering & Gnd Shock
 - Analysis
 - Nuclear simulation
- Debris

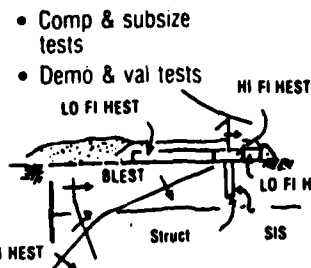


- Radiation & EMP

Simulation & Instrumentation

Silo Basing Hardness

- Simulation
- CARES

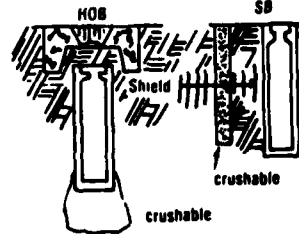


- Demo & val tests
- Instrumentation
 - Surface overpressure diagnostics
 - FF stress & motion
 - SMI interface stress
 - Structural motion strain & deformation

Silo Hardening

Advanced Silo Hardening

- Advanced concepts
- Concept screening



- Concept testing (small scale)



- Add-on tests (50 in 0)



Fig. 7 — Advanced Silo Hardening Program — Instrumentation and Simulation

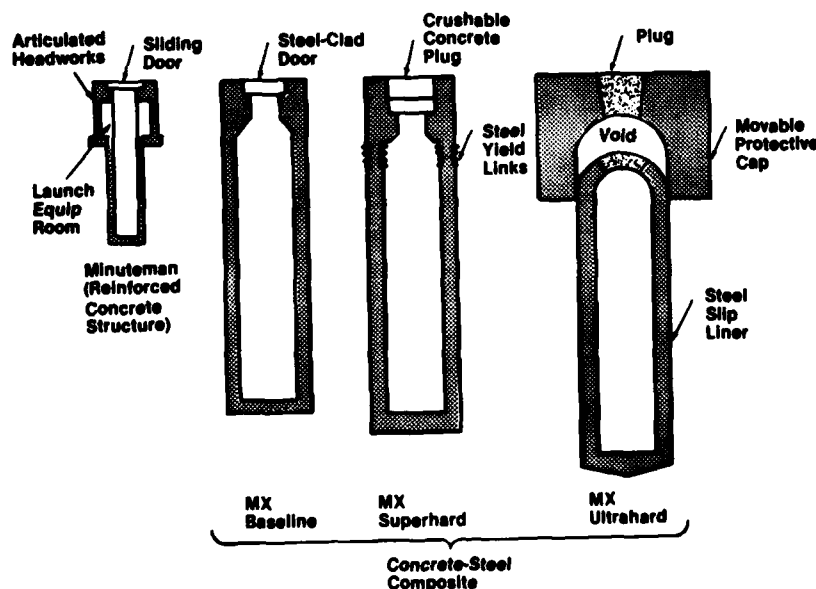


Fig. 8 — Historical Trends in Silo Hardening

ULTRAHARD - DECOUPLED CONCEPT

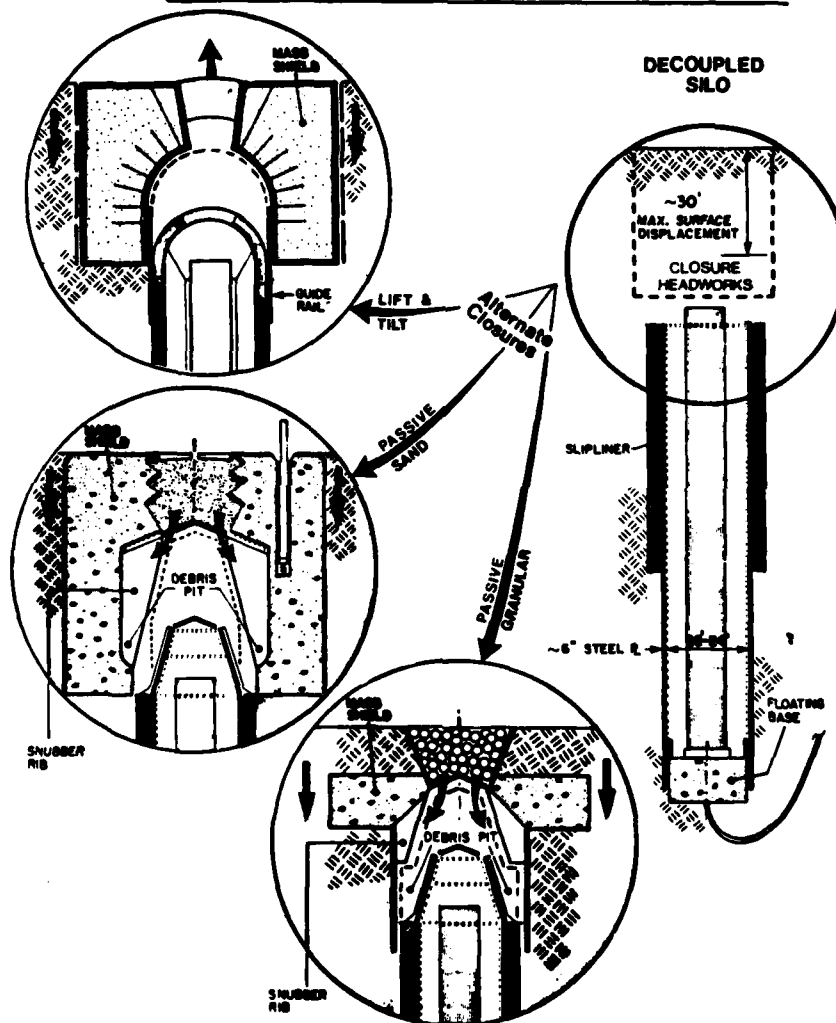


Fig. 9 — Concepts for Ultra-Hard Silos

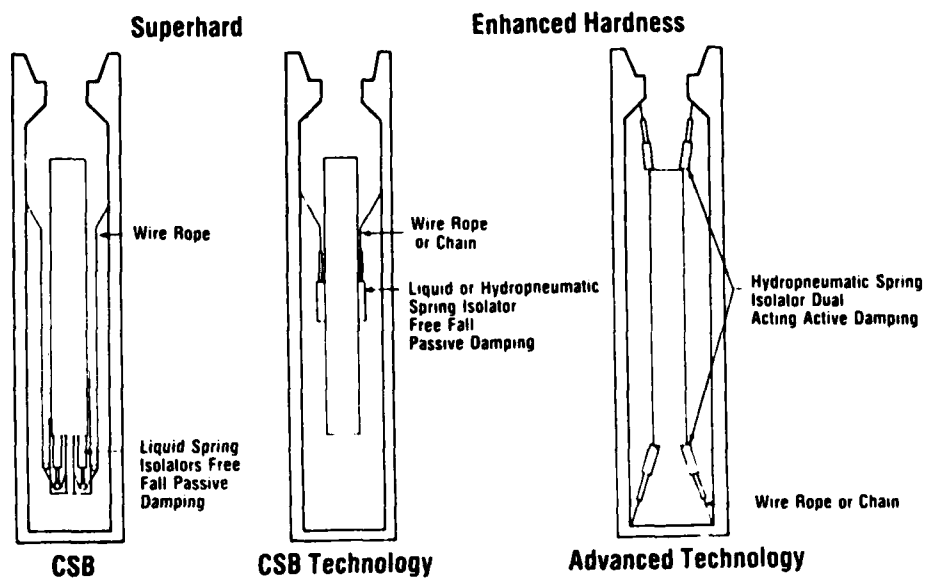


Fig. 10 — Vertical Shock Isolation Concepts

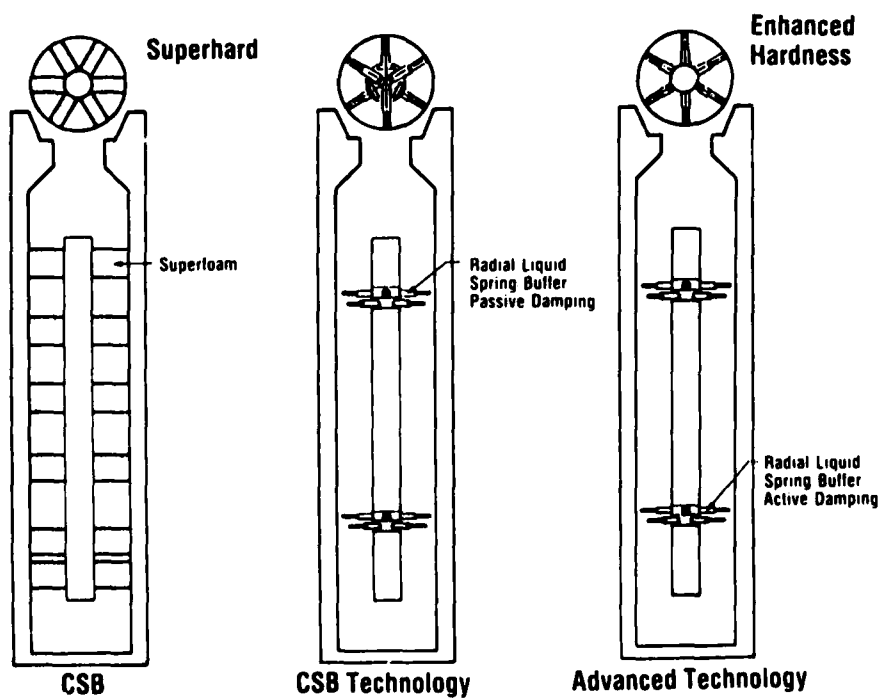
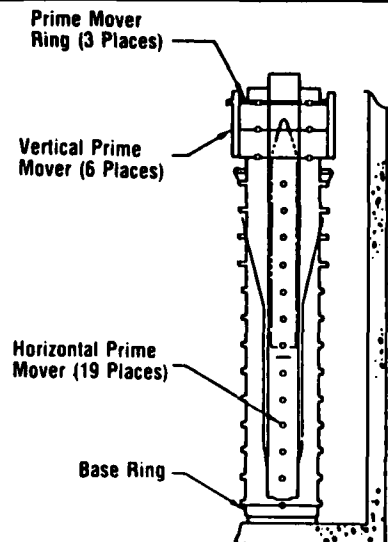


Fig. 11 — Lateral Shock Isolation Concepts

System Studies and Analyses SIS Shock Test Facility



- Conduct system studies to evolve concept for large scale shock isolation system test facility
- Provide concept point-of-departure, preliminary requirements document and costs

Ground Shock Test Machine

Fig. 12 — Large Scale Shock Isolation System Shock Test Facility

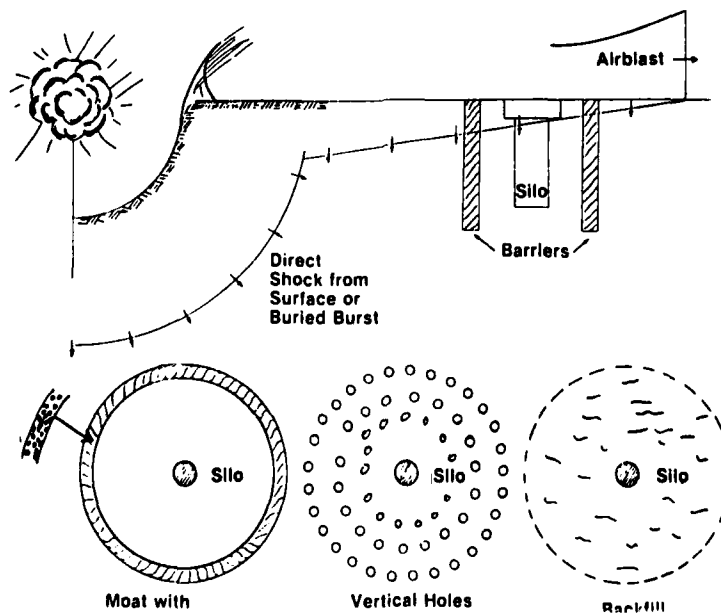
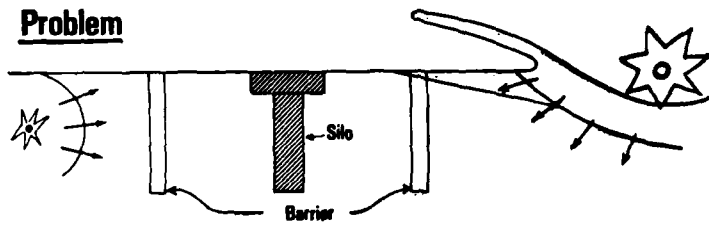


Fig. 13 — Barrier Concepts for Shock Isolation of Silos

Problem



Program

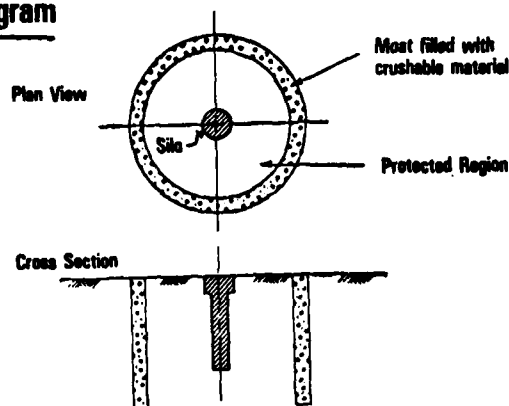


Fig. 14 — Crushable Material Barrier for Shock Isolation of Silos

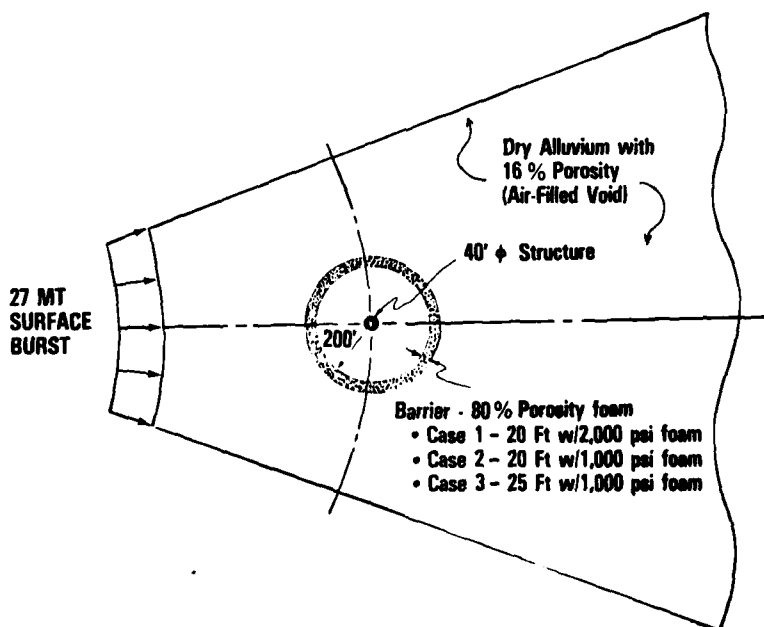


Fig. 15 — Analytical Model for Ground Shock Isolation Study

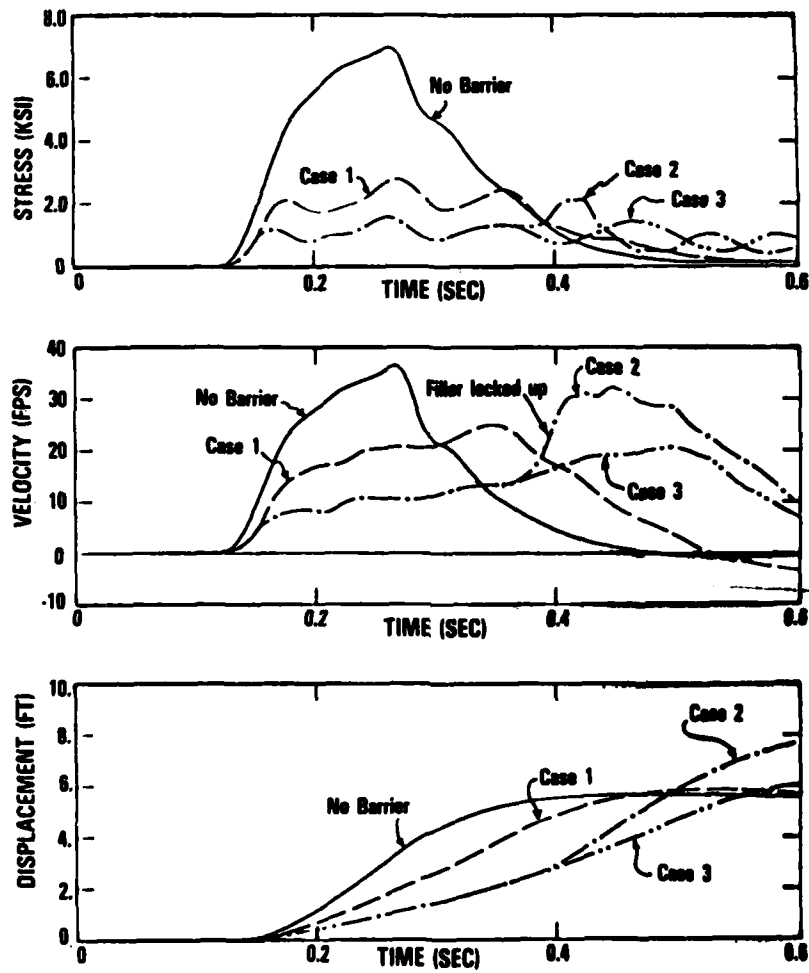
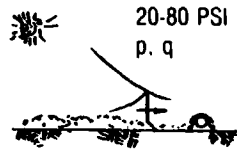


Fig. 16 — Effectiveness of Three Different Types of Barriers in Attenuating Ground Shock Induced Motions and Stresses

Environments

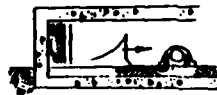


20-80 PSI
p, q

- Non ideal surfaces
 - Dust
 - Vegetation
- Thermal radiation
 - Single burst
 - Multi burst (dust shielding)
- Non ideal airblast
 - Dust boundary
 - Thermal precursor
- Ground shock, EMP, & radiation

Simulation & Instrumentation

• Simulation



- Concept evaluation
 - MT TNL (reinforced)
 - Reusable new fac
- Concept development
 - Source (gas, H.E.)
 - Facility (sml, lrg)

• Instrumentation

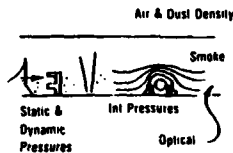


Fig. 17 — Hard Mobile Basing Environments

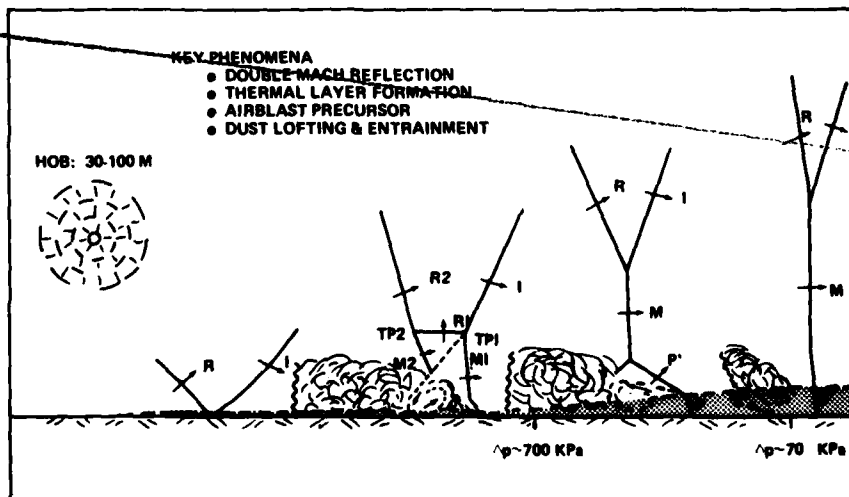


Fig. 18 — Height of Burst Phenomenology

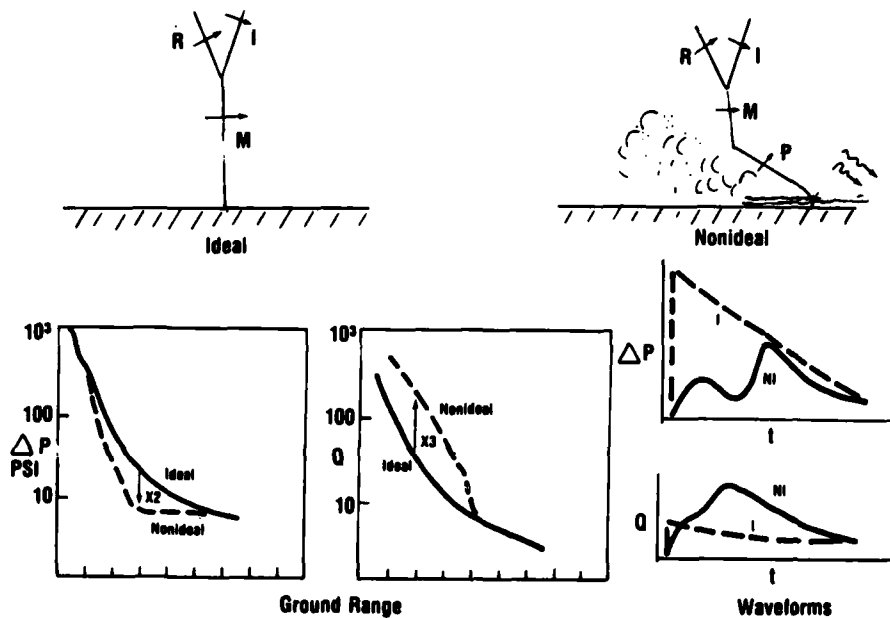


Fig. 19 — Nonideal Airblast Environments

Component Validation and Concept Verification

Data Requirements and Test Options

	Laboratory Test						Field Test				
	Primary	Shock Vibr'n	Shock Tube	Thermal Rad'n	X-Ray Sim	"Old" Nuclear	HE point Source	HE Area Source	SLED	UGT	DABS
Airblast precursor											
Aerodyn. forces											
Cratering/ejecta/dust											
Missile											
Canister/SIS											
Transporters											
EMP/Gamma											

Hardness Validation Guidelines:

- Use experiment to validate theory and provide "point validation" of design
- Use analysis to extend "point validation" over design envelope
- Exhaust lab test options before proceeding to field test
- Test simulation fidelity essential if theory questionable

Fig. 20 — Weapons Effects Test Technology Development and Simulator Requirements

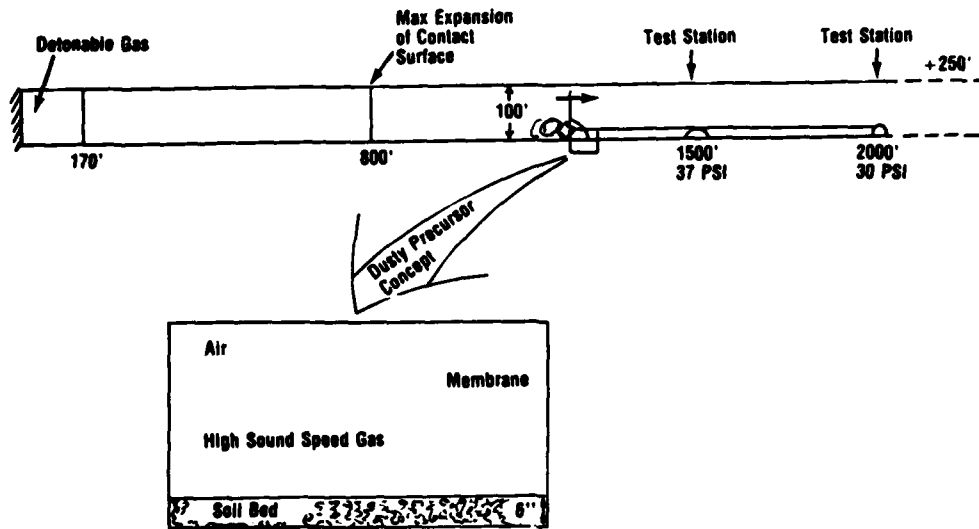


Fig. 21 - Large Scale Reuseable Airblast Test Facility Concept

Transporter Overturn / Sliding

1. "Nuclear Gust" aerodynamic force coefficients?
2. Ground roll arrival time & acceleration: seal degradation?
3. Shear reaction limited by terrain shear resistance or friction
4. Lift force developed if base seal leaks

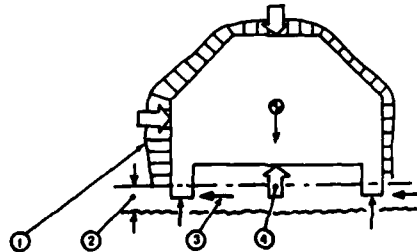


Fig. 22 - Dynamic Loading on the Transporter

Precursor Implication: HML Pressure Loads

Plumbbob - Priscilla Airblast Loads Data
Translational & normal pressure histories
37 KT/700 ft HOB/6×6 ft block at 2000 ft

- Reduced Hold-Down Force
- Increased Drag Force
- Reduced Stability

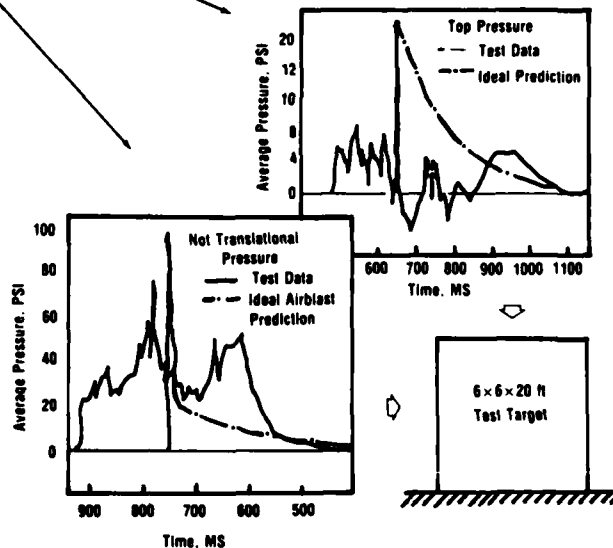


Fig. 23 — Airblast Pressure Loads on Hard Mobile Launcher

Seal Candidates

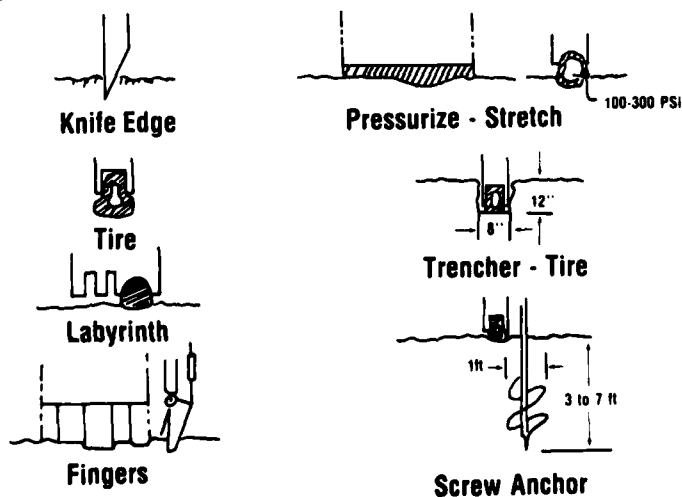


Fig. 24 — Candidate Designs for Seals for Hard Mobile Launcher

Program Schedule

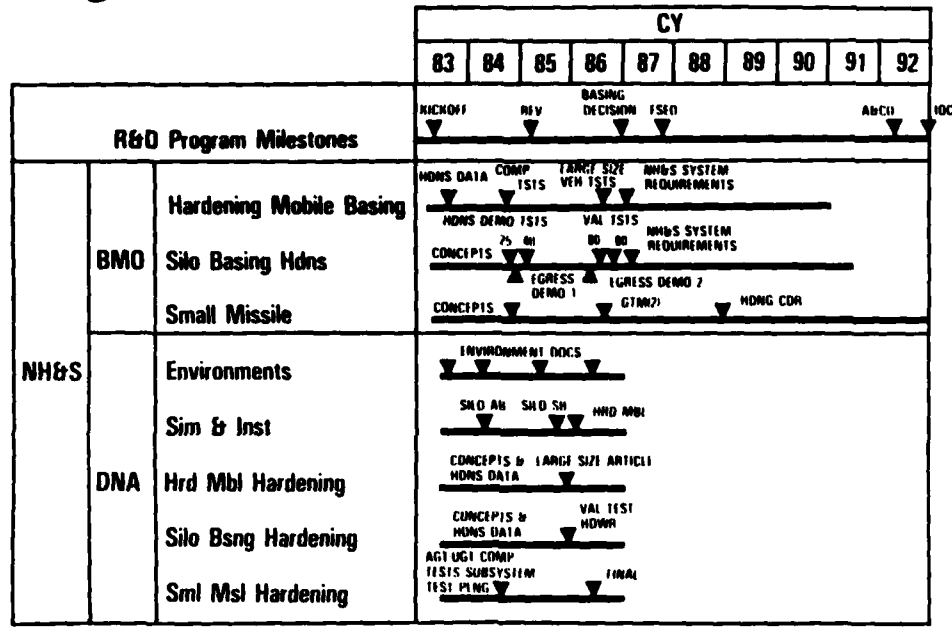


Fig. 25 — DNA-BMO R&D Program Milestones

SOME DYNAMICAL ASPECTS OF ARMY MISSILE SYSTEMS

James J. Richardson
U.S. Army Missile Command
Redstone Arsenal, AL

INTRODUCTION

An indication of the range of dynamic phenomena involved in Army missile systems can be found in considering the PERSHING II and the VIPER. The PII, the latest member of our inventory, is a 16,000 pound missile, while the VIPER weighs in at around three pounds. Obviously, the structural design philosophy and dynamic environments are at great variance between the two systems. As we "fill the gap" between them with other missile systems, the great diversity of size, employment, and flight parameters yield an accordingly diverse set of dynamics problems. It is the aim of this paper to provide some indication of the range of these problems. A few general remarks will be followed by examples taken from recent experiences.

It is difficult to categorize dynamic loading on missile systems since there are so many environments and conditions to consider. Perhaps one approach is to look at phases of system life. During transportation, the missile and accompanying ground support equipment are subjected to various modes of transport each of which produces its own shock and vibration spectrum. During this phase, the missile is generally carried in a container which furnishes some cushioning. The launcher/missile combination, during the employment phase, may be carried on wheeled, tracked, or winged vehicle, or for that matter, may be slung on a soldier's back. This obviously represents a wide range of imposed forces. The large temperature spectrum, generally -25°F to $+125^{\circ}\text{F}$, to which military equipment is exposed (particularly during this phase) frequently affects material and mechanism response to dynamic loading.

Launch loads are sometimes quite severe. Ignition shock on HAWK, for example, has been measured as high as 1,500 g's at a pulse duration of .15 milliseconds. Only recently have we adopted a shock spectrum representation rather than half sine to reflect this environment in the system specifications. Obviously, it is not easy to analytically predict the effect of such a force on a conceptual rocket. Detent, blast impingement, spin-up, and friction forces occurring during this time must also be accounted for.

Besides thrust/inertial loading, flight conditions include vibration from the propellant, particularly if there is a danger of uneven burning. Fortunately, the new propellants produce a much smoother environment than some of those in the past, and unlike the liquid propellant motors of NASA, the Army's solid motors avoid slosh and POGO problems. Flight aerodynamics produce both constant and fluctuating forces which can lead to divergence and flutter. Structural, and sometimes flight path destabilizing loads, can result from roll acting on mass offsets along the rocket's length or from the thrust acting on a CG misalignment with respect to the centroidal longitudinal axis.

The problems which occur during various phases of the acquisition cycle and fielded life of a missile system are certainly related. The conceptual and design difficulties are generally of a more basic (but no less complex) nature than those of later phases. They typically involve trade-offs between performance and design requirements. Effort is directed toward developing ideas to overcome overall performance barriers. Less time is dedicated to cost and to system qualification to military specifications subordinate to the requirement documents. During these early phases, the dynamicist needs a good deal of experience to predict environmental conditions and to find ways to withstand those conditions. Problems solved or avoided at this time save immense amounts of redesign work and testing farther on in the cycle. Performance problems surfacing during development testing usually lead to test-redesign-retest cycles. Care must be taken not to solve one dynamical or kinematical problem at the expense of creating another.

Once fielding occurs, the true operating conditions will be imposed on the system and the abilities of the dynamicists to predict environmental loads, design to them, and synthesize them in laboratory tests will be revealed. Merit lie the grounds for a good many struggles. Predictions of levels and durations of shock and vibrations must be made early in the design stage (usually during advanced development) when little may be known of the system characteristics or the environment in which it must operate. The tendency is to

become very conservative in formulating design and test specifications, a tendency which is frequently resisted by program managers and designers alike. If the resulting early shock and vibration specifications are too high, the missile system is overdesigned, but since contractual design requirements may be relaxed, there is little controversy. If, on the other hand, the specifications are found to be too low, the program manager must raise design requirements, often just as his contractor is preparing for production. As can be imagined, the ramifications are severe. Additionally, as will be demonstrated later, production changes (both directed and inadvertent) keep problem solving engineers in business throughout the life of the weapon.

So, from the very beginning, the tools of the dynamicist: loads definitions, stress and response modeling and analysis, failure criteria development and both specification and diagnostic testing, must be sharp. This is becoming even more true as demands on our designs grow. Today, in general, we are looking to faster missiles, weighing and costing less, with higher performance and more mission versatility. In addition, new materials are being employed in order to meet these demands. Composite materials, for instance, solve many design problems at low cost and weight. Unfortunately, however, the structural dynamicist is frequently unsure of their strengths, their response to dynamic excitation, and the deterioration of their mechanical properties after exposure to moisture, ultraviolet energy, and long term storage. Inconsistent properties of these materials from manufacturer-to-manufacturer or even lot-to-lot also haunts the structural engineer.

By no means a comprehensive categorization of dynamic environment and problems, the above does outline a few challenges which we face in our field. I would like now to take a page from our bretheren in the MBA world and present some "case histories" to illustrate some of the generic problems described above.

As we turn our attention to specific examples, it is perhaps appropriate to begin with the VIPER (Figure 1), the program recently cancelled by the Army. VIPER's history began and ended with a lack of appreciation for its complexity. Complexity not in electronics or systems, but in the physics involved. Just imagine an unguided, three-pound rocket which attains a velocity of nearly the speed of sound after just five feet of travel, flies in a flat trajectory for a half kilometer and destroys nearly any mobile armor made. In order to do this, the rocket must reach an acceleration of 8,000 g's from a thrust level approximately equal to that of the PERSHING II. At the same time, it must transfer a set of forward-folding fins from a large to a smaller tube -- all within inches of a soldier's ear. The accomplishment of such a feat (and no other

system has matched it) is even more amazing when one considers the other safety and performance demands imposed. Now, this is not to say that mistakes were not made on VIPER, nor even that the program should not be cancelled. It is important, however, that as engineers, particularly in the "mechanics" end of engineering, we recognize the complexities imposed by physics and not only those resulting from interactions among numerous electronics and mechanical components. This case history involves a design problem resulting from the need to transfer the fins from the large diameter rear tube to the smaller diameter forward tube. The launch tube telescopes closed in order to reduce carry length, but opens prior to firing to provide the required guidance length. The fins, shown in Figure 2, must therefore remain fully closed as they move from the larger into the smaller tube. The original design is shown in Figure 3. The hold down device was simply a ring of arms made from foamed plastic. An arm was inserted between each fin and the inside surface of the outer (rear) launch tube. When the face of the inner (front) tube was impacted, the foam crushed, leaving the fin to continue alone down the smaller tube. Since the velocity at that point was 500 to 600 feet per second, the hold down device offered little resistance. Unfortunately, the foamed material was subject to environmental deterioration. The catastrophic results of its failure to hold the fin down prior to transfer is obvious if one imagines the fin tip encountering the face of the inner launch tube.

One proposed solution was the adoption of aft fins similar to those in Figure 4. In fact, g loads affiliated with launch would insure that the fin would remain down without a locking device, although one is shown in the figure. This solution was discarded due to loss of static margin (the fin pivot being further forward). So, back to the forward folding fin. Several hold down ideas were generated. These are shown in Figure 5. The GEM clip was adopted. As a fail safe, it was demanded that the fins transfer without hazard even if the clips were missing or broken. This was accomplished by flaring the inside diameter of the inner tube into a ramp and cutting the fin top off as illustrated in Figure 6.

The HELLFIRE missile, shown in Figure 7, is airborne launched. In its captive flight phase (when it is being carried on a helicopter) it must operate under the shock and vibrational environment produced by the aircraft. This environment, depicted in Figure 8, is in the form of a complex harmonic, the sum of sinusoidal functions at the rotor blade pass frequencies, with a floor level random spectrum. Until fairly recently, it was impractical to impose such spectrum in the laboratory. The advent of Fast Fourier Transform controllers and vibration analyzer, however, have made this possible. Further, we

can now control response rather than input with a fair degree of accuracy. The combination of these advancements allow a much more accurate synthesis of reality. Even so, we must evaluate the cumulative damage caused by the administration of this cyclic load. Although techniques have been developed to accomplish this, none appear to be sufficiently accurate. More work is needed in this area.

The ANSSR, an acronym for Aerodynamically Neutral Spin Stabilized Rocket, was proposed by Emerson Electric Corporation as a replacement for the venerable 2.75 inch rocket, now known as HYDRA-70. Its chief improvement was accuracy, which is attained by gyroscopically stabilizing the rocket during flight. This meant spinning the rocket at 12,000 rpm. As can be imagined, a 100-pound rocket spinning at this rate on your right wing at 1,000 feet can be unnerving. The first launch was made on the ground. The rocket spun-up on the launcher and broke up 50 meters downrange. The structural failure occurred at the pedestal joint, where the ogive is threaded into the motor case. Two facts were obvious: pieces of the joint picked up on the range indicated that the threads were stripped by the extremely large bending moment and there were no external forces present in the system which could have produced such a failure. The later of these facts led us to suspect a resonance condition; the former directed us to investigate roll/pitch interaction. Roll/pitch interaction occurs when spin induces transverse forces on the rocket creating bending moments along the longitudinal axis. These forces are caused by mass offsets along the rocket's length. Their magnitude is $F = m_e w^2$ where:

m_e = effective mass offset

e = distance from m_e from longitudinal centroidal, axis of the rocket

w = spin rate.

Of course, the frequency of this forcing function is that of the spin rate and, as that frequency approaches a natural bending frequency of the rocket, resonance occurs. We delivered our verdict, suggested a modal scan to determine the rocket's natural frequency in bending, and were ignored by the project management. Their solution was to apply "locktite" to the joint and retest. This time the rocket covered 75 meters before breaking apart. We were allowed to conduct modal scan! The rocket was suspended from its first bending mode nodal points using "bungee" cords and vibrated in the horizontal plane, at its center of gravity. The results: the first bending mode frequency was 12,000 cycles per minute, exactly at the spin rate. The spin rate was lowered to 10,500 rpm and no further structural problems were encountered.

This incident triggered a general investigation of the effects of structural response on launch and flight dynamics. Figure 9 illustrates how a rocket may be subjected to

large dynamic forces at launch if its spin rate excites a second mode response on the launch and a first mode response in the free flight phase -- and it could not happen at a worse time. Detailed models now allow us to predict the effects of structural response throughout the launch and flight phases on the accuracy of the rocket.

The HELLFIRE system experienced exactly the reverse of the ANSSR roll/pitch interaction. A system requirement demands that HELLFIRE's seeker gyro be tested under captive flight vibration. Figure 10 shows the gyro rotor in black. The spin-up history, shown in Figure 11, must insure that the rotor reach its operating spin rate of approximately 65 RPS within 30 seconds. The vibration test chosen to represent the helicopter environment was a sinusoidal sweep from 5 to 500 Hz. The gyros always began their spin up at the start of this sweep and well within 30 seconds (while the sweep test had progressed at a relatively low g level to 6 or 7 Hz) the operational spin rate was reached (shown by gyro spin-up history curve number 1 in Figure 12). The laboratory began testing three seekers at a time, however, allowing each to reach 65 RPS before starting the next one. This resulted in the third seeker starting its spin up while the sweep test was in the 10-20 Hz region, the high g portion of the spectrum shown in Figure 13. The first two seekers spun through 65 RPS with no difficulty, but the third would not progress beyond the frequency of the traverse vibration. In fact, the spin rate increased with the sweep -- an interesting case of pitch/roll lock-on, when the forces generated by the transverse vibration dominated the torque supplied by the spin motor. Incidentally, the solution to this "problem" was to decide that it was not a problem. Since no helicopter imposes a forcing function with a single frequency, we suggested a complex harmonic with the stronger components of those shown in Figure 8. The resulting distribution of energy among more than one frequency was more realistic and allowed normal spin-up.

The final case history involves what must be termed our most successful missile system. The HAWK has been fielded for 25 years. Nearly every allied country has employed it, and yet it is far from trouble free. Indeed, some of our most interesting mechanical problems have been encountered on this "stable" and relatively reliable system -- a great consolation for engineers concerned about losing their jobs after their product's development cycle is over. The HAWK launcher is zero-length. Since there is no guidance rail or tube, the missile must be held in place until the thrust is sufficiently high to insure flight stability. This is accomplished by a "forward sector" which holds the missile at point A in Figure 14 until 2,800 pounds of thrust rotates it out of the way and thereby releases the missile. Unfortunately, a few launchers experienced sector rotation during tracking missions or

launcher azimuth and elevation exercises, manifesting in the dropping of a live 1,400 pound missile on the ground. Embarrassing! Obviously, the azimuth and elevation movements of the launch somehow produced 2 g's in the longitudinal direction, thus overcoming the sector. In order to determine the parameters affecting this force, we instrumented a launcher in the laboratory with a single accelerometer located at point A in order to measure F_g under various conditions. We found, for example, that the highest g levels at A occurred when driving the launcher arms from 50 mils elevation to 0 mils (Figure 15). Further, the two missiles mounted on the side arms were cushioned due to the flexibility of those arms, while the more rigid center launch arm saw much higher g's. Figure 16 shows that the azimuth was an important parameter as well. Lowering the launcher from 50 mils to zero in elevation produced the highest g levels at 800 mils azimuth, due to a stiff outrigger at that location on the launcher bed. Hydraulic shock absorbers (or "snubbers") arrest the launcher arms at +5 mils elevation in order to prevent bottoming out. If hydraulic fluid was low, g levels increased. Even under the worse conditions, however, we were able to induce no more than 1.5 g's. Armed with this understanding of the dynamic response characteristics of the launcher, we took our measuring system and procedures to the field. We strapped dummy missiles to several launchers which had previously dropped missiles. G levels above 2 were measured on these launchers signifying a unique problem. Engineers familiar with the hydraulic systems identified the culprit, -- a valve designed to control the hydraulic flow in the lines powering elevation motion. The real culprit, however, was good intention in the form of a depot worker who had been "saving the government money" by rebuilding these valves against the advice of the manufacturer. Figure 17 portrays the effects of replacing ineffective snubbers and out-of-tolerance valves on the g levels. Since purging the Army's inventory of bad valves, there have been no dropped missiles due to high launcher accelerations.

If these case histories carry any lesson, it is that dynamicists must frequently look beyond their primary field of interest if they are to discover the cause of problems. The need for experience, I believe, is also obvious. Experience coupled with a constant pressure to continue upgrading our collective analytical and experimental tools are musts in being able to solve problems which are frequently very complex in nature.



Fig. 1 — Viper: A Feat in Physics

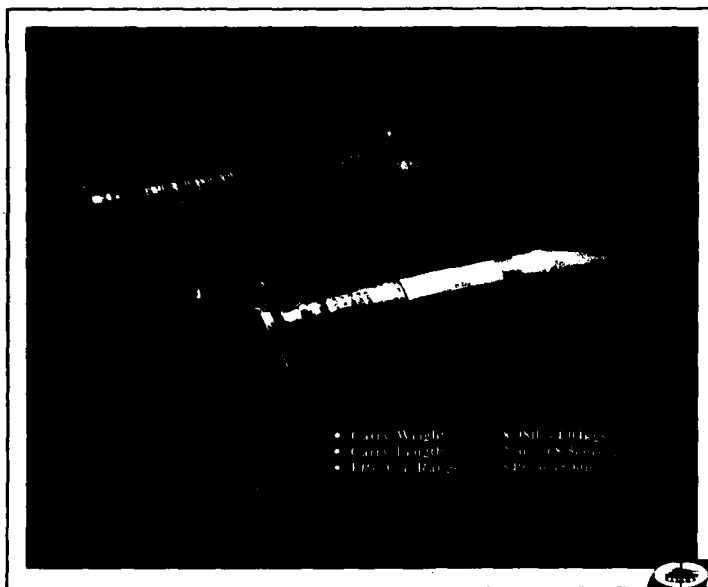


Fig. 2 — Viper Fin Arrangement

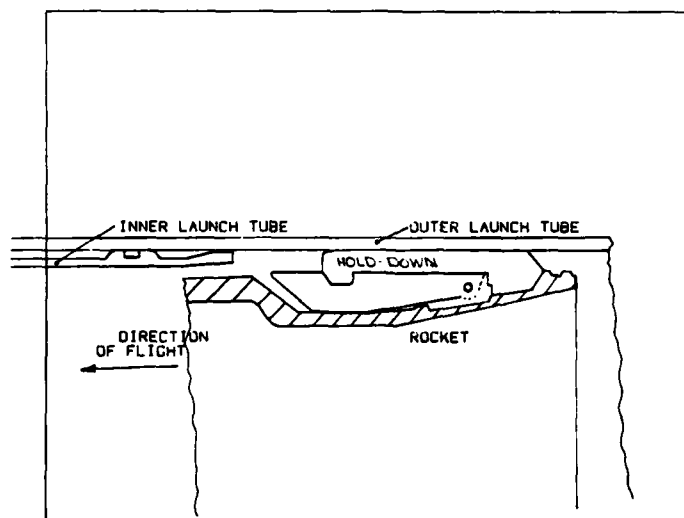


Fig. 3 — Fin Transfer Problem

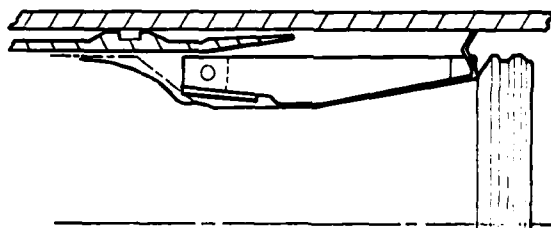


Fig. 4 — AFT Folding Fin

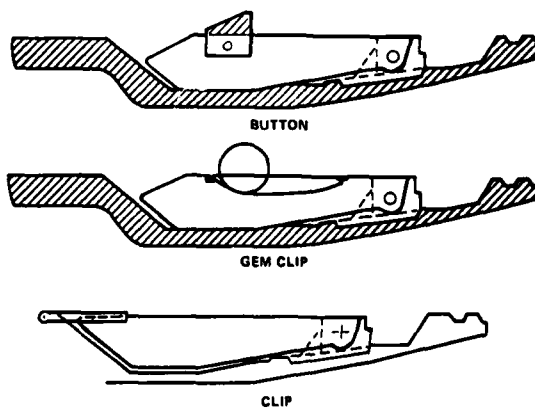


Fig. 5 — Simple Concepts



Fig. 6 - Unassisted Transfer



Fig. 7 - HELLFIRE

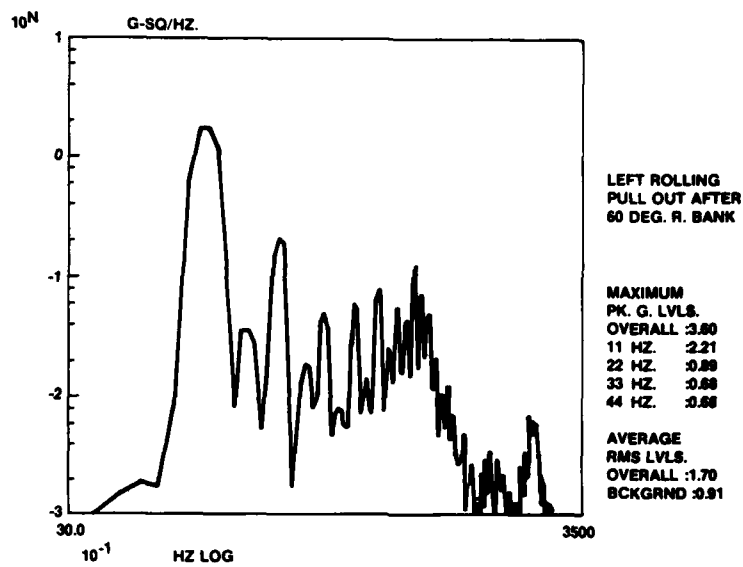


Fig. 8 — Power Spectral Density of AH-1

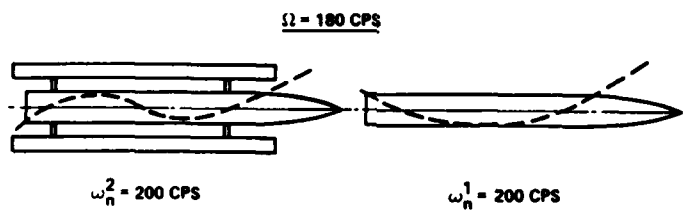


Fig. 9 — Transition from Launcher Constraint to Free Flight

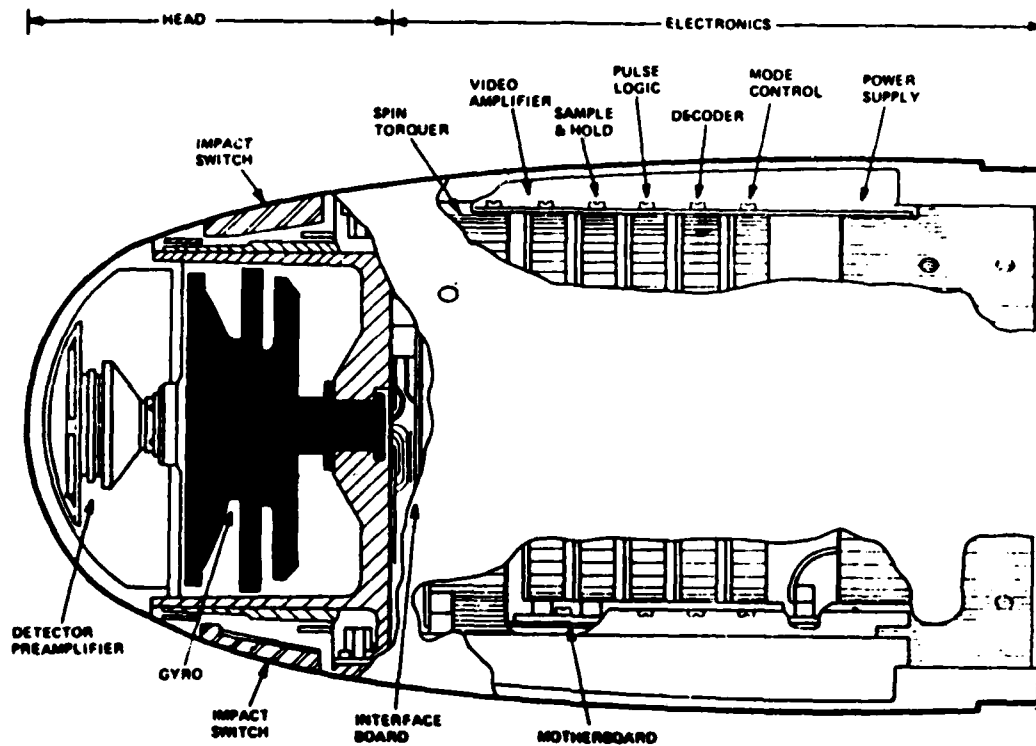
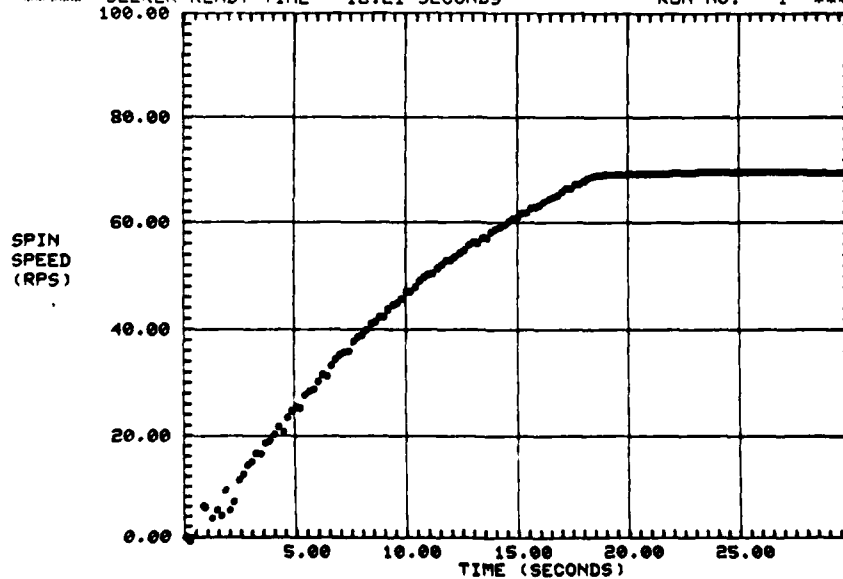


Fig. 10 - HELLFIRE Laser Seeker

```

***** LASER SEEKER          SPIN SPEED TEST          *****
***** U.S.ARMED MISSILE COMMAND (ALSPES DATA) (205-876-8521) *****
***** SSN:HX-0014          CODE: 488          DATE:20 SEPT, 1983 *****
***** SEEKER READY TIME 18.21 SECONDS          RUN NO. 1 *****
  
```



NOTE: RUN#1--CHARACTERIZATION TEST
SAMPLE TIME 30 SECONDS

Fig. 11 - Characterization Test

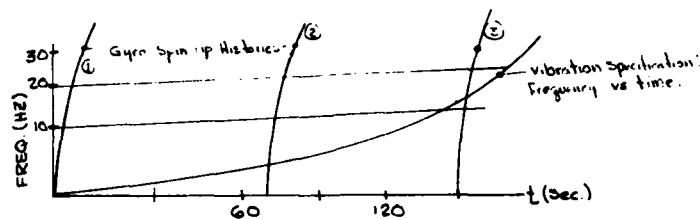


Fig. 12 - Spinup Histories vs Vibration Frequencies

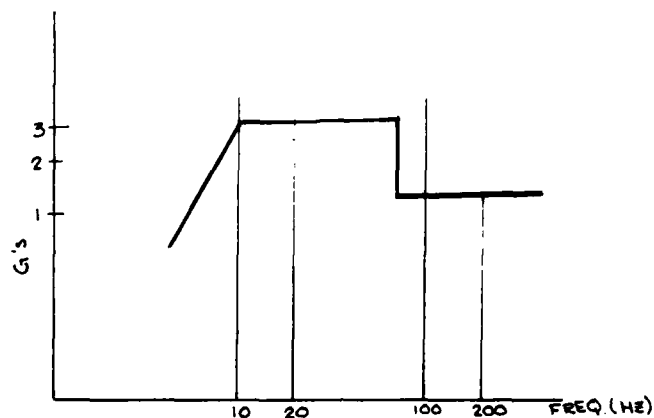


Fig. 13 - Vibration Specification

$$W = 1400^*$$

$$\theta \approx 9^\circ$$

$$\Sigma F_x = W \sin \theta - F_s = W q_x$$

$$F_s = W [\sin \theta - q_x]$$

$$F_{s_{cr}} = 2800^* \Rightarrow q_{x_{cr}} = 1.84g$$

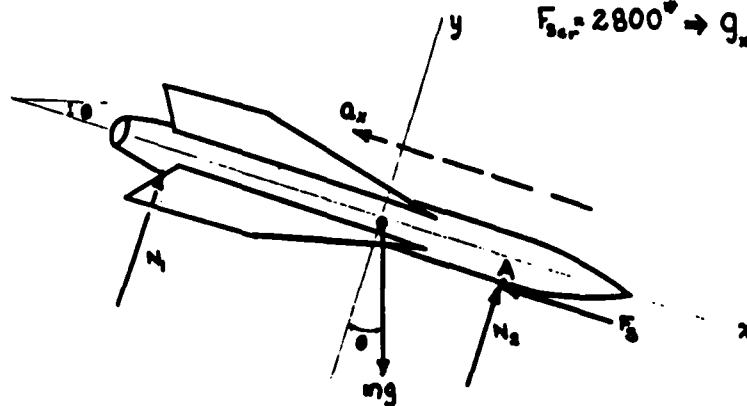


Fig. 14 - Hawk/Launcher Free-Body

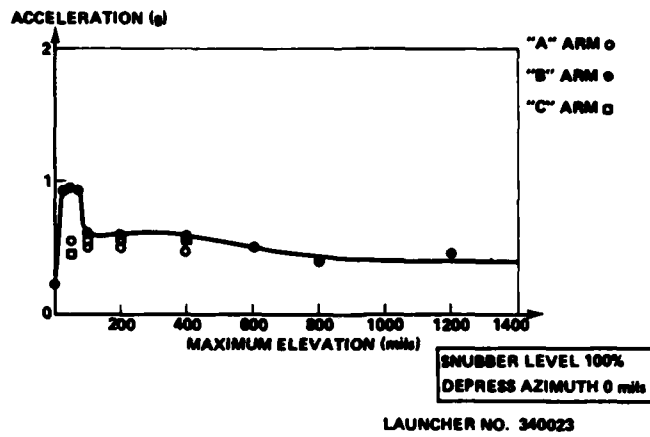


Fig. 15 - Maximum Elevation vs Acceleration

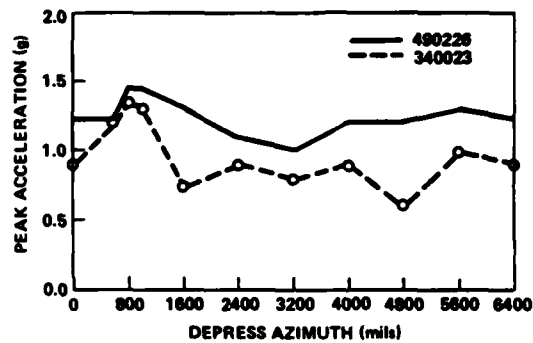


Fig. 16 - Peak G's at Various Azimuths for 50 mils Starting (Maximum) Elevation

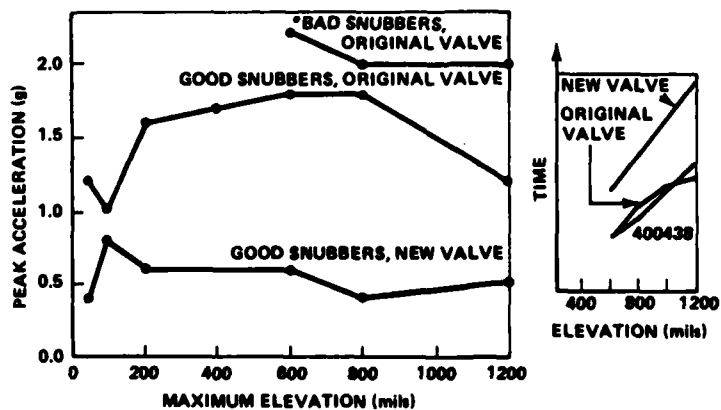


Fig. 17 - Effects of Snubbers and Flow Regulator Valve

AIR FORCE SPACE TECHNOLOGY CENTER SPACE TECHNOLOGY - EMPHASIS 84

Colonel Frank J. Redd
Air Force Space Technology Center
Kirtland AFB, NM

INTRODUCTION

The Air Force Space Technology Center (AFSTC) was activated at Kirtland AFB, New Mexico, on 1 October 1982. Its stated mission is to centralize the planning and execution of space technology in support of future Air Force space systems mission requirements. The successful accomplishment of this mission is heavily dependent upon a vigorous planning effort which provides guidance for investments in space technology not only at the Space Technology Center but throughout the Air Force and DOD laboratory structure. This effort is key to successfully managing the Air Force space technology base and insuring a cohesive, integrated Air Force space technology investment strategy. The AFSTC does not intend to establish a space laboratory structure; rather, its purpose is to utilize the existent Air Force laboratory structure to meet space technology development goals. In some cases, the AFSTC will directly contract for selected advanced development tasks and/or demonstrations.

Upon its activation, the AFSTC was assigned command/manage-ment responsibility for three Air Force laboratories; the Air Force Rocket Propulsion Laboratory, the Air Force Geophysics Laboratory, and the Air Force Weapons Laboratory. As part of its mission, the AFSTC is responsible for protecting the non-space related technology activities at these laboratories to preserve the Air Force-wide technology support base.

The parent headquarters for the Air Force Space Technology Center is the Air Force Space Division in Los Angeles, California. AFSTC is thus an Air Force Systems Command organization. Close ties to the Air Force Space Command are provided through Space Division Office of Plans and through the Space Division Commander who also serves as the Vice Commander of Space Command.

SPACE TECHNOLOGY PLANNING

The core of the AFSTC planning effort is the Military Space Systems Technology Model (MSSTM). Expanding upon a NASA concept, the MSSTM provides a structured, systematic process

for linking future technology needs to project mission requirements. Wide distribution through the auspices of two American Institute of Aeronautics and Astronautics sponsored space technology workshops has insured broad government and industry participation in the model development. This participation has served to enhance the MSSTM's value as a communications tool among government and industrial agencies.

The recently published MSSTM is a five volume work which begins with a description of military missions which are or can be performed in space. Systems concepts to meet mission requirements are then derived and the technologies needed to enable the concepts are identified. The comparison of required technology figures of merit with current state-of-the-art and trend forecasts yields shortfall assessments which provide the basis for a technology plan designed to alleviate those shortfalls. Volume V of the current edition contains a series of technology roadmaps designed to meet the assessed technology needs. Unconstrained by available dollars, these roadmaps provide a beginning point for the prioritization effort which will follow in Volume VI.

The July 1983 kickoff meeting for Space Systems Technology Workshop III initiated the next MSSTM planning cycle which will sponsor a Workshop at Kirtland AFB in March of 1984 and culminate with the publication of the third edition in August 1984. This edition will include the Volume VI prioritized investment plan which will become the Air Force Systems Command Technology Plan for Space. As part of this effort the AFSTC is automating the modeling process to provide an automated, interactive process for deriving an optimal investment strategy based upon mission priorities and cost, risk, and schedule assessments.

AFSTC SPACE TECHNOLOGY - EMPHASIS 84

Recognizing that the MSSTM Volume VI prioritization process would not be complete until 1984, the AFSTC launched a major effort to construct an FY-84 investment strategy designed

to emphasize those technologies which had already emerged as critical to future space missions. Begun in the Spring of 1983, this effort was designed to provide an integrated investment strategy for FY-84 and a well constructed, defensible input into the FY-86 budget process. The remainder of this paper will be devoted to a description of the technology goals that emerged from that process and the rationale supporting their input into the AFSTC program. The technologies included on-board processing, autonomy, space prime power, surveillance and advanced military spaceflight technology.

ON-BOARD PROCESSING

Heavy Air Force emphasis on the survivability of future military space systems has generated a great need for increased capabilities in on-board processing. The ability to perform expanded signal and data processing tasks in space will reduce dependence upon vulnerable ground systems and enable the future use of mobile terminals. The problem, of course, is that the modern electronics elements needed to increase on-board processing capacity and speed are vulnerable to the radiation hazards present in the space environment. The challenge is to capitalize on the rapid, dynamic advances in electronic circuit design and manufacture while introducing hardening techniques to insure their survivability in space. In the immediate future, the AFSTC is concentrating on the development of an 8 bit 1750A hardened generic processing unit with a processing speed of 600-750 KIP's. This element is hardened to 10^6 RAD's total dose and is essentially immune to single event upsets. We are also initiating work toward the provision of a 2 million instructions per second (MIP) generic space qualified VHSIC processor by 1987 with a further advance to a 10 MIP processor in the early '90's. Accompanying component development will expand from the present 16K hardened RAM effort to 256K RAMS and 12M bubble memories in the late 1980's. We are joined by a multitude of government agencies in broad based effort to design and produce hardened electronic components for future space processing requirements.

SATELLITE AUTONOMY

With the increase in space on-board processing capacity and a speed, reduced satellite dependence on ground processing facilities becomes a real possibility. The ability to manage satellite health, determine spacecraft position and attitude and process mission related signals in space will diminish the frequency and scope of required ground contacts to enable the use of small mobile terminals. Such capabilities will increase both the survivability and endurance of space systems.

The current AFSTC efforts in autonomous systems include the ARMS program at the Jet Propulsion

Laboratory which is concentrating on autonomous health management, and the Multimission Attitude Determination and Autonomous Navigation (MADAN) Program, which focuses upon autonomous navigation and attitude determination. The former effort is primarily directed at algorithm development while the latter seeks to provide a space qualified solid state star tracking system.

ADVANCE SPACE PRIME POWER

Perhaps the most critical enabling technology for all future space systems is power. Nearly all the future concepts in the MSSTM call for power increases; some by 10's of KW, some into the 100KW to MW range. The AFSTC Technology Emphasis 84 program addresses this need with continued investments in solar power emphasizing gallium arsenide solar cells and nickel hydrogen batteries. New areas of solar power investment will include cascaded cells, concentrators, and high energy density rechargeable batteries. We are also hoping to kick off a new project to develop and test a high voltage, high power distribution system. These efforts project a doubling of current solar cell efficiency with power density increases into the 40-60 w/lb range. Similar increases in battery power density are projected with a potential factor of seven growth in watt-hr/lb by the early 90's.

Power demands in excess of 50-80 KW will exceed the practical limits of solar systems and require new approaches. While chemically-driven turbo alternators promise power output into the megawatt range, the duration of that output is limited by the fuel which can be carried into space. Space nuclear reactors present the only solution for the long duration, high (100KW-MW's) power systems needed for such applications as high power, space-based surveillance systems, high power jammers, anti-jam communications, electrically propelled orbital transfer vehicles and weapons. The thrust of the AFSTC FY-86 new start initiative for space nuclear power is to provide a follow-on commitment to the present DARPA/NASA/DOE SP-100 program with sufficient funding to make the commitment real.

STRATEGIC SURVEILLANCE

Despite some setbacks in infrared surveillance technology funding, the AFSTC continues to consider this an area of critical need to future space systems. We thus have proposed a strategic surveillance technology program designed to establish a useful data base of background and target signature data; to begin work on background suppression techniques; and to provide sensor/focal plane technology in wavelengths of interest to include the cooler technology needed to enable focal plane sensitivity. To support the latter objective, we envision two demonstrations of integrated focal plane assemblies, one MWIR and one LWIR.

ADVANCED MILITARY SPACEFLIGHT TECHNOLOGY

For several years, the AFSTC and its predecessor organization at Space Division have been attempting to launch a program designed to identify, develop and test long lead technologies for a quick response, multimission spaceflight vehicle. The program envisions investment in the key enabling technologies (propulsion, aerodynamics, structures/materials, avionics, etc.) needed to support the future development of such a system. FY-84 budget cycle saw the program funding survive until the very last when it was zeroed by the House/Senate Conference Committee. FY-85 POM/BES money still survives, however, and favorable support from SPACECOM and SAC gives us hope that we can "stay alive in '85."

SUMMARY

In the year since its activation, the AFSTC has become a viable force in focusing Air Force space technology efforts in support of future mission requirements. Its primary tool for providing this focus, the Military Space Systems Technology Model is progressing toward completion and has already played a key role in the construction of the AFSTC Technology-Emphasis 84.

The current AFSTC technology thrust as displayed in the FY-84 investment plan and the FY-86 POM submission emphasizes integrated investments in on-board processing, autonomy, prime power, strategic surveillance and advanced military spaceflight technologies. We plan to defend our budget packages vigorously with the hope that we will obtain the necessary funding to implement them as we have planned.

REFLECTIONS ON TRENDS IN DYNAMICS
- THE NAVY'S PERSPECTIVE -

Henry C. Pusey, Consultant
NKF Engineering Associates, Inc.
Vienna, Virginia

INTRODUCTION

At the First Symposium on Naval Structural Mechanics in August, 1958, Captain James A. Brown [1] of the Bureau of Ships suggested that "...the world is racing along the path of technological advancement at what appears to be an accelerating pace. Each bit of new knowledge spreads the light over wider vistas." Captain Brown was alluding to the rate of technological expansion "today" as opposed to the time it took in earlier periods. Keep in mind that Captain Brown's "today" was 1958; this is 1983. As an example of slower development in the past, he used the steam boat, the first of which sailed down the Fulda River in Germany in 1707. It wasn't for another hundred years that Robert Fulton's CLERMONT became the first steam boat to be a commercial success. Finally in 1871, England, the leading naval power of that era, abandoned sails as a standby means of propulsion for major warships.

Compare this with the changes in our Navy over the past 35 years. We have introduced revolutionary new hull forms (ALBACORE) and nuclear power (NAUTILUS) in our submarines. The ENTERPRISE was our first nuclear propelled aircraft carrier and the LONG BEACH our first nuclear propelled surface ship. Our weapons, and those of our potential adversaries, have evolved from simple guided missiles to precise targeting and tracking systems and to ICBM's with multiple warheads with independent targeting capabilities. At the same time, we have developed complex electronic systems as countermeasures against these advanced threats.

It is clear that this rapid technological change has had a great impact on warfare. It is equally clear that as weapons, ships, and planes become obsolete and new combat systems are developed, there is a significant demand on the technologies supporting these developments. Dynamics, that branch of mechanics which deals with forces and their relation to the motion of bodies, is an area that is important to designers of equipment for all types of systems. In this paper we will take a look at one area in Dynamics from the perspective of the

Navy. We will examine some trends in that area and highlight some problems that will require attention in the future.

DYNAMICS AND THE NAVY

Like the other Military Departments, the Navy uses aircraft, weapon systems, electronic systems and, with the creation of the new Space Command, space systems. With the latter the Navy has come full circle, since the Navy's Vanguard Program formed the nucleus for NASA's Goddard Space Flight Center. Just as the Army uses ships, the Navy uses tanks and other ground vehicles in its Marine Corps. The design of all these vehicles and systems involves significant dynamics problems, but in the interest of brevity, this paper will concentrate on the one system unique to the Navy, the naval warship.

At the present time, Navy planners are insisting on sophisticated missiles, superior electronic devices, higher speed, greater endurance, greater depths for submarines, more diversification in types of ships and more capability in each type. There is a demand for lighter ships so that they can carry more payload in weapons and equipment. At the same time, we are required to build structures which are stronger and more rugged. Meeting these requirements is not an easy task; it is placing increasing demands on the ability and ingenuity of the designer.

If we consider the problems facing the engineer who is designing the structures or equipment for the ships of our modern Navy, we may well conclude that his problems are in some ways similar to a designer concerned with airplanes. Both have to design complex structures or equipment which defy exact analysis, and both have to design to maximum dynamic loads that are difficult, if not impossible, to determine exactly.

To meet the difficult challenge brought forth by this complex set of dynamic design requirements, the engineering must be

continuously more innovative in his approach. At the same time he cannot rely only on his own capabilities. He must stay abreast with the latest developments in his field. He must seek out ideas, advice and counsel from his peers. And he should take advantage of every opportunity for interchange of information at symposia such as this one. In 1957, Dr. Elias Klein [2] reported on ten years of progress of the organization now known as the Shock and Vibration Information Center (SVIC). He very aptly stated that "The rate at which application of science is being made in weapons programs today demands that engineers and scientists have ready access to current developments related to their work. Hence the channeling of pertinent and new knowledge to the working scientist becomes of vital importance to the defense program. The information disseminated must be live and relevant; it must be communicated with directness and dispatch." Dr. Klein's remarks are even more valid today than they were then, and engineers are fortunate in that they still have SVIC available as a valuable central source for dynamics information. Further, with this 54th Symposium we now have more than 36 years of reported progress in the shock and vibration field.

In a paper like this, one cannot hope to cover all areas of dynamics pertinent to Navy interests. A look at the program for this symposium gives an indication of the breadth of such an endeavor. Not only would all the dynamic environments, such as shock, vibration, and acoustics, need to be covered, but topics such as instrumentation and measurement, data analysis, specifications, design methods, isolation and damping, dynamic analysis, and testing would have to be addressed for each of these environments. Mechanical shock has therefore been chosen as the area to be examined.

MECHANICAL SHOCK

Harris and Crede [3] define mechanical shock as a nonperiodic excitation (e.g., a motion of the foundation or an applied force) of a mechanical system that is characterized by suddenness and severity, and usually causes significant relative displacements in the system. The source of the shock excitation on ships is usually either an underwater explosion or a blast from the ship's own guns. Interestingly, shock damage from underwater explosions probably came first. It has been reported that the Confederate semi-submarine C.S.S. DAVID equipped itself with a long boom attached to its bow with a 60-pound charge at the end of the boom. The DAVID was maneuvered so that the charge contacted the hull of a Union ship. The charge was exploded, destroying the Union ship, but the shock from the explosion disintegrated the cast iron engine of the DAVID. It is understood that gun blast-induced shock damage was reported during World War I. Numerous incidents of shock damage were reported

[4,5] following the development of non-contact mines and bombs for attack on ships in World War II.

Shock Tests

During the winter of 1939-1940, the Germans laid large quantities of magnetically actuated ground mines in the waters surrounding the British Isles. Having no protection against this new weapon, many British vessels were sunk or disabled by the explosions of these mines. The mines being large (500-2,000 pounds of explosives) and the explosions occurring not in contact with the hull, many cases of complete disablement of a ship's machinery due to the shock effect were reported. In a large number of these cases, the hull damage was not of serious consequence, so that if it had not been for the shock damage, the vessel would not have been disabled as a combat unit. The cause of the majority of the serious cases of shock damage could be traced to the general use of cast iron in the British Navy; the second most pronounced defect was shown to be the susceptibility of the electrical equipment to shock.

Early in 1940, the Admiralty initiated a program to increase the shock resistance of machinery and equipment on British ships, and shortly thereafter a similar program was begun in this country. As a part of their program, the British developed a shock machine for testing equipment weighing not more than a few hundred pounds. Late in 1940, the United States Navy obtained the design of the British machine and, after a few modifications, it became the High-Impact Shock Machine for Lightweight Equipment. For testing heavier items, the High-Impact Shock Machine for Mediumweight Equipment was designed in 1942. The first test on this machine was conducted at the Naval Engineering Experiment Station in Annapolis, Maryland. The upper limit for equipment that could be shock tested remained at a practical limit of 4,500 pounds until the development of the Floating Shock Platform (FSP) in the early 1960's. Equipment to be tested is mounted on the platform and the shock loads are produced by a nearby underwater explosive charge. A large version of the shock platform (FSP) is now operational for testing equipment weighing more than 300,000 pounds. A submersible version of the FSP, called the Submarine Shock Test Vehicle (SSTV) has also been introduced to test submarine equipment and systems. Clements [6] provides a thorough description of the Navy shock test devices and their operation, although the report was written before the development of the last two facilities.

Various research studies relating to underwater explosions have been conducted over more than a century with particular emphasis beginning about 1940. Numerous technical papers indicate that the underwater explosion phenomena are well understood [7,8,9]. The classic paper

by Keil [9] gives an excellent insight into the shock environment on a ship subjected to an underwater explosion. Yet, as Keil pointed out in this 1961 paper, "The actual accomplishment of shock hardening or shock toughness is demonstrated by shock tests. The equipment can be installed (for the test) either on the anvil of a shock machine or in a ship." To this day, shock testing is the preferred method for confirming the shock resistance of equipment.

Ship Shock Tests

During and following World War II, the Navy embarked on a program for the shock testing of full-scale ships. Initially, these testing efforts were rather exploratory in nature and aimed primarily at improving our understanding of underwater explosion phenomena and the relationship of these phenomena to ship vulnerability. These early tests also served to put shock resistance goals in proper perspective relative to other ship capabilities and limitations. Later research-oriented ship tests as typified by the KILLEN and FULLAM series, have been largely devoted to the development and refinement of the shock design and test criteria now specified for new construction by the U.S. Navy.

Research type ship tests are most often performed at shock severities ranging from moderate to severe, thus necessitating the use of older, expendable ships as tests targets. While such testing procedures offer many advantages, it is becoming increasingly necessary to reconfigure older ships extensively in order to acquire certain information directly applicable to today's more sophisticated warships. Reconfiguration of ships for shock testing purposes is both expensive and time consuming, making it more difficult to obtain approval for tests in this category.

Full scale ship tests conducted shortly after World War II demonstrated that combatant ships' mission keeping capabilities could be lost or seriously impaired at very low levels of attack severity. This revelation prompted concern for the safety of vital operational ships and led to the evolution of what is now known as the "routine ship test program". This program provides for "routine", standardized shock testing of the first ship of most new classes, and also for testing of selected representative ships from older operational classes.

Operational ship tests performed under the routine ship test program are not conducted primarily for research purposes, but rather to permit identification of items critically deficient in shock resistance due to improper design or faulty workmanship. Once isolated in this manner, the conditions responsible for inadequate shock resistance can be corrected by backfit shock hardening.

Shock Design

Early shock design procedures were to a large extent empirical. Designers relied on judgement and good engineering practice insofar as possible, and out of their experiences grew a number of qualitative guidelines or "rules of thumb" for shock design. Welch [5] was among the earliest to provide written guidance for the design of shockproof equipment. One of the design rules that evolved from design experience was the "static g" method. Using this method the designer was told that the equipment should be designed for static loads equal to N times its own weight, with N (the number of g's) varied according to orientation and weight. This procedure was made a part of the specification for equipment too large or too heavy to be tested on the shock machine. The procedure had its drawbacks in that it did not account for the differences in mounting/foundation frequencies, locations in the ship, or ship types.

Recognizing the deficiencies of the "static g" method, the Naval Research Laboratory in the 1950's sought procedures which would promote more realistic shock design. This resulted in a method to evaluate equipment design based on dynamic loads, now known as the Dynamic Design Analysis Method (DDAM). DDAM [10] requires that a mathematical model be made of a piece of equipment and that its response under dynamic load be determined, using realistic inputs provided by the Navy. The inputs are made possible by the data from ship shock tests and do take into account the type of ship and the location on the ship. The inputs are specified as design shock spectra, and the analysis is made possible by the digital computer. The failure criterion is basically the effective yield strength of the material together with a factor that takes account of the efficiency with which the material in the number being analyzed is utilized. DDAM is now specified as the acceptance method for shockproof items of equipment which are nontestable. Ample Navy guidance documents on the application of the method are provided [11].

Shock Spectra

In 1943, Biot [12] defined a quantity called the "effective acceleration of the earthquake for the period T". From this the present concept of earthquake spectra evolved. In 1949, Walsh and Blake [13] applied the earthquake spectra concept to the mechanical shock problem, resulting in what is generally known as the shock spectrum. Various authors have used the term "shock spectrum" in different ways. From the Navy's viewpoint, a shock spectrum is a plot of the maximum absolute values of the relative displacements of a set of damped (in general) single-degree-of-freedom oscillators with negligible mass which have been subjected to a shock motion versus the natural frequencies of the oscillators. In some cases

response of single-degree-of-freedom systems to the applied foundation motion.

O'Hara [14] introduced the concept of the design shock spectrum, the form used to describe the inputs for a DDAM analysis. A design shock spectrum is a plot of the values which enables an analyst to predict the stresses in a selected structure for a specific type of excitation such as an exploding mine. This special kind of spectrum is a mathematical concept rather than an easily measurable quantity.

G. J. O'Hara and R. O. Belsheim were the developers of the Navy's Dynamic Design Analysis Method. Together and separately they contributed greatly to the advancement of Naval shock design and analysis. O'Hara [15], for example, introduced what is called the "shock spectrum dip". He explicitly showed that structures on nonrigid foundations, when excited by a shock motion, feed back forces into the foundation which effect the motion in such a fashion that the spectrum values of major interest for a shock tend to lie in the region of a valley rather than in the vicinity of a peak of the plotted spectrum. O'Hara's work demonstrated that overconservatism in design can result from incorrect usage of shock spectra.

There are many examples of breakthroughs in shock analysis too numerous to cite here. Suffice it to say that great strides have been made in the use of dynamic analysis to assist in and confirm shock design.

Shock Measurement

This discussion of ship shock would not be complete if it did not include a few words about instrumentation for shock measurement. This is especially true because of the great advancements that have been made over the last 50 years and because of the importance of measurements to provide data for rational design.

In 1943, Vigness [16] described the shock measuring instruments generally in use at that time. They were quartz crystal type accelerometers, high speed moving pictures (up to 3,000 frames per second), and wire type strain gages. The quartz accelerometer, described as the best instrument for measuring impact accelerations at that time, was fraught with errors from zero shifts, phase shifts, and cross axis sensitivity. There was also a British velocity meter available in 1943, but it was bulky (about 35 pounds) and good for a travel of less than one inch.

By 1960 Vigness [17] described quartz as having been rendered obsolete by barium titanate as a piezoelectric sensing element for accelerometers. Relatively compact seismic type velocity pickups were in common use with less than five per cent error if operated in the range above three times its natural frequency.

High speed photography had reached speeds up to 15,000 frames per second. Also by 1960, both analog and digital computers were available to calculate shock spectra from the measured data.

Today we have extremely accurate, highly reliable transducers coupled with very sophisticated computers. It is safe to say that we can make the measurements that we wish and massage the data to present it in almost any form imaginable. It is not uncommon in a full-scale ship shock test to have several hundred channels available for taking data. However, in my view many measurements are taken without a clear understanding of why, and with no preconceived notion of how the data will be used. In spite of this, we are improving in this area, as evidenced by the more precise definition of inputs for DDAM analyses.

Shock Isolation

If one has the objective of improving the shock resistance of a piece of equipment, the usual first thought is to use some sort of resilient mounting so that a cushioning effect is provided to the equipment. Although the use of shock isolators often produces the desired ruggedness, they often complicate the design, increase the overall weight and add additional maintenance problems. Furthermore, the design of a shock isolator for shipboard equipment can be a tedious problem. First of all, the isolator must have an adequate stiffness and permissible deflection to respond to the maximum shock motion in a way that reduces the severity of the shock as it is transmitted to the equipment. At the same time it must have a stiffness adapted to isolate the vibration of the structure of the ship in response to the shock; it must also have a means to prevent excessive vibration of the equipment as a result of propeller-induced vibration (either damping or a relatively high natural frequency).

For part or all of these reasons, the Navy's policy over the years has been to produce intrinsically shockproof equipment through design without resorting to shock isolators. Shock mountings have been employed only for delicate and complex equipment for which a shockproof design was not feasible. There was an excellent early guidance document on the use of shock mounts on ships by Crede [18], and reference [3] provides an able treatment of the principals of shock isolation.

STATUS REPORT ON SHOCK

We have now examined, however briefly, some of the important facets of mechanical shock as it related to ships. It is appropriate now to provide a brief status report on our progress related to shock. To give meaning to such a report, we need a starting point.

In 1940, an intensive program of development and investigation related to ship

an excellent early guidance document on the use of shock mounts on ships by Crede [18], and reference [3] provides an able treatment of the principals of shock isolation.

STATUS REPORT ON SHOCK

We have now examined, however briefly, some of the important facets of mechanical shock as it related to ships. It is appropriate now to provide a brief status report on our progress related to shock. To give meaning to such a report, we need a starting point.

In 1940, an intensive program of development and investigation related to ship shock problems was initiated. That program, with modification from time to time, has continued to the present date. Captain Ron Trossbach described the latest major thrusts related to the shock hardening program in another paper at this symposium. The shock program, as defined at its inception, had several phases, all of which by necessity have been carried on concurrently. These phases were as follows:

1. The development and application of methods of improving the shock resistance of presently installed equipment.
2. Redesign of equipment for new construction to accomplish inherent shock resistance.
3. Development of shock testing machines to simulate the type of shock encountered aboard ship, and the installation of a large number of these machines in the plants of manufacturers and naval laboratories.
4. Experimental and theoretical investigations of the nature of shock and shock failures, including the development of instruments to measure shock. This phase has included a number of full-scale shock tests, from which the majority of fundamental data on shock has been obtained.

This program, taken as a whole, has produced firm shock hardening goals. Pursuing these goals has resulted in significantly improved shock resistance for most shipboard equipment. Although it is true that poorly designed equipment still slips "through the cracks" because waivers or extensions have been improperly granted, much of this "weak" equipment is being exposed during routine ship shock tests.

Our capability in the shock test area has expanded far beyond the wildest dreams of the shock program pioneers. Our capability for shock measurement is probably limited only by our ability to apply the data. Rapidly advancing analysis techniques coupled with better modeling procedures has resulted in significantly improved design methods. In

general, we can be satisfied with our progress in the ship shock area, but we cannot be complacent. Research must proceed so that it will continue to lead to improved, refined and more diversified methods and techniques of shock hardening.

SOME FUTURE NEEDS

It would be impossible to list all future needs relating to ship shock; the following offers only a few suggested items that need attention. Some of these are from my own observations, and some are drawn from the suggestions of associates.

- There is a need for more diversified testing techniques such as structural scale model testing.
- Fixtures for simulation of shipboard installation characteristics need improvements and further studies.
- There is a need for a central data bank, with easy and efficient access and retrieval systems which will reflect past experience with machine and barge testing.
- New equipment and systems to be developed and introduced into the fleet pose the problem of their capability to be shock hardened. An excellent opportunity to test and possibly harden the system is on a new ship concept; The Test and Evaluation Ship. This ship concept is based on a dedicated platform, new or conversion, for the sole purpose of testing and evaluation. A study should be performed to assess the feasibility of the concept and its impact on the ship hardening program in terms of effectiveness versus cost.
- More experimental work is needed in the dynamic yield of structural materials.
- More research and application-oriented work is needed in the area of plastic design methods.
- Exceptional analyses techniques should be developed and adopted for special cases such as structural analysis of underwater appendages subjected to the direct shock wave (e.g., propellers, sonars, rudders, fins, etc.).
- Methods for evaluating nonlinear structures subjected to shock motions should be developed (e.g., large deflections, nonlinear mounts, etc.).
- Research work is needed in the analysis of an entire ship subjected to shock pressure waves. This may lead to the

development of shock design gradients tailored for a specific ship.

- Although great improvements have been made in the area of availability and access to technical information, there is still a need for continuous efforts in this area. There is a need for a data bank which will include all pertinent information with regard to shock.
- There should be more emphasis on training in the ship shock area.

REFERENCES

1. J. A. Brown, "Problems Related to the Design of Structures for Ships of the U.S. Navy," Proc. First Symposium on Naval Structural Mechanics, J. N. Goodier and N. J. Hoff, Eds., Pergamon Press, 1960.
2. Elias Klein, "Progress in Shock and Vibration during the Last Decade - Correlation and Dissemination," Proc. of ONR Symposium on a Decade of Basic and Applied Science in the Navy, March 19 and 20, 1957.
3. C. M. Harris and C. E. Crede, Editors. "Shock and Vibration Handbook," 2nd Edition, McGraw-Hill Book Company, 1976.
4. J. M. Fluke and H. W. Hartzell, Shock Damage on Naval Vessels," Proc. Symposium on Shock, Electrical Equipment Shock Committee, Bureau of Ships, 30 October 1943.
5. W. P. Welch, "Mechanical Shock on Naval Vessels," NAVSHIPS 250-660-26, Bureau of Ships, 1946.
6. E. W. Clements, "Shipboard Shock and Navy Devices for its Simulation," NRL Report #7396, July 1972.
7. E. H. Kennard, "Report on Underwater Explosions," David Taylor Model Basin, Report #480, October 1941.
8. "Compendium on Underwater Explosions Research," Vols. 1, 2 and 3, Office of Naval Research, 1950.
9. A. H. Keil, "The Response of Ships to Underwater Explosions," David Taylor Model Basin, Report #1576, November 1961.
10. NAVSHIPS 250-423-30, "Shock Design of Shipboard Equipment, Dynamic Design Analysis Method," May 1961.
11. NAVSEA 0908-LP-000-3010, "Shock Design Criteria for Surface Ships," Naval Sea Systems Command, May 1976.
12. M. A. Biot, "Analytical and Experimental Methods in Engineering Seismology," Trans. ASCE 108:365, 1943.
13. J. P. Walsh and R. E. Blake, "The Equivalent Static Accelerations of Shock Motions," Proc. SESA 6(No. 2):150, 1949.
14. G. J. O'Hara, "Shock Spectra and Design Shock Spectra," NRL Report #5386, Nov. 1959.
15. G. J. O'Hara, "Effect upon Shock Spectra of the Dynamic Reaction of Structures," NRL Report #5236, December 1958.
16. I. Vigness, "Measuring Instruments for Shock," Proc. Symposium on Shock, Electrical Equipment Shock Committee, Bureau of Ships, 30 October 1943.
17. I. Vigness, "Instrumentation, Analysis, and Problems Covering Shock and Vibration," Proc. First Symposium on Naval Structural Mechanics, J. N. Goodier and N. H. Hoff, Eds., Pergamon Press, 1960.
18. C. E. Crede, "Shock Mounts for Naval Shipboard Service," NAVSHIPS 250-600, Bureau of Ships, June 1944.

ELIAS KLEIN MEMORIAL LECTURE
MODAL TESTING - A CRITICAL REVIEW

Strether Smith
Lockheed Palo Alto Research Laboratory
Palo Alto, CA

When I was looking at the title I had chosen for this paper, it made me think about what a critical review really means. You don't do critical reviews on a science because a science is something that is fact; unless, of course, it is wrong. You do critical reviews on art. I will discuss modal testing as an art form, and start from very early history where it all came about.

In the beginning God said, "I will create heaven and earth, and in my image I will create man. From his rib I will create woman, and call them Adam and Eve." And he told them of the rules of living in the Garden of Eden; they could eat from any tree except the Tree of the Knowledge of Good and Evil. It is little known, but God also told Adam and Eve that they should go out and create structures to house themselves, to carry them over the earth and even to the stars. They should be characterized by certain properties, namely, linearity, reciprocity, and distributed damping. We all know what happened at the end of the story. Adam and Eve got themselves in deep trouble and were thrown out of the Garden of Eden. Obviously, their descendants didn't pay any attention to the other things that God had said either. Actually, I think that Eve has been wrongly criticized for her actions in the Garden of Eden. I think she actually ate of the fruit of the Tree of Knowledge absentmindedly because of her preoccupation with reciprocity and matters of that sort.

Well, man has since built many structures. Some of them were successful, and some of them were not so successful. Figure 1 shows one built by Otto Lillienthal that obviously didn't follow the rules for linearity and such things very well. In fact, in the end, Otto got caught in a structural failure; and he was killed in a mechanism similar to this. To compensate for this fall from grace, Adam and Eve's descendants have attempted to develop sciences to explain the behavior of nature. They have tried to follow the word of God to develop structures or to characterize structures in a manner that is ideal. One of those is what we will discuss today - modal analysis.

Modal analysis is the art, or the science as I have already discounted, of characterizing the dynamic behavior of a structure in terms of its normal modes. The fundamental characteristics of this science are that if you use normal mode theory, the motions of a structural system can be described as the sum of a discrete set of independent and predictable motions (Fig. 2). These characteristics are called normal modes, natural modes or characteristic modes of vibrations. Each of these modes is characterized by only three parameters: the natural frequency, the damping behavior expressed in some simple sort of characteristic, and a deflected shape of each mode. Using this simplistic theory, we are able to do many things. First of all, if we attempt to predict the behavior of a structure due to some input, we can make up an analytical model which predicts the response of the system as a sum of mode shapes; from that we can predict the vehicle response to any set of inputs. Using the concept of modal characteristics, we can also perform an experiment which will determine the characteristics which we can compare with an analytical model. This is the primary use of modal testing: to substantiate analytically-derived models of structural behavior.

So, with these things in mind, it is really tempting to assume that a structure is linear, has distributed damping, and it isn't influenced by other characteristics that aren't included in this relatively trivial theory. The results of not using these assumptions are extremely painful; they require us, at least with our present science, to use integration techniques to predict structural behavior which are relatively expensive compared to the simple straight forward modal analysis theory. Also, they are generally not particularly reliable.

Let us consider what is behind the theory that we are discussing. The assumptions are that the stiffness of the structure is a constant. If you put a load on it, it deflects a certain amount; if you double the amount of load, the deflection will also double. It also assumes that the damping is linear and, for this discussion, we assume that the damping is

viscous, i.e., the forces due to damping are proportional to the velocity. Under these conditions we can use a relatively simple expression to show the input-output response between two different points on the structure.

$$H_{ik} = \sum_{n=1}^N \frac{Y_{in} Y_{kn}}{K_n \left[1 - \left(\frac{\omega}{\omega_n} \right)^2 + 2j \zeta_n \left(\frac{\omega}{\omega_n} \right) \right]} \quad (1)$$

The frequency response function H_{ik} relates the forcing input at point i to a response at point k which is simply given as the summation of the response of the "N" modes of the system. The mode shape of the n th mode at the i th and k th locations, the stiffness of the n th mode, the natural frequency of the n th mode, and the damping associated with the n th mode give the frequency response function, or characteristic response, as a function of the driving frequency, (ω) .

Under these assumptions we have a large number of well developed tools that are all based on normal mode analysis. We have a huge science of finite element eigensolution methods for theoretical prediction of modal behavior. We can do linear superposition of responses. If we know what a characteristic response is at one point, we can add another forcing function, and we know what the response to that is. We can use the concept of reciprocity to extend and to verify testing results. The structure can also be characterized by fitting measured data to a simple theoretical model, specifically, the expression in equation 1. Also the system modes are real.

The concept of superposition is used to model complex systems. We add all of the modes of the system, and we add all of the generalized input functions to the system. We know what the responses are for a complex system as long as they are linear and follow the basic ground rules. Figure 3 demonstrates this concept, where the total response of the system can be made up of the responses of various modes of the system times their generalized input forces.

The concept of reciprocity means that if we force the specimen at one location, and measure the response in another location, we will measure a frequency response function that is identical to the one determined by swapping the forcing point and the response point (Fig. 4).

A further implication of linearity is that the mode shapes that we determine are independent of the excitation location.

The concept of curve fitting for linear structures is shown in Figure 5 where we have some experimentally measured data (Fig. 6) on an extremely linear structure which is very well fit with analytical expression in Equation (1).

We have also said that measured modes should be real for a linear structure. That means that all modal response should lie on the

same line in the imaginary plane. But even for extremely linear structures, this is seldom the case. Figure 7 shows an example taken from a very stiff linear structure, where phase differences approaching 30° at points of significant response are apparent. The first column is the normalized modal amplitude for a bending mode. One can see errors of 27° at one point and -15° at another point for a spread on the order of 40° .

As it turns out experience shows that failure of the mode-shape realness and reciprocity criteria are common occurrences, even for relatively linear structures. As a matter of fact, the only structure that I have ever seen approaching real modes and good reciprocity is on the extremely stiff and intentionally lowly damped Space Telescope Simulator (Fig. 8). Most of the structures we have to live with are not so nice to handle.

Deviations from the ideal come from many sources. First, nonlinear material behavior is something that is becoming more and more important. Solid rocket motors inherently have nonlinear material behavior in their propellant. Joints that have nonlinear springs and dampers are becoming commonplace. In general, structural joints produce most of the damping, and this violates the distributed damping criteria that we started with. Energy dissipation doesn't follow viscous damping, structural damping or any of the nice models that we like to work with. Obviously large deflections cause errors due to the "theta-equals-sine-theta" criteria for linearity. Temperature, phase of the moon, and poor modal testing techniques are other problems that plague us.

A typical example of a structure which has nonlinear joints is the antenna shown in Figure 9. Structures of this type get most of their stiffness and most of their energy dissipation from their joints around the root. The joints are inherently nonlinear in that they have significant slop to come up against stops and cause all sorts of problems with our basic analysis theory. Other types of structures that have nonlinearity problems are solid rocket motors with difficulties with the propellant, and there are also strong nonlinearity problems with the nozzles that are associated with rocket motors in the guidance systems, again because of joint and actuator nonlinearities.

The kinds of things that we expect to see from these deviations are modal frequencies that depend on the excitation function that we use. Obviously, for a nonlinear structure, reciprocity is going to fail, and if we are not very careful with our test techniques, we probably will not get repeatable results. An example of nonlinear jump behavior which occurs in sine testing of nonlinear structures is shown in Figure 10. As we sweep upwards in frequency, a "jump" in amplitude, whose frequency and magnitude are drive-amplitude dependent, occurs

and then the response slopes off. The superimposed line is a least squares fit to the data using linear theory which is obviously a poor representation of the measured response.

Despite these problems, the objective of modal analysis is to attempt to characterize the dynamic behavior in terms of normal modes because if we get some other results, we don't have the required science to do anything with it anyway. We can characterize nonlinearities in some ways, which I will discuss, but in general, we don't have the science ready to do it.

Figure 11 shows most of the available modal testing options analysis. The left-hand side shows a selection of excitation functions that are available; multi-frequency functions include ambient, twang or impulse, random and chirp; sinusoidal excitation includes broadband sinusoidal sweeps with single exciters, narrowband sinusoidal sweeps with single exciters and multi exciters. All of these methods can also be used with multiple exciters. The multi-frequency methods produce time series data. The sinusoidal sweep techniques produce spectral data directly. Time series may be converted to spectra using the FFT and frequency response functions are calculated by dividing the response spectra by the forcing spectra. We have a variety of tools to do either frequency domain analysis to get modal parameters, or to do a time domain analysis using impulse functions, calculated by inverse FFT of the frequency-response.

Impulse or twang testing is the cheapest choice. It has the advantage of being inexpensive and it is convenient, but it has the disadvantage of not being able to attain high amplitude response which is often important. It requires skill; and without skill, it produces a poor excitation spectrum. I don't recommend that technique except in certain highly linear structural cases, or if you have no other choice.

Sine-sweep excitation, which comes and goes in popularity, is one of my favorite techniques. It has the advantage of, if you do it the right way using the SWIFT algorithm that was developed at Lockheed about ten years ago, of being extremely fast if you have a large number of channels. It becomes the fastest technique when you have more than 64 channels. You can get any data density in the spectral domain that you want. It concentrates exciter power at the frequency that you are driving, and it has the capability of characterizing nonlinearity as shown in the plot in Figure 11. Its primary disadvantage is that it is horribly slow if you don't have a lot of channels and that sine excitation is not a realistic simulation of service histories of the structure.

Chirp (fast sine sweep) excitation is the next step along. It has advantages shared with other, multi-frequency excitation techniques, in

that it is fast for a small number of channels. It is sensitive to nonlinearities. You can get different results from an up-chirp and a down-chirp. In fact, you can force the data to look linear by averaging the results of up- and down-chirps. You can't characterize the nonlinearity. Also, the excitation is not a simulation of real-world input.

Most people's favorite is random excitation. It is the primary excitation technique used by most of the commercial systems. It has the "advantage" of making nonlinear systems look linear. It gives average parameters for nonlinear structures, whatever average means. It is fast for a small number of channels. One of its worst disadvantages is that it overestimates the damping of nonlinear structures, which is a non-conservative result.

Let us briefly discuss the difference between sinusoidal and random excitation with a nonlinear structure, specifically a softening spring system. If we sweep from a low frequency to a high frequency at a low amplitude, both the sinusoidal and the random excitation tests will give the same frequency response function. As we increase the amplitude, the indicated frequency will start to slide to lower frequency for the sinusoidal test. When we get to a high enough frequency, where the so-called "back bone" turns over, we will get jump phenomena similar to the data that is in Figure 11. Random excitation, on the other hand, can't see this kind of behavior, and you will see a slight sliding of the frequency to lower values as the drive level is increased. In general, it looks like a nice linear amplitude characteristic.

A new excitation concept has just hit the streets. Multiple input random excitation extends the random-input concept to allow excitation at many locations simultaneously. If we applied this to friend Otto's aircraft, Figure 1, it means that we put shakers on several locations on the structure, somewhere between three and six, and we measure the response in many different places. We then reduce the data using matrix algebra techniques to determine the frequency-response-function matrix. The claimed advantages of this technique, when compared to the single exciter techniques, are that the data are more consistent. Well, that's nice. You only have to do one test. You are using all of your shakers at once. You don't have to worry about the fact that when you shook with shaker A, and got one result, and then you shook with shaker B, you got a different result. You don't have to worry about that problem anymore which makes the management a lot happier. The reciprocity is also greatly improved. The data are more realistic because the excitation responses are higher, and they are distributed over the structure, which is probably a good representation of what the structure will see in real life. The test has to be done only once, so in principle, it takes less time.

Its disadvantages are that it takes a significant computer to perform the analysis. You can't see your frequency response functions in real time with present-day hardware. You have to record your data, go away in a corner, and come back with your results later. It's a little scary if you can't keep your test specimen. By standards relative to the commercial data acquisition world, a fairly large data acquisition system is required; at least eight or 16 channels are needed to make it worthwhile. Not to be facetious about it, but the funny thing about this concept is that if the structure is as linear as the test technique would lead you to believe, then you don't need the test technique to do the testing.

That is not to say that I don't think it is a really good technique. The techniques need to be developed and implemented on hardware that can reduce the data in a reasonable (1-2 hour) time period. Then they must be used with full knowledge of the technique's tendency to "linearize" the structure.

The bottom line is that many excitation functions are available, but most of the time you can't pick the one that you would like from a mathematical or other rational standpoint. You should pick the one that makes the most sense. For instance, we would not take an electrohydraulic shaker out into eastern Oregon to test a transmission line structure, and we would not take an electrohydraulic shaker out into space. But if we have a clean laboratory environment and a simple structure, then it would be nice to have all of the excitation tools available to us. Fortunately, I have been involved with a system that allows us to use any of these techniques that we've discussed.

We have discussed what the excitation techniques are, and how we measure the frequency response functions, but, we still need to analyze them for modal properties. Starting in the late 1940's, people started to worry about how they would analyze frequency response functions to get modal parameter information. In the 1940's, Lewis and Wrisley developed a technique called the tuned dwell by which they used multiple shakers and sinusoidal excitation tuned to produce a response that they felt was a single mode. This technique has endured to this day, and it was used on a test this year.

This technique used several shakers and sinusoidal excitation to excite one mode while attempting to suppress all others. In fact, attempting to suppress all others is the important phrase, because if you suppress all others, it does not make any difference how well you do with the mode you are interested in if it is the only one going. The problems with the Lewis and Wrisley technique was that at that time, there were no objective tuning procedures, no purity criteria, and the technique produced undetectable errors.

Kennedy and Panu went off in the

sinusoidal sweep direction realizing that there were better ways to reduce the data if you could see your entire frequency response function. They realized when you plotted the real and imaginary response of the frequency response function in a complex plane, it produced a circle. They came up with the idea, which is still used in some systems, of doing circle curve fits to the frequency response data to do a modal analysis and to determine the modal response.

At Lockheed, we extended that technique to a global approach which allowed us to determine the modal frequency and damping based on global behavior of the whole structure as long as the mode was properly tuned. We used the technique of Asher with some success to develop a rational tuning criteria to do the modal analysis. This technique was the one originally installed on the 256-channel Modalab system.

At about the same time, Klosterman was developing the idea of multi-frequency excitation and using circle fit analysis to do the modal analysis. Shortly thereafter, Richardson and Potter came up with the idea of doing a genuine curve fit to the data analysis, which is the technique that is used in almost all of the commercial systems today.

More recently, work that was started by Sam Ibrahim and continued by him and quite a few others, developed the concept of time domain eigensolution techniques. The impulse response function is determined, either by a free decay of the structure, "random-dec," or by an inverse Fourier transform of the frequency response functions. Then an eigensolution analysis is performed to determine the modal frequencies, damping and the mode shapes. This is presently the "Cinderella" concept. It has been used primarily in research environments and it is claimed to enable the extraction of high density, highly damped modes which are impossible to extract with a standard curve fit analysis. It is also claimed to be relatively noise insensitive. A problem with this technique is that processing the data requires significant computer capacity. It is not commercially available, and it is still being proven. However, I think we will see it commercially developed before very long.

One of the concepts that I want to stress is, with modern modal testing ideas we need to separate the problem into two categories. The first step is to measure the frequency response functions and get a clean set of transfer functions by any means that seems appropriate for the equipment and environment that you have at the time. Then you can transform either to the time domain or frequency domain, and then you have a whole family of analysis techniques that you can use to get the modal behavior. So one should not be stuck with one path through the modal technique map (Figure 11) by any hardware that they care to buy.

What do you have to worry about? What are your objectives? Do you want your specimen to behave linearly? Do you think your specimen really is linear? For instance, we were pretty sure the Space Telescope was linear. All we had to do was prove it. The chances are good your structure will not be linear, and you will know it. Then you must decide how you want to characterize it. Do you want to try to characterize the nonlinear behavior of the structure? Practical considerations include: (1) what kind of excitation capability you have, (2) how many channels of data acquisition are available, (3) how many channels of transducers do you have. In the analysis area, what do you have in the way of software to do the data analysis?

An interesting modal test was recently performed at the Jet Propulsion Laboratory (JPL) on the GALILEO Spacecraft. It was interesting because the people at JPL were fortunate enough not to have to make any real decisions about which test technique would be used; they used them all. The following test methods were used: tuned dwell-decay, single shaker swept sinusoidal excitation, single shaker "chirp" excitation, single shaker periodic random and decay, tuned sinusoidal sweeps using the SWIFT algorithm and multiple input random excitation. The people who participated in the test are analyzing the data using the following methods: Dwell-decay measurement analysis, frequency domain curve fitting, time domain eigensolutions, and frequency domain eigensolutions. I am sure the results of this test will be the subject of many papers for years to come at symposia like this.

The modal testing techniques that we have discussed have been primarily driven by hardware capabilities, particularly in the area of mini-computers. Fortunately, it looks as if we are coming into a new generation of hardware and software systems that are very exciting, and that will revolutionize modal testing and analysis. New display technology and intelligent systems techniques will help us enormously in deciding what approach to take as we are doing both the modal data acquisition task, and data analysis. New computer hardware using extremely high-speed buses and distributed processing will allow real-time processing of data from multiple input random tests with a large number of channels.

The state-of-the-art in data-acquisition and analysis hardware is that commercial systems are now available that are capable of acquiring 500,000 samples per second to disc; this is something that was only done in special test systems just a few years ago. The cost of signal conditioning, which is the most expensive part of a large modal testing system, is about to be cut by a factor of ten to \$100/channel by the concept of switched capacitor filters. This will mean that relatively large scale modal analysis systems will be much more realizable than they have been in the past. In the area of

extremely high speed data acquisition systems, Lockheed is presently building a system that will record 5,000,000 samples per second to disc for over three minutes, and at the same time, have real time data visibility to show the operator the status of the test.

What will next year's large modal testing system look like? For those of you who have heard of, or know of, the "Modalab" System, that was something that we built at Lockheed 10 or 12 years ago. At that time it was an extremely powerful system using the PDP-11/45 mini-computer. Until a few years ago it was able to support all of the modal testing techniques that were available. Now we are finding two things. First, the poor thing is old, tired and worn out despite adding hardware and software to it. But we have also decided that its computer power is not enough to do the kinds of things we want to do. So we are designing Modalab II, to be constructed in late 1984. It will be characterized by new technology. The host computer that we have tentatively selected is a VAX 750 (old technology) which will allow us a good, software-friendly system, but it will not do very much of the work. The work will all be done by a high speed bus system and input processors which will control the data acquisition and storage. An array processor will allow us to calculate frequency response functions on the fly in the input. We will have command generation, controlled by the high speed bus system which will allow arbitrary-function generation for long periods of time.

What will we be testing in the future? The most interesting one is large space structures which brings up a whole new problem area in modal analysis. First of all, somebody has the weird idea that we will test large space structures on the ground when they won't even hold themselves up. This means that we will have to come up with some sort of a suspension system that will hold them, and this will be the hardest part of that problem. The suspension system will have to have a long stroke, a very low basic frequency, and will hold up relatively small masses. Of course there is talk about doing modal testing in space, but talk is about as far as it has gone.

The art of modal testing and analysis is presently experiencing a renaissance. New techniques are being introduced regularly and many of them appear to be very promising. However, as with any technique, it behooves the investigator to investigate the underlying assumptions of the method and to assess their effects on the results.

DISCUSSION

Dr. Showalter (Naval Research Laboratory): Strether, if you were running for a plane, and if you only had five minutes, and you had to advise somebody how to pick a modal analysis system in five minutes or less, what advice would you give to them?

Mr. Smith: Wait more than five minutes! It gets back to the question of what your test is all about. There are many good commercial systems available. All of the vendors sell good systems. The systems all work for possibly a fairly restricted area. However, not knowing what your test is about, I would say buy the cheapest one.

Voice: I think everybody wants to know what a switch capacitor is.

Mr. Smith: Does everybody know how a filter works? A filter is a network of resistors and capacitors. One of the problems with them is that the natural frequency of filter is governed by the RC constants of that filter. The easy way to make a programmable filter in the past was to change the resistor. This had to be done either with an analog multiplier, or something like that, if you wanted continuous changing, but that restricted your dynamic range drastically. Another way to get a programmable filter is to switch a whole bank of resistors. This means, for an eight pole filter, you have to switch eight resistors for each frequency setting that you want. So you need a couple hundred resistors and a couple hundred switches for each of the filter channels. That's why they cost \$1,000.00 per channel. A switched capacitor filter, in essence, changes the capacitor by using a time-sharing technique. They essentially turn the capacitor on and off at a high rate to change its duty cycle. It has some disadvantages one of which is noise. It allows a continuous change in cut-off frequency dependent on an input clock frequency, which is extremely convenient, because it is a clock we are probably using for something else. These modules are sold by several semi-conductor vendors for on the order of \$30.00 per channel for an eight pole filter.

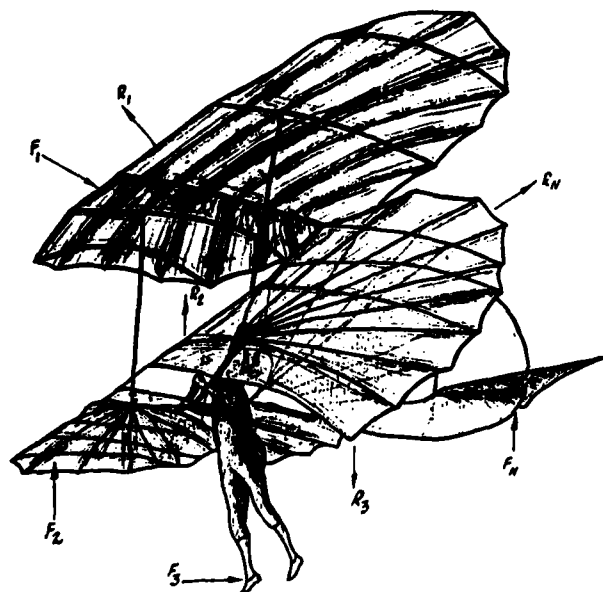


Fig. 1 — OHO Lilienthal's Aircraft

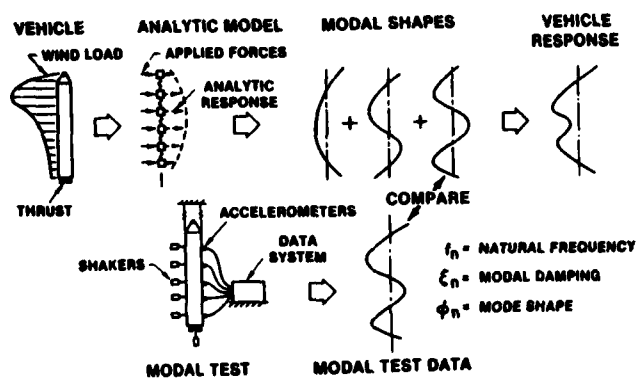


Fig. 2 — Concept of Modal Testing

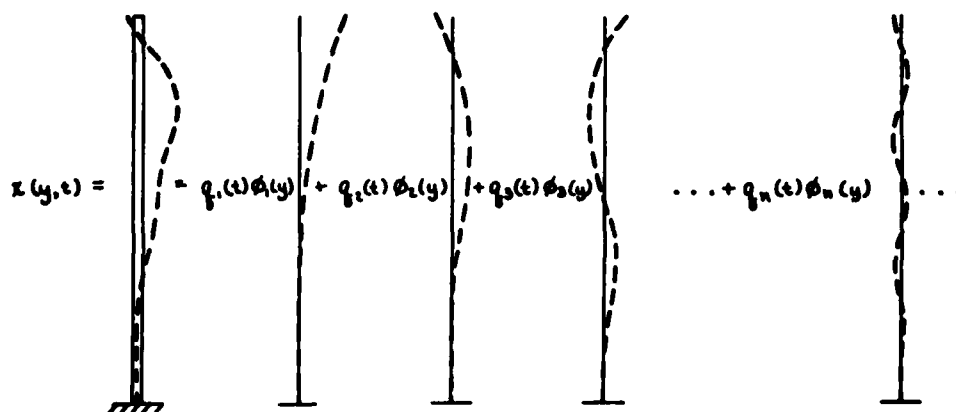


Fig. 3 - Fundamental Modal Behavior

		Excitation					
		1	2	3	4	5	6....
Response	1	A11	A12	A13	A14	A15	A1...
	2	A21	A22			A25	
	3	A31	A32				
	4		...	.	.	.	.
	..	..	..	..	..	..	

>>>> A_{ij} = A_{ji} <<<<

Fig. 4 - Reciprocity

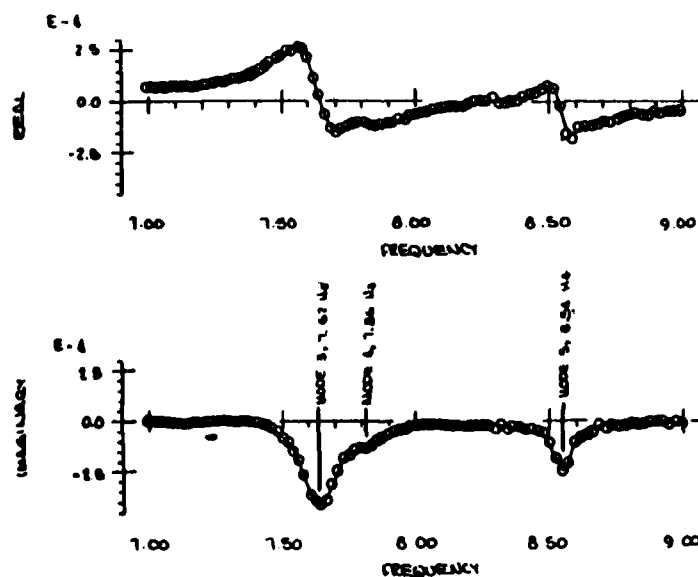


Fig. 5 - Curve Fit Results from a Modal Test on a Linear Structure

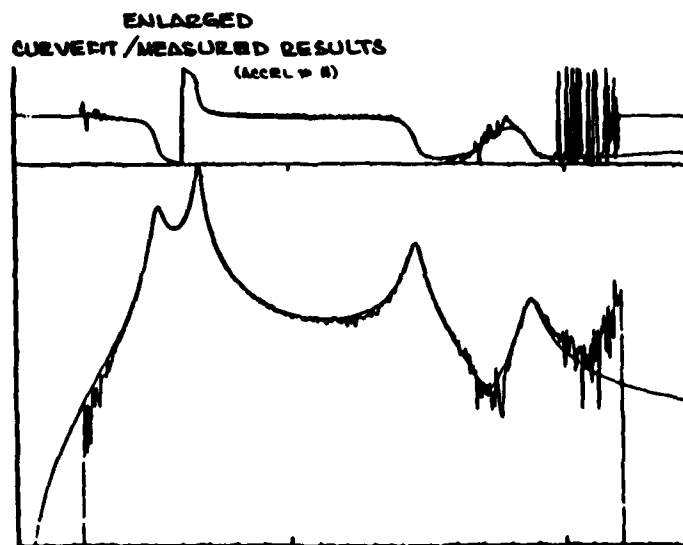


Fig. 6 - Enlarged Curve Fit Measured Results

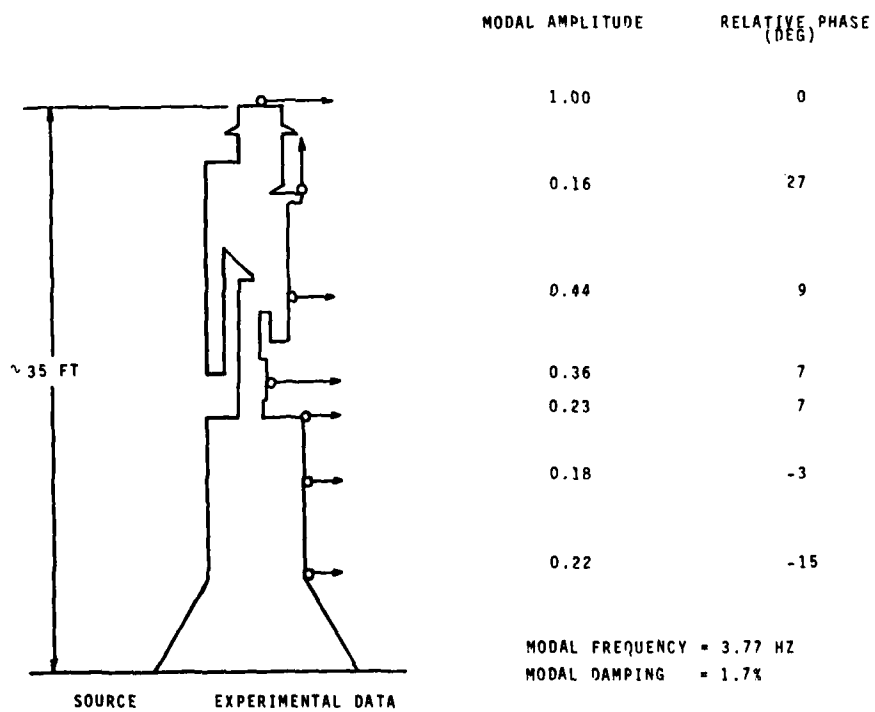


Fig. 7 - Example of Phase Angle Variations for Structural Modes

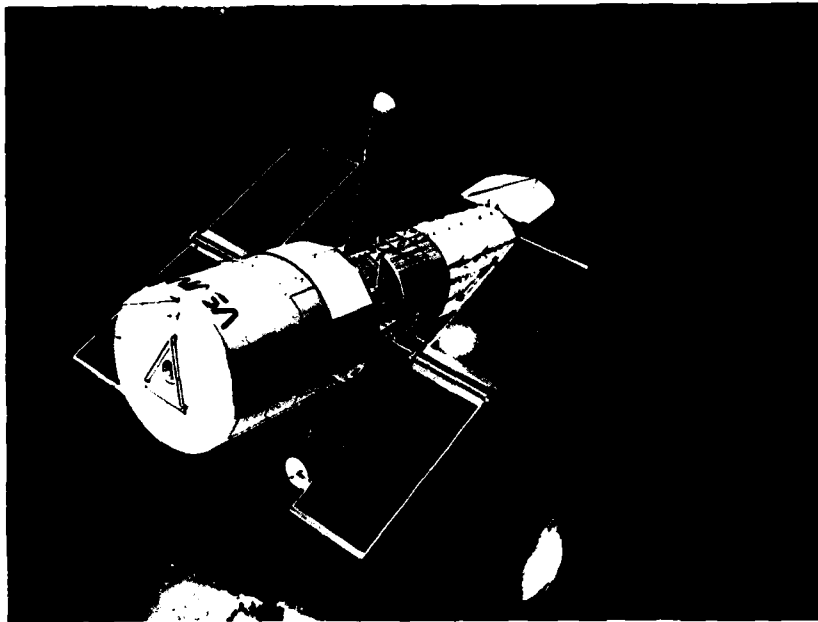


Fig. 8 — Space Telescope

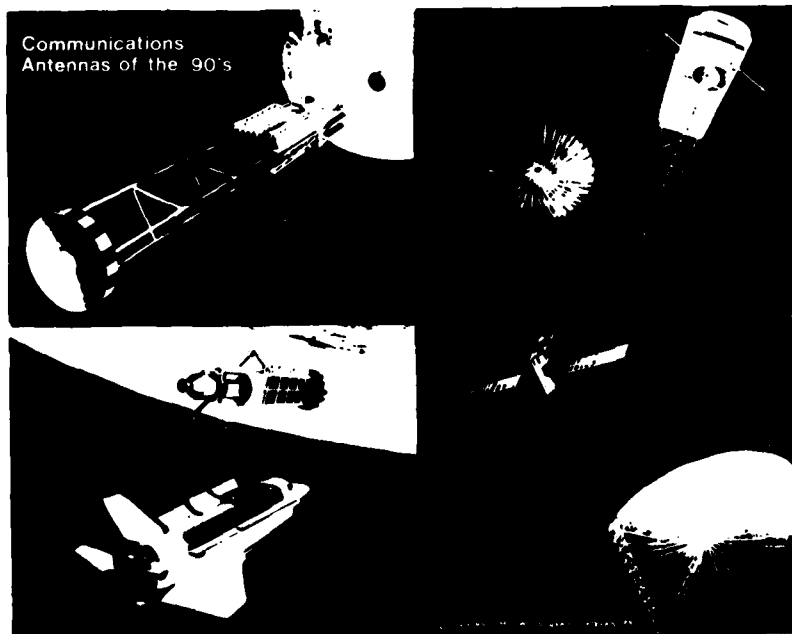


Fig. 9 — Space Telescope Antenna

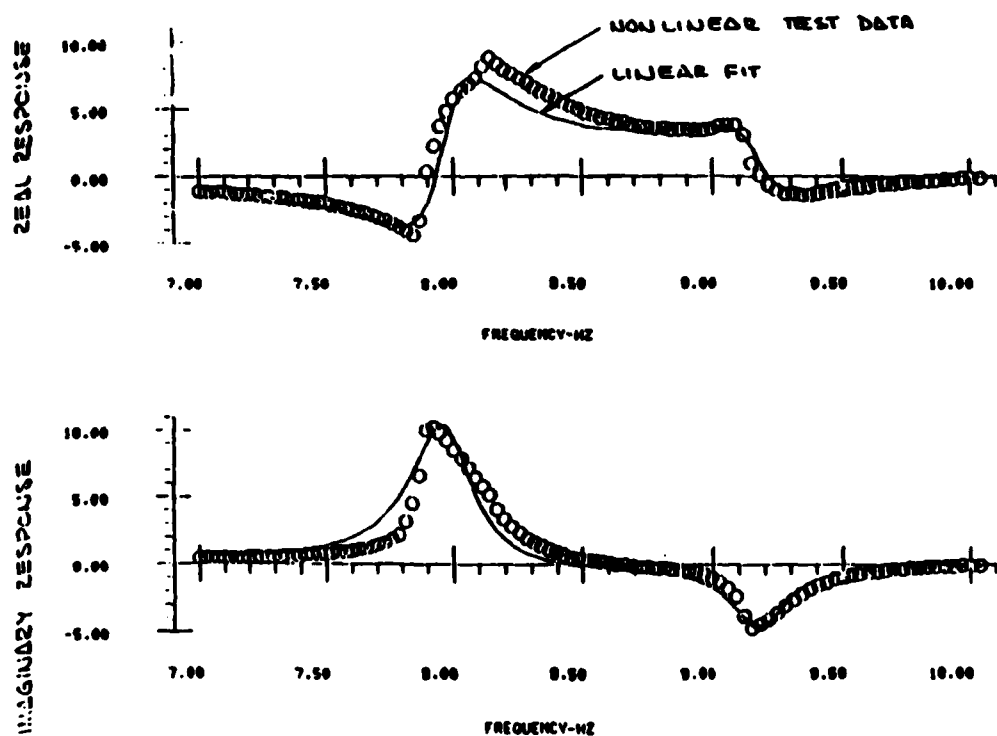


Fig. 10 — Curve Fit Results from a Modal Test on a Nonlinear Structure

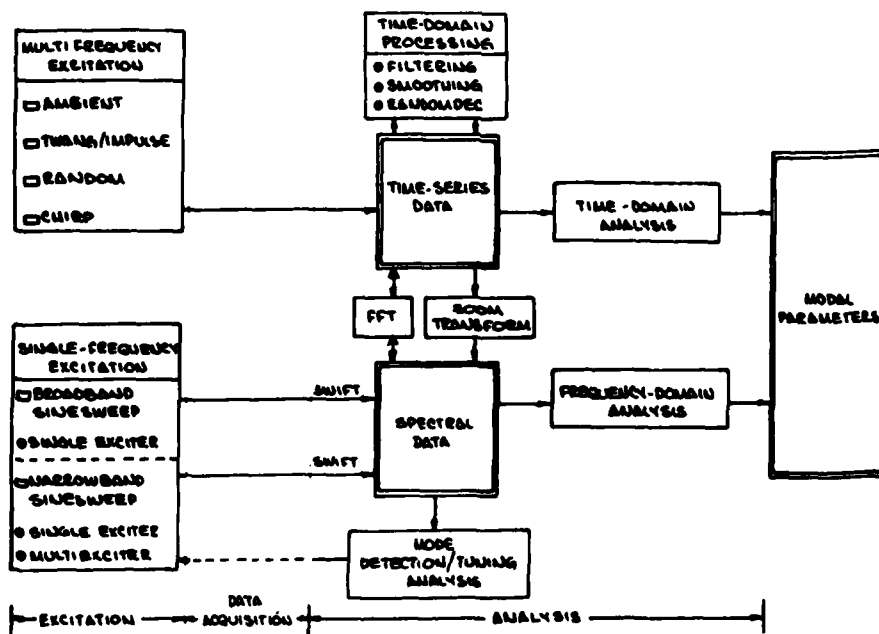


Fig. 11 — Modal Testing Methods

SOLUTIONS TO STRUCTURAL DYNAMICS PROBLEMS

G. Morosow
Martin Marietta Corporation
Denver, Colorado

My subject for today is "solutions to structural dynamics problems." The subject is controversial because it deals with philosophy, and therefore, with opinions. I suspect there will be a fair amount of discussion, so I have reserved ample time for this after my presentation.

To introduce the subject, I will set up the hypothetical situation of an interview for an engineering opening in a small but energetic aerospace company. The company is in transition from a more conventional spacecraft business to one opening new frontiers. These frontiers embrace a variety of new system feasibility studies for advanced space structures.

Mr. Charlie Bigwig is supervisor of the dynamics department. He is trying to update his capability in the analysis area to be able to face the challenges of future projects. He has ample analytical tools for handling the everyday problems involving spacecraft, but he is worried about future projects and types of analyses to support them. His applicant is a mature engineer with seven years of experience named Joe Databank.

"Good morning, Mr. Databank, please sit down and make yourself comfortable," says Charlie Bigwig. "As I explained in my letter, I need someone experienced enough to take care of our expanding business. We have, at this point in time, a series of fixed or "black-box" codes that do a marvelous job for our present projects. What we see coming in the near future, however, is beyond our capability both technically and in terms of our computer programs. For example, I see some nonlinear antenna deployment studies, retargeting of large space antenna, very-fine and precise structural jitter control studies, design of highly damped structures, large space structure nonlinear response studies, and so on. So tell me a little bit about your experience and what you think you can do for us."

Your applicant, Mr. Databank, sits down slowly, folding his large frame into the rickety chair standing in the corner of Bigwig's office. "My experience," he begins, "consists mainly of working on two large programs having

to do with the design of fighter aircraft, including trajectory analysis of an ejection seat. Our tasks, for the most part, were well defined; something that has been done for several years. In addition to our own programs, we used several nationally known commercial programs, including NASTRAN, to do the modeling of structural components. The closest I came to the type of tasks that you mentioned was during my work on the ejection seat. I developed a rigid-body analysis and FORTRAN program to handle the case."

"In terms of your future work," Mr. Databank continues, "I strongly believe that effort is needed to develop some kind of programming capability that would maintain flexibility -- i.e., one that would be able to adapt to different problems, be modular, and would perhaps use miniature programs or subroutines coupled with an executive language. This approach gives flexibility and fast response to new and unusual problems."

Charlie Bigwig puffs on his fat cigar and stares at the ceiling for a long time. "Do you think a system could be developed at a reasonable cost and in a reasonable time to handle most of our anticipated problems?" he asks with a monotonous voice designed to disguise his keen interest. "Yes, I firmly believe some sort of system can be designed," replies Mr. Databank. "After all, NASTRAN has a D-MAP modifier; there are other programs that have matrix algebra subroutines. The system I would spec would contain (1) matrix algebra abstraction subroutines, and (2) special operations on matrixes that are not normal matrix operations but might be extremely helpful, such as element multiplier."

"Then I would have a number of special routines, all compatible with matrix routines, such as numerical integration algorithms, mass properties generators, element stiffness generators, eigenvalue solutions, etc. the last one, the eigenvalue solver (or solvers) are really small, single-purpose programs, but are compatible with the rest of the system, and therefore could be considered as king-sized subroutines."

Mr. Bigwig's eyes light up. "Now you are talking! When can you start work?"

At this point, we will terminate the interview. We have established that if we seriously consider what has been said, there really might be a good reason to have this type of system. Some of you probably think, what is he trying to tell us? We have been using this type of approach for years!

But perhaps your system is not as "perfect" as the one I will define. So you may want to take home a few tidbits in the way of new ideas.

From the interview, it is apparent that a modular system is not a cure for all ills, but is, nevertheless, an extremely convenient medication to have available.

Let us examine this type of system in more detail. I remember back in the 1950's, FORTRAN was unknown, and programming of electronic digital computers by engineers was not only discouraged, it was virtually forbidden. A typical digital computer of that time was an installation that was only rivaled by the distribution control room of a central power station. The heat generated by this monster required special air conditioning, and its capability was no better than today's personal computer. Programming in the 1950's was a black art, and only a privileged few belonged to the club. It was a strict no-no for an engineer to attempt to program using the computer language, whatever it was at that time. FORTRAN was not recognized as a universal language. Normally, engineers would define the problem, submit it to the programming department, and a few weeks or months later they would get the program. Meanwhile, invariably some changes would occur, and the program would have to be updated.

We engineers finally got tired of this and submitted a dozen short programs involving matrix algebra. Each was a program in itself. Then, we convinced the programming staff to develop for us an executive pseudolanguage that would provide a continuity between the subroutines. It would call individual programs and couple them. This was the beginning of structured programming. We broke through the barrier of closed-door programming in a climate where programming by engineers was virtually impossible.

As time passed and FORTRAN became available, we began setting up a complete system in FORTRAN. The system consisted of a number of subroutines from basic matrix operations to fairly complex subroutines or miniprograms. We also decided that it would be advantageous to have a smart executive system do all the drudgery of housekeeping chores. Because of the complexity of such a system, we finally settled for use of subroutines in terms of call statements, using FORTRAN commands to provide communication between call statements when necessary. The name of the system is FORMA

(Figure 1).

In FORMA, each subroutine has a series of arguments defining size of matrixes, names, etc. There are three basic categories. The first is matrix algebra, the second is special routines that represent small single-purpose programs, and the third is housekeeping subroutines like listing, writing, plotting, etc. These subroutines were developed, through the years, as need occurred. They were coded by engineers on the job. Later, these routines were incorporated into the system, put on tape, and locked so they could not be modified without authorization. If a person needed a modified subroutine, he could copy it, modify it, give it a different name, and use it to his heart's content. If it showed a general usefulness, it would be incorporated into the FORMA system. For example, there are a great number of modal (eigenvalue) subroutines, MODEL, 1A, etc., that originated this way.

The system grew and developed in response to needs that existed at the time. For that reason, it had to be a simple, inexpensive system. There never was enough money to sit down and plan in totality a Cadillac-type system. FORMA's biggest claim to fame has been its versatility, but a set of subroutines by themselves will not do. These are the tools to execute the commands of the analyst, but what about commands by themselves? It takes a certain breed of analyst to really use the system to its limit. I believe quite strongly that any new analyses should be done in an exploratory way. That is, set up the simplest possible problem and use approximations, but be sure they are realistic. If three degrees of freedom are not sufficient to describe the situation, use six. But not 600! More degrees of freedom do not necessarily guarantee a better model. Incidentally, most of our problems executed on the FORMA system do not exceed 100 degrees of freedom. However, the system has large-degrees-of-freedom capability by using so-called partition logic.

I would be remiss if I did not mention fixed or black-box program techniques or at least compare to modular approach. If I take several large finite-element programs, in a general sense they all exhibit certain similar characteristics. Their usefulness lies in their ability to do "standard" problems quickly and efficiently. Their disadvantage is their lack of transparency, or their inability or difficulty to add modules. Therefore, one has to work the problem with whatever modules are available.

One more important item is checks. If one understands the equations in the problem, it is relatively easy to develop a continuous check through the problem. In a modular approach, these checks should be easy to implement. If a fixed program does not offer checks, there is not much one can do.

Interestingly, the modular and fixed programs can be related to the structure of various languages. For example, the Latin alphabet, which forms a basis for a great many modern languages, makes use of the concept that each sound can be expressed by one letter, or at the most, by a group of two or three letters. If one analyzes these sounds, one comes to the conclusion that (depending on the language) there are approximately 30 different sounds that make up all the words in the dictionary. In the English language, some sounds result in using two or more letters together, therefore reducing the number of letters to 26. Some languages go the other way and have a letter for each sound. We can equate it to a building-block approach. Twenty-six to thirty subroutines that generate a specific sound and are written in a serial order are all that is necessary to communicate.

Now, let us look at a converse situation. Some of the Asian languages take another approach. A single, fairly complex symbol represents a word or concept. Sometimes more than one symbol is required to describe the concept. An example is the Chinese language. Figure 2 shows a comparison. Six Chinese characters versus 13 different Latin characters; obviously, a clear advantage for the Chinese language for a short message. If one considers a long article, in Chinese, one may easily use 2,000 to 3,000 symbols. In English one uses 26 characters and no more (Figure 3).

Well, that is 26 versus 2,000 symbols. Each symbol represents a program, a concept compared to an alphabet or a building-block approach that represents something much more basic, a sound.

I do not intend to compare languages in terms of efficiency, or speed of communications. The point is that there are at least two entirely different approaches for achieving the same objective.

The same statement can be made when we talk about computer programs. The programs are a means of communication between the analyst and the computer. Most of us are used to more or less "special-purpose" black-box programs that have been used as a tool for a variety of analyses. The reason I said "more or less special-purpose programs" is that some of the black-box codes are called general-purpose programs. This means they are capable of performing a number of related analyses. These programs can not, by any means, be called general-purpose programs. The only practical way to program something new and unusual is to use the building-block approach, unless one wants to invest a considerable amount of time and money to develop a special-purpose program.

The 1970's and early 1980's can be considered a time of drastic changes, new developments, and significant technological advancements. There is appearing on the horizon

a class of new problems that were completely unknown only a few years ago. Some of these problems may be characterized by large space structures with their deployment and control problems, requirements for highly damped structures, nonlinear structures (membrane), and structure-fluid interaction in propulsion tanks at zero-g.

To make it easy for you to see the advantages that the modular approach provides in many cases, I have prepared a short, simple example. You are invited to see how you would handle it on your system.

A simple system with two degrees of freedom is a rigid bar on two springs (Figure 4). To make it really easy, we will stay with statics only. Suppose we would like to write the equilibrium equation between applied forces F and ensuing displacements x . First, we write the relation between the reaction forces R and displacements:

$$\begin{matrix} R_1 &= & k_1 x_1 \\ R_2 &= & k_2 x_2 \end{matrix} \quad (1)$$

or in a matrix form

$$\begin{bmatrix} R_1 \\ R_2 \end{bmatrix} = \begin{bmatrix} k_1 & 0 \\ 0 & k_2 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix}$$

or:

$$\{R\} = [K] \{X\} \quad (2)$$

Next, we develop a transformation matrix between the reaction forces R and applied forces F .

$$\begin{matrix} R_1 &= & \frac{b}{1} F_1 + \frac{d}{1} F_2 \\ R_2 &= & \frac{a}{1} F_1 + \frac{c}{1} F_2 \end{matrix}$$

or in matrix form

$$\begin{bmatrix} R_1 \\ R_2 \end{bmatrix} = \begin{bmatrix} b/1 & d/1 \\ a/1 & c/1 \end{bmatrix} \begin{bmatrix} F_1 \\ F_2 \end{bmatrix} \quad (3)$$

or

$$\{R\} = [T] \{F\} \quad (4)$$

Now combining (2) and (4) we have

$$[K] \{X\} = [T] \{F\} \quad (5)$$

or

$$\{X\} = [K]^{-1} [T] \{F\} \quad (6)$$

Therefore, to compute deflections at the spring locations, we have a matrix equation involving a matrix product and an inverse. Let us see how a computer program would look like:

Deflections due to Static Forces

```
Call READ (K)
Call READ (F)
Call READ (T)
Call INVL (K,E)
Call MULT (E,T,G)
Call MULT (G,F,x)
Call WRITE (x)
```

The last letter of arguments containing two or three letters designates a matrix which is the result of the operation performed by the subroutines. The entire program consists only of call statements, nothing else. Now suppose we want to change loading from F_1 , F_2 locations to P_1 , P_2 locations. All we have to do is to express "F"'s in terms of "P"'s, which becomes:

$$\begin{Bmatrix} F \end{Bmatrix} = [T_2] \begin{Bmatrix} P \end{Bmatrix} \quad (7)$$

and we have

$$\begin{Bmatrix} X \end{Bmatrix} = [K]^{-1} [T] [T_2] \begin{Bmatrix} P \end{Bmatrix} \quad (8)$$

The modified program reads:

```
Call READ (K)
Call READ (P)
Call READ (T)
Call READ (T2)
Call INVL (K,E)
Call MULT (E,T,G)
Call MULT (G,T2,M)
Call MULT (M,P,x)
Call WRITE (x)
```

As you can see, it is not difficult to change the problem. Also, the ability to do checks is extremely important. For example, one of the checks might be $K^{-1}E = I$. All that is required is to add the following statement:

```
Call MULT (E,K,I)
```

Now, how would you do this problem, if only fixed programs exist? Most likely, you would go through your library and find one that best fits the situation. Then, you would fit the problem, perhaps modify it, and then interpret the results accordingly. In other words, you force the problem to fit the tool at hand. You are constrained in your attempt to do that. Now, conversely, in modular approach, you develop the tool to fit the problem -- infinitely more freedom!

In the future, there will exist complex problems that will not fit any of the so-called general-purpose programs. No amount of shoehorn squeezing will fit the problem to your tools.

The best world of all is the one where both fixed and modular tools are available. There is no reason why these may not have common interfaces and why the fixed programs could not be considered as giant subroutines.

AABB	CHKZER	INV1	MULT	READ	TERP1	UMAM1
AAPBB	CKMAS1	INV1A	MULTA	READIM	TERP2	UNITY
ALOD1	CKSTF1	INV1B	MULTB	READO	TERP3	
ALOD2		INV1NP	MULTAD	REORD		
ALPHAA	COLMLT	INV2	ONES	REVADD	TIMCHK	VCROSS
ASSEM	COMENT	INV3	ONRBM	REVISE	TRAE2	VDOT
ATXBAD	COMPAR	INV3A		RNUM1		VM1
ATXBA1		INV4	PAGEHD	ROWMLT	TRANS	VMTR1
ATXBB		INV5	PDCOM	RTAPE	TRANSA	
ATXBB1		INV6	PLOT1		TRMM	WRITE
ATXBB2	DCOM1			SELADD	TRSP1	WRITIM
AXB	DCOM2	LTAPE	PLOT2	SELECT		WTAPE
AXBA1	DCOM2X	MASS1	PLOT3	SIGMA	TRSP1A	XLORD
AXBA2	DIAG	MASS1A		SMEQ1		
AXBA3	DIFF1	MASS2	PLOT4	SMEQ2	TRSP1B	ZERO
	DIFF2	MASS2A		SRED1	TRSP1C	ZEROLH
	DISA	MASS2B	PLOTSS	SRED2	TRSP2	ZEROUH
	DTOZ	MODE1	PUNCH	SRED3	TRSP3	ZZBOMB
BABT		MODE1A	PUNCHD	SRED4		
BAB7A	EIGN1	MODE1B		SRESP3		
BAB71	EIGN1A	MODE1X	RBTG1			
BABTA1	EXPON	MODE2	RBTG2	START	TRSP2A	
BSOL2X	FR1		RBTG3	START2		
BTAB	FRAE1		RBTG4	STIF1		
BTABA		MODE2A	RBTTAB	STIF2		
BTAB1				SYMLH		
BTAB2				SYMUH		
BTAB2X						
BTABA1						
BTABA2						

Fig. 1 — FORMA Subroutines

快樂 } HAPPINESS
 及快樂 } AND
 之追求 } HAPPINESS-
 求 } HOW TO PURSUE

Fig. 2 — Equivalents in Chinese

WHERE IS THE REAL LITERATURE ON AIRBLAST AND GROUND SHOCK?

W. E. Baker
Southwest Research Institute
San Antonio, Texas

"The pursuit of truth will make you free -
even if you never catch up with it."

Clarence Darrow

Introduction

When Rudy Volin called me to ask if I would give this paper, he suggested that I could repeat a paper given at an Air Force Symposium this past May. The theme of that paper was that a limited library of broad, unclassified references was very useful to engineers or scientists engaged in studies of nonnuclear weapons effects. In agreeing to present this paper, I promised that I would not repeat the earlier paper, but would instead present some other data and opinions on the sources of literature for study of airblast, ground shock and their effects. The present paper is similar because I limited my reference list to relatively few broad, English language sources, and to unclassified sources.

What are the Potential Literature Sources?

Our topics are airblast, ground shock, and the effects of these manifestations of explosions. The potential literature sources fall into four general classes:

- Books
- Periodicals
- Technical reports
- Proceedings of symposia & colloquia

Books can include those published by many commercial publishers, and those published by government printing offices, most notably the U.S. Government Printing Office. Periodicals can include peer review journals, other engineering and technical society periodicals, industrial periodicals, and some government publications (Shock and Vibration Digest is an example). Technical reports in this field are issued by a number of U.S. and foreign agencies. The Department of Defense agencies which are the best sources for reports on airblast and ground shock are the tri-service Defense Nuclear Agency; The U.S. Army

Laboratories: ARADCOM Ballistic Research Laboratory, Waterways Experiment Station; The U.S. Navy Laboratories: Naval Surface Weapons Center (both White Oak and Dahlgren Laboratories), Naval Weapons Center, and Naval Civil Engineering Laboratory; and the U.S. Air Force Laboratories: Air Force Weapons Laboratory, Air Force Armament Laboratory, and Air Force Engineering and Services Laboratory. Many technical reports in this field are prepared by contractors from industry and academia, and they usually appear as Contractor Reports distributed by the appropriate DOD agencies.

Department of Energy agencies are also fruitful sources of the report literature in this field. The most prolific are Sandia National Laboratories, Los Alamos National Laboratory, and Lawrence Livermore National Laboratory.

The other U.S. Government Agencies who generate an extensive report literature, some of which is quite useful in airblast and ground shock studies, are NASA and the Bureau of Mines. NASA also contracts a number of pertinent studies, and publishes results in Contractor Reports.

None of the foreign laboratories or agencies publishes as extensively on these topics as the U.S. agencies, and of course, many of their reports are written in foreign languages. We have found the best sources there to be Royal Armament Research and Development Establishment in Great Britain, Norwegian Defence Construction Service in Norway, National Defence Research Institute and Royal Fortifications Administration in Sweden, Technological Laboratory TNO in the Netherlands, Ernst Mach Institute in West Germany, and Ernst Basler & Partner in Switzerland.

The proceedings of symposia and colloquia which contain useful literature on this topic include those which recur on a regular basis, and those which are offered once or perhaps as a

series of several on the source topic. The former are most useful, and include the minutes of the biennial Department of Defense Explosives Safety Board Safety Seminars, the ballistics symposia sponsored by the American Defense Preparedness Association, and last but not least, the Bulletins of the Shock and Vibration Symposia, such as the one you are attending.

How do we classify the various literature sources?

I'm not trying to confuse security officers by the word "classsify." Instead, I'm using the dictionary definition of assigning references to a category.

The classes for separating the various literature sources are only two: 1) the open literature, and 2) all the rest. I will somewhat arbitrarily call the second class the report literature, because it is in number and content dominated by technical reports.

When one considers the various literature sources I've discussed, it is sometimes easy to assign a reference to one of these two classes, and sometimes difficult. As an example, NASA publishes a variety of documents, some of which are clearly open literature like their Special Publication (SP) bound books, and some of which are clearly very limited reports like Technical Notes. But, their Contractor Reports are considered to be open literature if they are low-numbered CR's, and not open literature if they are high-numbered CR's. This may seem to be a petty distinction, but over 3,000 copies of low-numbered CR's are printed and distributed, while only perhaps a hundred copies printed for high-numbered CR's. To make the classification clear cut, I have assumed that all publications readily available for purchase without ordering them through the National Technical Information Service (NTIS) are open literature. This includes a large number of government publications advertised by the Government Printing Office. Conversely, any publication which must be ordered through NTIS is consigned to the report literature. This includes all NASA CR's, even if NASA claims the low-numbered ones are open literature.

Even so, my classification of the open literature is probably much more liberal than that of many of my university friends. The prevalent attitude there is often "If it doesn't appear in a peer-review journal, it does not belong in the literature." I disagree with that viewpoint.

Survey of Literature Cited in some General References on Airblast and Ground Shock.

For this survey, I chose twelve references. They are listed in the reference list at the end of the paper, and some data and statistics regarding literature citation in these references appear in the table.

All but two of the twelve references are themselves open literature by my definition (Refs. 1 and 6 are voluminous reports), but one can see from the table that most of them rely very heavily on the report literature for their material. The topics covered in the references include many aspects of airblast, ground shock and their effects; dynamic response of structures to airblast, ground shock and impact; theory and experiment in shock waves and airblast; theory and practice of dynamic scale modeling; and theory and practice of dynamic impact. Classification by type of reference is shown in the table. Authors' affiliations represent a spectrum from industry, universities and government.

Although all of the twelve references are broad ones, the thoroughness of referencing varies widely, from a minimum of 14 citations for Ref. 4 to a maximum of 779 citations for Ref. 1. Reliance on open literature versus report literature also varies widely, with Refs. 1 and 7 being the extremes. Because Ref. 1 is a summary report of World War II research on effects of impacts and explosions, it cites the report literature almost exclusively (748 of 779 citations). On the other hand Ref. 7 shows the strong preference of its (university) author by primarily citing references from peer review journals (103 of 130 citations). But, not that all references included a mix of open and report literature citations. Citations in government reports and books were weighted toward the report literature.

Before starting this survey, I already knew that the report literature was as essential as access to the open literature, for the title topic. I also thought that I would be able to show a strong bias of university authors toward the open literature, as opposed the report literature. Instead, I found that the bias may be more a matter of personal preference and training. While Prof. Oppenheim in Ref. 7 indeed leaned strongly toward citations in peer review journals, Profs. Courant and Freidrichs in Ref. 2 used nearly as many report references as open references in a somewhat similar topic. In comparing citations in Refs. 8 and 9, which are truly on the same topic, we see that two authors (or set of authors) from industry slanted their reliances on report versus open literature quite differently. But, perhaps the clearest indication of preference of an author occurs in the citations in Ref. 12. The table shows that the literature citations in this reference are extensive, and that there is a rather strong emphasis on open literature citations. A chapter-by-chapter study of this reference reveals that the majority of the open literature citations are listed by only one of five coauthors, Dr. Ted Nicholas. Eliminating his citations would both drastically reduce the literature cited in the book, and would shift the emphasis from the open literature to the report literature.

Discussion

It would be foolish to draw any sweeping conclusions from this brief survey. But, some points are evident. If your field of interest is airblast, ground shock, or their effects on a variety of "targets," then, as for any other technical topic, you must survey the literature on these topics to avoid repeating the successes and failures of the last (my) generation. If you are in a university or strongly academically oriented, you will probably gravitate toward a review of the open literature. Do not ignore the report literature, because that is where a lot of the action is and has been. If you do not know that NTIS exists and what the initials stand for, you are missing at least half of the pertinent literature.

If you are an engineer or scientist working in one of our excellent government ballistic or ordnance laboratories, you may have a strong tendency to read, use and cite only reports from your own and sister laboratories. Again, do not ignore the open literature because you can often find much pertinent work reported there. I know that the internal report review and printing process is often very lengthy and traumatic in your laboratory, and there is seldom much management incentive there to have that same work published in the open literature. But, my personal opinion is that you should make that effort, even if (horrors!) you have to write the papers on your own time.

For my cohorts in industry, I suspect that, because you are still doing business, you have already learned that you must review both the open literature and the report literature, both in the U.S. and abroad, to be at least reasonably sure that you have discovered most of the pertinent work in airblast and ground shock. Keep up the contacts and don't throw away all of the technical reports that may automatically come your way through distribution lists.

References

1. M. P. White (Editor) (1946), "Effects of Impact and Explosion, Summary Technical Report of Division 2, National Defense Research Council, Vol. 1, Washington, D.C., AD 221-586.
2. R. Courant & K. O. Friedrichs, (1948), Supersonic Flow & Shock Waves, Interscience Pub., NY.
3. C. H. Norris, R. J. Hansen, M. J. Holley, Jr., J. M. Biggs, S. Namyet, & J. K. Minami (1959), Structural Design for Dynamic Loads, McGraw-Hill Book Co., NY.
4. G. F. Kinney, (1962) Explosive Shocks in Air, MacMillan, NY.
5. W. E. Baker (1973), Explosions in Air, Univ. of Texas Press, Austin, Texas.
6. M. M. Swisdak, Jr., (1975) "Explosion Effects and Properties. Part I - Explosion Effects in Air," NSWC/WOL/TR 75-116, Naval Surface Weapons Center, White Oak, Silver Spring, MD, 9 Oct. 1975.
7. A. K. Oppenheim (1970), "Introduction to Gasdynamics of Explosions," Udine, Springer-Verlag, NY.
8. W. E. Baker, P. S. Westine & F. T. Dodge, (1973) Similarity Methods in Engineering Dynamics, Hayden Books, New Rochelle, NJ.
9. D. J. Schuring, (1977) Scale Models in Engineering, Pergamon Press, NY.
10. S. Glasstone & P. J. Dolan, (Editors) (1977), The Effects of Nuclear Weapons, Third Edition, U.S. Government Printing Office.
11. F. T. Bodurtha (1980), Industrial Explosion Prevention and Protection, McGraw-Hill Book Co., NY.
12. J. A. Zukas, T. Nicholas, H. E. Swift, L. B. Greszczuk, & D. R. Curran, (1982), Impact Dynamics, Wiley-Interscience, NY.

MIL-STD-810D

TAILORING INITIATIVES FOR MIL-STD-810D ENVIRONMENTAL TEST METHODS AND ENGINEERING GUIDELINES

David L. Earls
Air Force Wright Aeronautical Laboratories
Wright-Patterson Air Force Base, Ohio

Previous editions of MIL-STD-810 emphasized environmental qualification tests conducted at worldwide climatic and dynamic environmental extremes. The tests were essentially rigid worst case requirements, presented in a cookbook style, offering few alternatives for individual applications. In contrast, new MIL-STD-810D provides engineering tasks to determine life cycle environmental histories of equipment so that tests can be formulated and tailored to the individual equipment applications. The engineering tasks, leading to more realistic testing, include the development of an overall Environmental Management Plan; a Life Cycle Environmental Profile, including environmental conditions for an equipment from its release from manufacturing to its retirement from use; Environmental Design Criteria and Test Plan; and Operational Verification Plan. These engineering tasks, documented and applied by Data Item Descriptions (DIDs), provide the data for developing and tailoring individual environmental tests. MIL-STD-810D also aids in selecting tailored environmental tests by providing environmental test criteria, rationale, and background as a new, separate section of each test method. MIL-STD-810D, by means of the new engineering tasks, implemented by DIDs, effectively defines actual environmental conditions as encountered in the real world and bridges the gap between environmental criteria and environmental tests.

INTRODUCTION

MIL-STD-810 was initiated and published as an Air Force document in 1962. Subsequently, the other two military services, Army and Navy, made their technical contributions to the standard, principally incorporating tests peculiar to their needs, and MIL-STD-810A was published as a Tri-Service coordinated military standard in 1964. Additional technical advancements in environmental testing techniques were incorporated, resulting in 20 natural and dynamic test methods in MIL-STD-810C. These test methods were based upon environmental extremes in order to restrict practical laboratory test time to a minimum compared to the years of service that an equipment would experience in the actual environment. Innovative technical approaches to tailoring were introduced into the dynamics test methods of MIL-STD-810C, where engineering parameters and calculations were required to arrive at test levels. These tailoring concepts were utilized in the random vibration, acoustic noise and gunfire test methods.

In recent years, DoD has placed increased emphasis on environmental specifications and standards because of their broad application to all forms of military hardware. This broad utilization results in a significant overall cost. Changes in the way the specifications

are used and applied have therefore been investigated with the intention of making them more cost effective. One of the more significant means of achieving cost effectiveness has been to increase the practice of tailoring to individual applications in order to avoid blanket use of standards.

The tailoring approach permits selective application of realistic field environments in the laboratory and prevents over testing or under testing, both of which are costly.

TRANSITION IN METHODOLOGY

MIL-STD-810D has been completely transformed from previous editions; it is essentially a new environmental testing standard. Most of the previous individual test methods were very rigid, applying a step-by-step test procedure, normally with only one maximum environmental stress condition, which was based on worldwide climatic extremes or maximum dynamic measurements. It was strictly a cookbook style document. No alternatives to the specified test conditions were offered, and no rationale or explanation was given. This led to a lack of credibility or confidence in the test conditions, since they were often found to be inappropriate for specific equipment. In the

drive for increasingly cost effective tests, it became apparent that a new approach to environmental testing standardization was urgently needed. It was decided to formulate new General Requirements and restructure the individual test methods to incorporate tailorable environmental criteria and guidance for their application.

SCOPE AND GENERAL REQUIREMENTS

Tailoring, as applied to MIL-STD-810D, is the process of choosing or altering test procedures, conditions, values, tolerances, and measures of failure to simulate or exaggerate the effects of one or more environmental condition which an item will be subjected to during its life cycle. Broadly speaking, it also includes the engineering tasks and preparation of planning documents to assure proper consideration of environments throughout the life cycle. This concept sounds good in theory, but presents considerable work in practice. Much of the MIL-STD-810D engineering effort over the past few years has been directed toward reducing this tailoring concept to practical application. This requires precise information on details of the actual environmental conditions to be experienced by an item throughout its useful life. This is considerably different from past ways of doing business, where environmental data were usually very general. For instance, atmospheric or natural environments were presented as extremes to be encountered worldwide. The natural environments need to be known by regions, or specific locations where a weapon system is employed, in order to be precise. Furthermore, natural environments are altered by the platform in which an item is installed; also, natural and dynamic environments are induced by the platform itself. They vary from platform to platform and also vary with location in the platform. In short, a much more detailed knowledge of precise environmental conditions with respect to the specific application is necessary. This is being accomplished by two approaches. The first approach is being done through new General Requirements; the second approach by newly structured test methods.

The new General Requirements information is essentially a series of environmental engineering tasks which can be accomplished by the procuring activity engineering offices or by the contractor. These tasks are oriented toward the large weapons system procurements, usually manned by a variety of engineering disciplines. When the Air Force makes a commitment to build a new airplane, for example, there are many plans and tasks put on contract. There is also a lot of data already available concerning mission profiles, locations to be deployed, amount of flight hours per month, design life, maintenance concepts, repair depot locations, all of which can directly relate to environmental conditions. Engineering tasks have, therefore, been developed for MIL-STD-810D which directly relate to this procurement concept and provide for development of realistic environmental conditions for

the weapon system. The following tasks implement the tailoring concept.

1. Environmental Management Plan. This plan has been established to provide overall control of the environmental program. It includes consideration of manpower requirements, scheduling, life cycle environmental conditions, test tailoring, test performance, analysis of results, corrective actions, and actual field environmental conditions. Plans for monitoring, assessing, reporting and implementing the entire environmental program are addressed.

2. Life Cycle Environmental Profile Task. This task is formulated to document the life cycle history of events and associated environmental conditions for an item of equipment from the time of its release from manufacturing to its retirement from service. Phases of the life cycle to be considered include handling, shipping and storage prior to use; phases between missions, such as standby or storage, or transfer to and from repair sites; geographical locations of expected deployment, and platform environments during and between missions. The environments and combinations of environments that an equipment will be exposed to during the various phases of its life are determined. This documented life cycle profile provides the necessary data base for establishment of detailed environmental design and test criteria.

3. Environmental Design Criteria and Test Plan. This plan defines the specific environmental design and test requirements, and includes an environmental test plan. It delineates the purpose and objective of the tests, the environmental conditions for test, test procedures and limits, test instrumentation, failure criteria and facility requirements. This plan builds on the previous ones and is an essential engineering task required to obtain effective, properly tailored environmental tests.

4. Operational Environmental Verification Plan. This task includes plans for obtaining data on actual operating or field environments for comparison with design and test criteria. Field service measured data provides the basis for analyzing the adequacy of the environmental program.

These four plans increase the opportunity for environmental engineering to take place. The old General Requirements from previous MIL-STD-810 versions stifled the use of engineering judgment, leaving project engineers with an inflexible specification with no engineering rationale for decision making. The new engineering approach of MIL-STD-810D encourages technical assessment and determination of the specific environmental conditions applicable to the item being purchased. It also takes into technical consideration the interaction of a weapon system operating in the environment. It is necessary to consider the weapon system (referred to as the platform in MIL-STD-810D) and its effect in increasing or decreasing the environ-

mental response at installed equipment locations.

The four General Requirement plans are conveniently put on contract, when desired, by Data Item Descriptions (DIDs). A government program manager has the prerogative of selecting those plans or tasks that he considers appropriate for his procurement. He may have a mix of in-house engineering available which may do one or more of the engineering tasks, and then he simply implements the rest by DIDs, which form a part of the contract. For example, some life cycle environmental information may be available from in-house engineering studies conducted prior to contract award. Also, operational environmental data needed to verify design and test criteria can often be obtained from government test activities, such as the Air Force Flight Test Center or the Armament Test Center.

INDIVIDUAL ENVIRONMENTAL TEST METHODS

The bulk of MIL-STD-810D resides in the individual test methods. They have been totally reorganized, and new additional technical information has been added so that realistic environmental conditions may be determined for a wide variety of applications. There are still twenty test methods, however, three have been discontinued while three more have been added. Method 504, Temperature-Altitude was dropped and is now superseded by Method 520, Temperature, Humidity, Vibration, Altitude. The combined environment Method 520 also replaces former Method 518, Temperature, Humidity, Altitude. Method 517, Space Simulation, was discontinued, as this type of testing is now governed by military standards covering space applications. A new Method 521, Freezing Rain, has been added. Also, a new Method 523, Vibro-Acoustic, Temperature, is included as a test for aircraft external stores.

A major contribution to each of the MIL-STD-810D test methods has been the inclusion of new technical environmental material with guidance for tailoring it to a particular test requirement. Each test method has been divided into Section I and Section II. The first section includes major subheadings, as follows:

- PURPOSE
- ENVIRONMENTAL EFFECTS
- GUIDELINES FOR DETERMINING TEST PROCEDURES AND TEST CONDITIONS
- SPECIAL CONSIDERATIONS
- REFERENCES

A quick glance at the purpose of the test will often alert a technical person responsible for a particular procurement as to whether he needs to consider utilizing the test method. The environmental effects subheading is intended to show how a particular environmental condition will adversely affect military hardware. It shows effects that may occur as a result of exposure to the particular environment under

consideration. The Guidelines for Determining Test Procedures and Test Conditions subheading is the heart of the test method for those interested in tailoring the test to their particular procurement. It elaborates on applications of the test method, lists a variety of test procedures and explains each. For example, the shock test method lists nine procedures: functional, packaged equipment, fragility, transit drop, crash hazard, bench handling, pyrotechnic, rail impact, and catapult launch/arrested landing. It also includes the rationale and restrictions for each shock test. This subheading explains the test conditions to be used for each procedure. A project engineer can confidently pick an applicable procedure, since enough information is presented to enable him to make a rational decision. The special considerations subheading sometimes includes test interruption guidance in case of inadvertent, unscheduled breaks in test performance. Also included is such information as special facility considerations and unique failure manifestations expected. Section I of each test method ends with a list of references. This is invaluable information for further researching rationale, to understand more fully the supporting background information used to develop the test. This is sometimes necessary when more detailed engineering effort is needed to fully tailor a procedure.

Section II of each individual test method is essentially a step-by-step laboratory procedure for conducting the test. The environmental conditions, limits, and durations for the test are established from the criteria of Section I, or from the technical tasks of General Requirements, and are applied by Section II. This section is essentially directed toward the test engineer or technician who actually performs the laboratory test. It contains the following subheadings: Apparatus, Preparation for Test, Procedures, and Information To Be Recorded.

MAJOR CONTRIBUTIONS OF MIL-STD-810D

MIL-STD-810D establishes environmental engineering as a recognized technical part of the acquisition process. It establishes an orderly series of engineering tasks which can be readily applied under contract to attain optimum tailored environmental tests attuned to a specific weapon system. It also increases the credibility of testing by providing technically valid rationale and background for each test, and it facilitates the proper selection of environmental tests by the inclusion of new procedures that reflect the end use of the equipment to be tested. MIL-STD-810D is a major step forward in bridging the gap between environmental criteria and environmental testing.

ACCELERATION RESPONSES OF TYPICAL LRU'S
SUBJECTED TO BENCH HANDLING AND INSTALLATION SHOCK

H. Caruso and E. Szymkowiak
Westinghouse Electric Corporation
Baltimore, Maryland

Measurements were made on a typical LRU (Line Replaceable Unit) to determine the levels of shock associated with the bench handling edge-drop tests described in Mil-Std-810C/D. Measurements were also made on three LRU's mounted on slide rails to determine the shock resulting from typical seating operations during installation. The 4-inch drops made during the bench handling tests produced levels from 94 to 250 peak-g's with durations up to 8 milliseconds. Energy was concentrated in the 50- to 300-Hz frequency band, a region of particular importance for typical LRU structures. Installation shock pulses ranged from 6 to 16 g's with a duration of approximately 10 milliseconds, again, a region of concern for typical LRU's.

Peak responses measured for bench handling and installation shocks represent an energy input between that associated with the traditional basic design and crash safety shock tests of Mil-Std-810C. Therefore, these shock producing events should be given at least as much attention as those events that are traditionally considered, especially since bench handling and installation are far more likely to occur on a regular basis. In particular, special attention should be given to those classes of equipment which are not normally thought of as encountering significant shock or vibration environments in end-use or mission application.

INTRODUCTION

Traditional shock and vibration design criteria for electronic hardware are often based on the environmental conditions associated with its intended end-use or mission application. For example, the vibration criteria for an airborne radar system would typically be based on forcing functions and responses associated with high-subsonic, low-altitude penetration and air combat maneuver buffet, these being the most visible and dramatic mission phases. However, for some airborne and shipborne equipment, and for fixed-base ground equipment, in particular, the vibration environments experienced during end-use mission phases can be relatively benign or nonexistent.

Such an approach ignores the fact that for most hardware, a majority of its service life may be spent under circumstances that have no

direct relation to the environments associated with its intended end-use. These circumstances include shipping, storage at various levels, idle time, and troubleshooting or repair activities. The current emphasis on tailored testing begins to address this deficiency by requiring the test and design engineer to consider the environmental conditions associated with all phases of its deployment, rather than limiting consideration to only the end-use or mission phase. In support of this approach, Method 514.3 of Mil-Std-810D places considerable emphasis on defining and testing for realistic shipping and ground transport vibration scenarios.

On the other hand, bench handling and other shock producing situations associated with the movement and manipulation of hardware by support personnel have remained little used, poorly defined design and test environments. There are several reasons for this prevailing state of ignorance. First, there

is relatively little glamour associated with handling when compared with the more visible and spectacular dynamic events that occur in executing a specific mission. Perhaps more to the point, however, is the fact that the shocks resulting from these situations almost invariably involve a mistake that the responsible party would rather not publicize. In some cases, the shock results from either carelessness or ignorance on the part of an individual, a situation which assuredly will remain undocumented. In other cases, the so-called mishandling may result from a conscientious operator or handler who has the misfortune to confront an item of hardware that was "designed" by the producer to thwart any reasonable attempts at careful handling. These situations too will remain, for the most part, undocumented.

Mil-Std-810[2,3] in its various manifestations is now more than 20 years old, but the 4-inch drop bench handling test has remained relatively unchanged throughout the evolutionary process of this document. As noted by Junkers¹ in 1965, "It is doubtful if this test involves any environmental measurements. It appears, therefore, to be based on such factors as experience and apparently reasonable assumptions of shock possibilities." For the reasons cited above, this situation is unlikely to change.

However, even if a statistically satisfying description of rough handling shock circumstances remains beyond our grasp, there is no reason to remain ignorant of the resultant environmental conditions associated with these events. If we accept intuitively the assumption that a 4-inch edge drop is not an unreasonable event to occur at least several times during the lifetime of a given item of hardware, then measuring the environmental conditions that occur during such drops will contribute valuable information to the design and test tailoring process for nontrivial life-cycle events.

Another shock producing event which is not included in Mil-Std-810 and, to our knowledge, has not been formally documented, involves the installation of rail-mounted equipment. By installation, we are referring to the process in which a hardware assembly is pushed along rails or tracks to be seated within a parent enclosure or structure. If misalignment or resistance is encountered, or if operator attitude is somewhat "aggressive" on a particular day, then one should not assume that the impact awaiting the hardware at the end of the installation process will be benign.

The experiments described below contribute to characterizing these shock-producing events and raise questions concerning the adequacy of traditional design and test philosophies for environmental shock.

DESCRIPTION OF BENCH HANDLING TESTS

Test Setup

For the bench handling shock tests, the test item was an LRU (Line Replaceable Unit) representative of the large population of modular electronic hardware currently being produced throughout the industry. The LRU (figure 1) weighed 21 pounds with dimensions of 20" (length) x 8.5" (height) x 6" (width). The center-of-mass of the LRU was at the approximate center of the structure. All internal assemblies were installed with the exception of a printed wiring assembly and a fan, which were unavailable at the time of the experiment. However, the mass of the missing parts was negligible when compared with the overall mass of the LRU and their absence did not compromise the validity of the measurements.

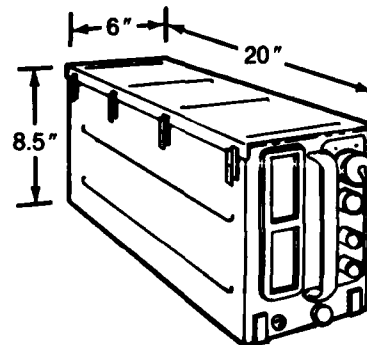


Figure 1. LRU used for bench handling shock tests

The acceleration response of the overall structure to the edge-drop shocks was measured with a vertically oriented accelerometer mounted on a structural hardpoint of the LRU frame. In each case, this location was close to the bottom of the vertical face furthest from the edge around which the LRU was pivoted (figure 2). Accelerometer responses were fed to a Hewlett-Packard 5451C Fourier Analyzer for storage and analysis. No measurements were made on internal components.

Test Procedure

The bench handling tests were conducted using the procedure described in Mil-Std-810D[3], Method 516.3, Procedure VI. Using each of the eight LRU edges (figure 3) as a pivot, the opposite edge of the LRU was lifted to a height of 4 inches or a point of balance was reached, whichever was reached first. The LRU was then allowed to fall freely on to a rigid impact surface. Additional drops were made from an edge height of 2 inches to better establish data trends.

RESPONSE ACCELEROMETER
LOCATION AND ORIENTATION

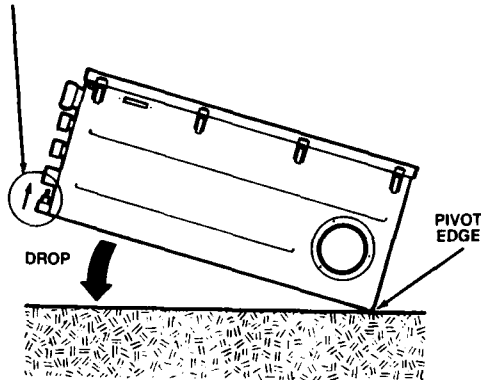


Figure 2. Accelerometer location for bench handling shock tests

Impacts were made on two of the six LRU faces, the top and bottom (Figure 3). Drops were not made on the front and rear faces because of their relatively small size (and resulting low drop height) and because of the LRU could not practically be oriented during servicing so that it would fall on either of these faces. Two different impact surfaces were used: a solid wooden slab resting on a concrete floor and solid wooden bench top surfaced with Micarta (a surfacing material similar to Formica).

BENCH HANDLING TEST RESULTS

What can one expect when an object is raised to some height and is then dropped to an unyielding surface? Raising an object to a height, h , implies acquiring potential energy proportional to that height. All of the potential energy is transformed into kinetic energy so that just before the object strikes the surface, the velocity of the object is $V = (2gh)^{0.5}$, where "g" is the acceleration of gravity. During the impact, the velocity returns to zero in some length of time dependent on the elasticities of the object

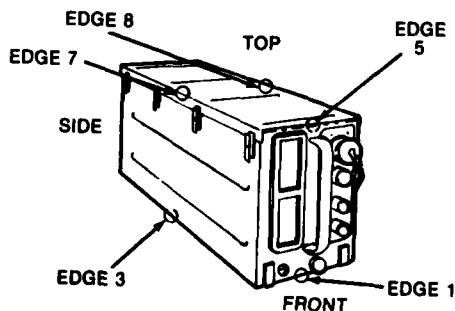


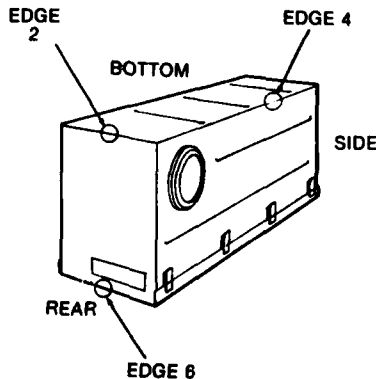
Figure 3. Pivot edge identification for bench handling shock tests

and surface. The resultant acceleration pulse height, shape and duration are functions of the energy dissipation process. In returning to rest, if the falling object is stopped in a very short time, the acceleration pulse height will be large; if stopping takes a long time, the pulse height will be small.

Bench handling tests with an LRU having the shape of a rectangular prism obviously result in a variety of drop heights and durations. For example, when pivoting around edges 1, 2, 5 or 6 (as identified in figure 3) a full 4-inch drop was appropriate. When pivoting about the long edges (3, 4, 7 or 8), the LRU is balanced with the center of gravity just within the pivot edge, so that a 4-inch drop was not appropriate in terms of the conditions described in Mil-Std-810D. Instead, with the LRU used, the drop heights were approximately 3 inches (edges 3 or 7) and 3.5 inches (edges 4 or 8).

Figure 4 shows an acceleration time history typical of 4-inch drops when pivoting about edges 1, 2, 5 or 6, with a peak height of about 245 g's, and a duration at 10 percent pulse height of about 3.3 milliseconds. Drops on the top surface of the LRU resulted in slightly smaller, longer duration pulses than similar drops on the bottom surface due to differences in elasticity. Since the number of drops was small, these differences were ignored. Similarly, there was no obvious difference in results for drops onto a workbench or a wooden slab, so these data are not separated in this report.

Figure 5 is a typical time history for drops pivoted around the long edges. In particular, figure 5 shows a 3.5-inch drop onto the top surface of the LRU and pivoted about edge 8. The acceleration amplitude is 14 g's peak, with a duration at 10 percent of peak pulse amplitude of 21.5 milliseconds. Due to differences in elasticity, drops pivoted about the long edges of the top and bottom structures resulted in larger pulse shape variations than drops pivoted about the short edges.



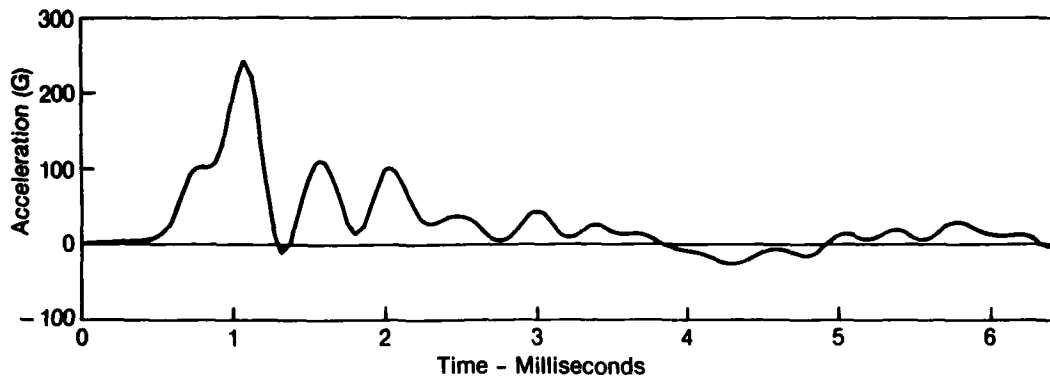


Figure 4. Representative bench handling shock time history (pivot edge 1, 4-inch drop)

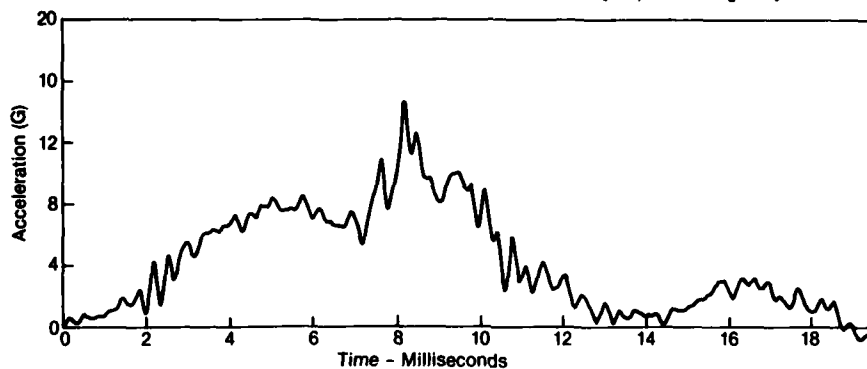


Figure 5. Representative bench handling shock time history (pivot edge 8, 3.5-inch drop)

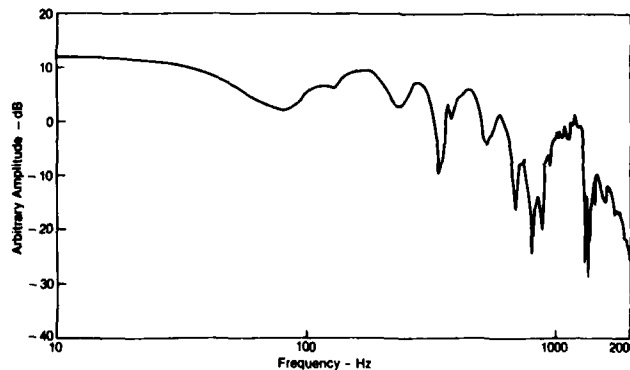


Figure 6. Representative frequency domain response spectrum for bench handling shock tests (pivot edge 1, 4-inch drop)

When the acceleration time histories collected for the bench handling tests are transformed to the frequency domain, the resultant spectra all have the general shape defined by the relation $\sin X/X$, modified by the structural response of the LRU. The major frequency components are concentrated in the 150 to 300-Hz range, with smaller but still strong components at higher frequencies. Figure 6 shows a representative spectrum demonstrating that the major frequency components coincide with the natural frequencies of typical printed circuit boards.

Table 1 is a summary of 24 drops of the same LRU, including drop height, pivot edge, peak acceleration, overall duration at 10 percent of peak pulse height, and the area under the transients. The major differences of drop height, pivot edge, and surface impacted are reflected in the resultant data as variations in pulse height and duration. Attempts to plot pulse height or duration as a function of drop height ended in a meaningless jumble of data points. A more sensible presentation was found to be based on the area under a transient (g-seconds), which is the

TABLE 1 Summary of Bench Handling Test Results

Drop No.	Drop Height (inches)	Pivot Edge	Peak Acceleration (g's)	Duration at 10% of Pulse Peak (msec)	Area Under Transient (g-seconds)
1	4	1	244.97	3.28	0.1906
2	4	1	257.86	3.87	0.2082
3	4	2	227.66	5.82	0.2057
4	4	2	175.51	4.45	0.1912
5	4	2	234.99	3.13	0.2074
6	4	5	175.51	3.62	0.187
7	4	5	94.07	7.84	0.1629
8	4	5	178.16	4.94	0.2133
9	4	6	223.88	1.91	0.1869
10	2	1	190.77	3.28	0.1306
11	2	1	109.61	3.76	0.1469
12	2	2	157.35	5.96	0.1430
13	2	2	111.36	4.49	0.1358
14	2	5	95.25	3.76	0.1321
15	2	6	117.52	3.47	0.1312
16	2	6	101.11	4.40	0.1542
17	3	3	51.28	8.11	0.0673
18	3	3	31.95	3.81	0.0747
19	3.5	4	138.39	11.48	0.1073
20	3.5	4	152.46	4.83	0.0913
21	3	7	42.21	11.93	0.0856
22	3	7	50.31	8.79	0.111
23	3.5	8	14.17	21.54	0.0815
24	3.5	8	8.90	24.96	0.0235

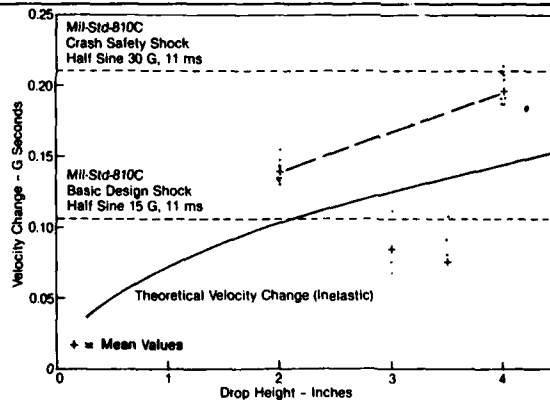


Figure 7. Velocity change vs drop height for bench handling shock test

resultant velocity change. In figure 7, scatter of data about the mean value for each drop height is relatively small. The data for 2-inch and 4-inch drops with pivot edges 1, 2, 5, and 6 parallels the calculated theoretical velocity change with no rebound.

DESCRIPTION OF INSTALLATION SHOCK TESTS

To our knowledge, installation tests are not described in any test standard. By installation, we are referring to the process in which a rail- or track-mounted assembly is firmly slid along its rails or tracks to be seated within a parent enclosure or structure. This situation most commonly occurs with ground-based or large airborne electronic systems in which modular electronic

assemblies are housed within cabinets or racks. If, during the installation process, misalignment or resistance is encountered, or if the operator's happens to be somewhat more "aggressive" than normal on a particular day, then the resulting impact when the hardware is seated may be nontrivial.

Three different LRU's (table 2) were subjected to the installation tests. These LRU's were production hardware in good physical condition. That is, the slide rails and dagger pins were properly aligned and no abnormal resistance was present during installation. To establish a degree of consistency in the force applied to seat the LRU's, installation was performed by personnel familiar with the assembly and disassembly of

the hardware using a "normal" amount of force. While this approach is admittedly unquantifiable, the mass of the LRU's involved probably tends to establish reasonably constrained upper and lower bounds on the force applied. In addition, since the personnel involved were using a "normal" push to install the LRU's, the measurements obtained most likely represent the lowest responses that might be expected in service.

Figure 9 is a representative time history after 700-Hz low pass filtering to show the faired (high frequency responses with insignificant damage potential filtered out) amplitude. The peak amplitudes of the filtered pulses were about 6, 16, and 7 g's peak for the filter, power supply, and power amplifier LRU's listed in Table 2. Durations of the seating pulses were 10, 10, and 9 milliseconds, respectively. Repeated installation resulted in remarkably consistent

Table 2.
LRU'S Used for Installation Test

Description	Height	Width	Length	Weight
Filter	8 inches 20 cm	11 inches 28 cm	18 inches 46 cm	71 pounds 32 kg
Power Supply	10 inches 25 cm	11 inches 28 cm	18 inches 46 cm	75 pounds 34 kg
Power Amplifier	15 inches 38 cm	11 inches 28 cm	18 inches 46 cm	114 pounds 52 kg

One accelerometer oriented longitudinally (the direction of the push) was used to measure the shock pulses resulting from installation. (The inside of the LRU was not accessible due to constraints imposed by production testing.) The accelerometer was located on a structural hardpoint on the lower edge of the front face of the LRU (Figure 8). The location was selected based on minimal attenuation between the accelerometer and the dagger pin contact point at the rear of the LRU, thereby providing a close approximation of the shock experienced at the dagger pins. (Space constraints prohibited installing the accelerometer on the rear face of the LRU.) Each LRU was "installed" 3 times.

INSTALLATION SHOCK TEST RESULTS

Time histories which resulted from typical LRU seating operations during installation show an initial low level portion 30 to 35 milliseconds long followed by the major seating pulse. The raw data showed high frequency "fuzz" at about 3 to 4 kHz due to friction between the sliding surfaces.

data, attributable to the numerous installations of those LRU's by the same personnel.

BENCH HANDLING AND SHIPBOARD EQUIPMENT

It is important to note that the bench handling test has been given increased importance in Mil-Std-810D. In addition to its traditional inclusion in Method 516.3, Shock, it now appears in Method 514.3, Vibration, under Category 9 for shipboard vibration. A sequence is recommended in which bench handling shock (or basic transportation vibration) is followed by a shipboard random vibration test. This sequence recognizes the fact that the most severe dynamic environment experienced by the majority of shipboard electronic hardware is transportation and handling. Low-level random vibration is performed with the equipment operating following bench handling shock (or transportation vibration) to verify that no physical damage has been sustained that would compromise equipment performance.

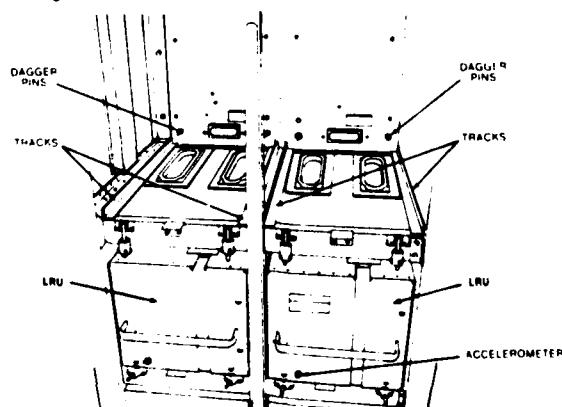


Figure 8. Accelerometer locations for installation shock test

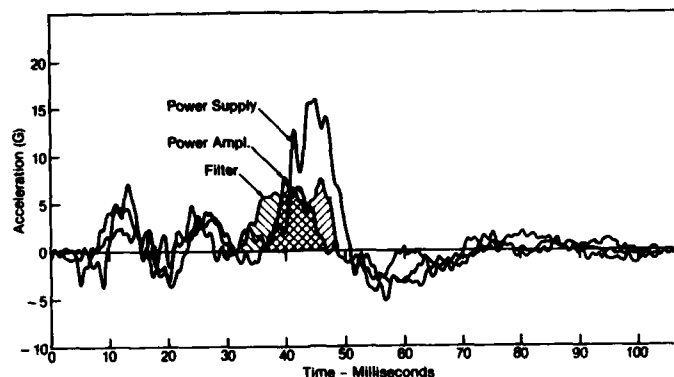


Figure 9. Representative time histories for installation shock test

CONCLUSIONS AND RECOMMENDATIONS

The measurements made during the bench handling and installation shock tests described above lead to several significant conclusions and recommendations:

1. The 4-inch drops made during the bench handling tests produced shock pulses with levels ranging from 94 to 250 peak-g's and durations up to 8 milliseconds. The majority of the energy associated with this response was concentrated in the 50- to 300-Hz frequency band, a region of particular importance for the majority of typical LRU structures.
2. The shock pulses measured during the installation shock tests ranged from 6 to 16 g's with a duration of approximately 10 milliseconds. Again, this response is in a frequency band of importance for typical LRU structures.
3. The peak responses measured for both bench handling and installation shocks represent energy that lies between that associated with the traditional basic design and crash safety shock tests of Mil-Std-810C[2]. This would seem to indicate that these shock producing events should be given at least as much attention during design and testing as those events that are traditionally considered, especially since bench handling and installation are far more likely to occur on a regular basis.
4. Serious consideration should be given to performing bench handling shock tests in place of more traditional

shock tests that require extensive fixturing and instrumentation. The similarity in energy levels between bench handling and traditional shock pulses suggests that equivalent effectiveness can be achieved with greatly reduced test time and expense.

5. Efforts should be made to establish a more rigorous characterization of the rough handling environment. Is the 4-inch drop height sufficient? How often are such events likely to occur in an equipment's lifetime? Is it likely that the equipment would be operating during such drops in service use? Such information would serve as a basis for updating the bench handling shock test in Mil-Std-810.

SUMMARY

Since shock producing events similar to those described in this paper are likely to occur on numerous occasions during the lifespan of an item of electronic hardware, the resulting acceleration levels and associated frequency bands make it especially important to consider bench handling and installation shock (as well as other forms of rough handling) in the design and testing of most electronic hardware. In particular, special attention should be given to those classes of equipment which are not normally thought of as encountering significant shock or vibration environments in end-use or mission application. In many cases, serious consideration should be given to replacing traditional shock pulse tests with bench handling shock tests that require less time to execute and need no special fixtures or instrumentation. There is considerable need for a more comprehensive definition of the rough handling environment in general.

REFERENCES

1. V. J. Junker "The Evolution of USAF Environmental Testing," Air Force Flight Dynamics Laboratory Technical Report AFFDL-TR-65-197, October 1965
2. Mil-Std-810C, "Environmental Test Methods," 10 March 1975
3. Mil-Std-810D, "Environmental Test Methods and Engineering Guidelines" 19 July 1983

DISCUSSION

Mr. Volin (Shock & Vibration Information Center): Do you think the fact that test engineers handled the LRV rather than field personnel made any difference in their bench handling and installation shock environments?

Mr. Szymkowiak: If one runs the Procedure VI test mentioned in the Standard, the tests are probably run by technicians. But there are things that happen in the field that are not documented. They are not in accordance with any kind of test. They just happen. You lay it down, but you are not quite on the table when you drop it. Or, you turn it over, and being that it is a heavier unit than you expected, you drop it and it gets a bigger drop than you expect. I assume things occur in the field which are sometimes more severe than in the laboratory. This particular test seems to be relatively close to the kinds of things that happen in the laboratory. We have had occasions where somebody accidentally dropped a special circuit board. Then somebody else said, "Hey, I want to know what went on there," so we did some drops to see. You get 600 G's on a circuit board when you drop it from waist high. Still, it is basically a sharp spike followed by the ringing of the circuit board because it bounces through the air and it is undamped for awhile. So things that accidentally happen in the field are quite often more drastic than any of the tests in the laboratory.

Mr. Volin: I would agree with you. Certainly a more rigorous characterization of this bench handling environment is needed. How did your measured bench handling shock levels compare with the shock levels or vibration levels that might be experienced in transportation? Would the transportation tests be more severe than the bench handling shock test?

Mr. Szymkowiak: That is a good question, but I don't have the answer for it. One of the things that I wanted to do, but I didn't get around to, was to put an LRU in a carton and measure the environment. There are many people that have to take equipment and move it in a truck from one spot to another. Method 514 has a number of cargo handling spectra, and these are measured spectra. I think if you look at those spectra, you can think of them in terms of shock spectra response, or Sheldon Rubin's method for comparing a shock test to a vibration test. There is a lot of energy at the low frequencies which should not do any damage, but it does. It bends corners, and it generates shock pulses in the item which ultimately damages crystalline structures. They are meant to pass their relatively benign end use environment, but if they haven't survived riding around in a truck, they may not have been tested enough.

Mr. Binder (United Technologies Corporation): Wouldn't you feel that the NAVMAT requirements for the random vibration burn-in or workmanship

type tests were more of a governing environment than the four inch drop relative to printed circuit boards and their dynamic responses?

Mr. Szymkowiak: In terms of circuit boards, yes. Definitely! .04 g²/Hz applies to a circuit board with a Q of 30 is disastrous. For a small box, hard coupled to an exciter, you can get rather high response, 50 to 60 G's on circuit boards. This is one of the reasons why some of us think one should not just arbitrarily apply NAVMAT P-9492. One should look in terms of the circuit board response on the highest Q board and keep the deflection down to a reasonable point, the threshold of damage point. Similarly, NAVMAT P-9492 just doesn't work for very large hardware. The question is, is NAVMAT P-9492 more severe than bench handling? I think it definitely is.

Mr. Binder: Do you still recommend the drop test if one faces a stringent burn-in requirement?

Mr. Szymkowiak: If one does a NAVMAT P-9492 orthogonal test, quite often you don't excite the circuit board as much as if you drove it at an angle. Again, it is a repeated thing. It lasts for a long time. You have many three-sigma excursions that bend the boards and stress the leads on the chips. So I think random vibration would be much more severe.

Mr. Silver (Westinghouse Electric Corporation): I think 200 g's is a lot more than the 18 g's that you get from the NAVMAT P-9492 procedure. So, if you get 200 g's, you could easily knock something loose that would never be knocked loose in the NAVMAT P-9492 test procedure.

Mr. Szymkowiak: True.

IMPACT OF 810D ON DYNAMIC TEST LABORATORIES

Dr. Allen J. Curtis
Hughes Aircraft Company
El Segundo, CA

At the request of SVIC, Allen Curtis delivered an informal assessment of the impact of 810D on Dynamic Test Laboratories to the Wednesday afternoon session, entitled "MIL-STD-810D, Session II, Implementation and Use." The following is a transcript of his remarks lightly edited for improved readability. The viewgraphs used on that occasion to structure his comments have been included as figures.

INTRODUCTION

I was asked to assess the impact of something we haven't used yet. Of course that means looking into a crystal ball. I am not sure how cloudy mine is. In trying to assess that impact, I couldn't completely ignore the new method of coming up with requirements, since that clearly feeds on down into the laboratory. I put a copy of MIL-STD-810C and a draft copy of MIL-STD-810D side by side and just went through the dynamic tests to try to compare them to assess this impact. Of course when I did that, an official copy of MIL-STD-810D was not available, but I think my draft was pretty current unless Dave Earls pulled a fast one on me.

I would like to tell you a little bit about how MIL-STD-810D evolved, although if you heard Dave Earls this morning, he described it more completely. Some of the impacts are sort of general, and they apply no matter what test method you are using. I call those overall impacts. Then I would like to look at the individual dynamic test methods. Under acoustic test methods there are Method 515 and Method 523. Method 523 is completely new. It is called Vibroacoustic Temperature, and it describes a method that is used at the Pacific Missile Test Center at Point Mugu, California, to reproduce the combined vibroacoustic and temperature stresses that external aircraft stores experience. Vibration is now called out in three methods. Method 514 is the old vibration method. Method 523, Vibroacoustic Temperature, of course calls out some vibration. Then there is the a mission profile test which is Method 520. So, if you want to vibrate something, you have several choices of method.

Finally, I would like to make a few summary remarks. "810 Dolly," as I fondly call it, has had a gestation period of about five years. As

Dave Earls mentioned, it has a new name. Not only is it Environmental Test Methods, but we have added "and Engineering Guidelines" to the title because now it is much more than just test methods. That is because of the tailoring concept which requires us to do some environmental engineering to try to arrive at more realistic design and test requirements. There are fall back numbers for most of the test methods in the standard.

The two-section format is very different, and I think if I gave MIL-STD-810D to one of the guys down in the lab, he will be overwhelmed. Section One, which tells you how to arrive at the test requirements, he won't understand at all, because you have to perform a life cycle analysis and you have to predict test levels. That is a little out of his ken. Then he looks at Section Two, which is really the test methods, and it is pretty brief. It says, "Do whatever you came up with out of Section One." So he will really be at a loss. If you just give a laboratory guy the new standard, he really has no idea of what he is supposed to do. That will be kind of a new experience.

I would like to broadly compare the two versions and point out where we have made some advances, where I think we have set ourselves a few traps, what this may do to how one runs a laboratory, and how to get the proper people and the proper facilities. I have listed some of the overall impacts that I see, by categories, in Figure 1.

OVERALL IMPACTS

PROGRAM SCALE

The first impact I would like to discuss is the scale of the program. Am I buying a black box or a whole airplane? If I am buying a major system, which perhaps is a whole airplane, we have been tailoring for years. I can remember

back when I was associated with the Phoenix program. Well, I guess Joe Popolo thinks the F-14 was the major system and the Phoenix was the major subsystem. Anyhow, we worked together on that. We did a lot of tailoring. We did random vibration on random vibration which we called stepped narrowband random vibration. We used response control when we tested large objects, and we used pulsed gunfire on the equipment. It bore no relationship to what one could find in the standard. In other words, we tailored and we will continue to tailor on the major systems as we have done in the past. I do think that one advantage which will fall out of this is that if I am the supplier of a major subsystem, and if the person to whom I supply that equipment is the major system contractor, I, as the subcontractor, will have a little more leverage to equalize the bargaining power between the buyer and seller in that case. Hopefully, it will still permit the subcontractor to do any tailoring which is appropriate and necessary because of the characteristics of what it is he is supplying rather than the characteristics of the vehicle. But I think when we get down to small equipment, for example, when I am selling one or two "black boxes" either to another contractor or to the Government, things will become a little more difficult to implement. It will be more difficult to get the data I need to do proper tailoring, and I am not sure that the smaller contractors will have the necessary resources to do it, both in terms of people with the proper experience and the proper skill level. Furthermore, the dollar value of the contract is likely to be a lot less, so there will be a reluctance to spend the necessary number of dollars to do the tailoring. When we have the production of a single black box to go into a number of vehicles, that in itself may make the tailoring a little bit difficult; I suspect it may also lead to a certain amount of over-engineering. Of course, all of these things will lead to more difficulties in the laboratory.

ENVIRONMENTAL ENGINEERING

As Dave Earls mentioned this morning, it requires a lot of environmental engineering to do the tailoring and develop test requirements. I sometimes wonder who all these people are that are running around with large scissors in their hands. Does the buyer do it? Does the seller do it? The Standard says something about it shall be done by the supplier when the contract so states. It does not say who does it when the contract does not so state. Beyond that, what is the background of the people who will be doing this? Do they tend to be people who are more experienced in structures who look at the world through a certain set of eyes and biases? Or is done by the people who are more used to environmental testing and who have their set of biases which are a little bit different than those of the structure people? Or perhaps reliability engineers might do it, and that could be a third

set of biases. If we get too analytical in our tailoring efforts, I think we could come up with some rather weird and wonderful requirements that perhaps would tend to bend the laws of physics. We must be pragmatic and empirical if we are to use this in a practical way. I can well imagine receiving some test requirements that would be pretty expensive to perform even if physically possible.

The good thing is we have these Data Item Descriptions (DID's), which hopefully will appear in the Contract Data Requirements List in the contract. David Earls described their purposes this morning; I think this will be a great leverage for people in our business. It will give us a reason to have the proper resources allocated to our efforts which we have lacked in the past when money and time were short. Sometimes we are considered frills that one can get along without. On the other hand, preparing Environmental Impact Reports is a cottage industry that has sprung up in the last few years. I hope we don't have a new cottage industry which has to do with interpreting or preparing DID's for MIL-STD-810D. We have to keep a balance and not go overboard.

TEST PLANS AND PROCEDURES

It is obvious that we must put more work into this area, and that will take some money. Test plans will have to be much more comprehensive to reflect the outcome of our efforts in tailoring the test requirements.

FACILITIES

If the test requirements get too complicated, we will need some new facilities, and I will discuss that in more detail later. We must worry about the lead time and the capital dollars required to procure the proper facilities in time to do these tests when they are required. There is a trap here because I remember years ago, going back to Phoenix again, we decided to do stepped narrowband random vibration testing. We got some money to develop the specialized facilities to perform those tests in our own laboratory; and since we had a couple of them, we had enough capability. But then we also subcontracted a couple of major "black boxes," and one of these was a computer system. First of all, the subcontractor did not understand the requirement when he first read them, and second, he certainly didn't have any facilities to perform these tests. So there is a trap when you generate these fancy test requirements. As they flow down through the tiers of contractors, things get more and more complicated and costly.

TEST COSTS

I suspect the cost of testing, on the average, will increase. As we go into MIL-STD-810D, we will run better tests because they will be better engineered. Hopefully, we will save as much or more money on the equipment side as

we have to spend extra on the testing side so that the overall program cost should not be impacted. Unfortunately, sometimes the overall costs are ignored. It is the cost in each individual pocket that is scrutinized.

COMPUTERIZATION

If I could offer one rather mild criticism of 810D, I think it is in the fact that it doesn't properly recognize the almost universal availability of digital test control equipment these days. It still tends to be written in the analog world and says, "If you have digital equipment, do something similar." It is too bad we didn't have time to do something about that.

Another fear I have with the computerization is that somebody will invent some fancy math models of the environment, and they will use it to tailor the test requirements. I suspect they will then come up with some very strange test profiles; and because of their unfamiliarity with what one can and cannot do in the laboratory, we could have a few problems to work out there.

SKILL LEVELS

Considering the skill levels of the people who inhabit laboratories, we will need many more and they must be better trained and have a higher skill level. We must spend some money developing those people, but I am not sure how you justify that expense.

SAFETY

As to safety, of course familiarity breeds contempt, but on the other hand, unfamiliarity also breeds pitfalls. I think as we do new and wonderful tests, if I may use a little slang, we will have to "kluge up" something to be able to run some of them. "Kluges" tend to be a little less safe than the productized test facilities. We will have a few accidents on our hands before we are through.

WITNESSING

As these tests get more complicated, we always have people hanging around our laboratory witnessing tests, and we have to prove to them that we did what we were supposed to do within the required tolerances. That may get to be a little more difficult with some of the things that I see in 810D. There is a funny one in this area that I'll come back to on gunfire testing in a few minutes.

IMPACTS OF SPECIFIC TESTS METHODS

ACOUSTIC TEST METHODS

Figure 2 is a comparison of the acoustic test requirements in MIL-STD-810C and MIL-STD-810D. In 810C we only have two requirements; one for internal equipment and one for external stores. Both tests were oriented towards

qualification testing or contractual compliance. In 810D we have four types of tests, one of which is called environmental worthiness. We have the old qualification test and we have two new tests. The environmental worthiness test is intended to be used in that situation where you have one or two flight articles. You are not trying to demonstrate contractual compliance, but you just want to run a test which represents the environment you think the flight articles will experience in the flight test program. You can tailor to that, and it specifies what I call an engineering development test which is something we have never had in MIL-STD-810 before. I think it will save the government a lot of money.

Mission profile tests show up in Method 515, and also in Method 523. The cavity resonance test is a specialized test using a sinusoid tuned to the organ pipe frequency of the cavity where the equipment is placed.

SHOCK TEST METHOD

Sheldon Rubin talked about the new shock test methods this morning, and he pointed out some things that I had not noticed as I had read through it. Figure 3 shows several new shock tests. For example, we have a test for equipment to be packaged. This is when you test the equipment without the package. The Fragility Test is new, the Crash Hazard Test is new, and the Catapult Launch/Arrested Landing Test is new. The others are pretty much carried forward from MIL-STD-810C to MIL-STD-810D.

The shock test requirements are innovative in several respects (Fig. 4). We have gone to the shock spectrum specification for these tests. From the laboratory point of view, assuming you have software for computation of the required shock spectra, this will be a vast step forward. Whether you do the test on a shaker or do it on a drop tower, but still show that the drop tower had produced a transient with the required shock spectrum, it is still a great advance. There is a sawtooth pulse fallback. Interestingly enough, there is no half sine test called out anywhere in the document that I could find. Also, we now have the option of deleting the shock test when we can show that it is demonstrably less severe than the random vibration test to which we are committed. MIL-STD-810D has sort of given us something with one hand and then immediately taken it back with the other. I say this because I think the random vibration requirement at the low frequency end of the spectrum will seldom produce responses equal to the shock spectrum in that same frequency range. So I think we need something in the document which says that the responses must be equal or greater at frequencies above the first resonance or something like that. Otherwise, we will seldom be able to take advantage of this generous offer.

Three exclamation marks appear after the

"Trapezoid for Fragility" (Fig. 4), because that is a new requirement. There is a rationale in the document which says, "Well, we gave you that wave shape because it is easy to calculate the velocity change," which seems to me is a rather weak argument. A trapezoid is a wave shape that is just very, very difficult to produce in the laboratory, especially when you have the same tolerances on the wave-form that were in 810C. To be quite honest, I think it was a step backwards to go back to a waveform, especially an even more difficult waveform than we have been trying to live with for several years.

The Catapult Launch/Arrested Landing waveform, without too much detail, says, "Well, you are supposed to do that with some sine wave bursts." I am not quite sure how to do that in the lab. I don't think there is much equipment around yet that could do that in any repeatable, safe fashion that I am aware of.

GUNFIRE VIBRATION

Method 519, the Gunfire Vibration (Fig. 5), is essentially unchanged except that now the pulse method of performing the test is given first choice instead of being second choice as it was before. However, one is allowed to perform the test with random excitation completely, using narrow band spikes at the gunfire shot rate, or with multiple superposed sine waves, as shown in Figure 6 which is reproduced from MIL-STD-810D. I am not quite sure how you put those four sine waves on top of a random spectrum if you have a digital controller. The trouble with digital controllers is that they are very perceptive. When you do something that is out of spec, they immediately tell you so. In the good old days, when we had analog equipment, we could fool the system and be apparently out of spec, but nobody knew it because the equipment was not smart enough to tell us. I would be delighted if someone could tell me how to do Figure 6 safely and repeatably and within tight tolerances with a digital system, and presumably some other ancillary equipment.

One way of doing gunfire with digital test control equipment is with what I call clock warbling, a way of fooling the system. The A/D converter has a clock, so that it samples at a certain rate. If you can get into that clock and change that frequency, you can, in effect, make the clock run faster or slower with an oscillator. For instance, if I were doing a six-thousand-shot-per-minute gunfire test, i.e., 100 Hz, I would lock into my A/D converter. If I really wanted to do a gunfire test of between 90 Hz and 110 Hz to look at the variation in the gunfire rate, I would merely change my oscillator a little bit up and down the scale. But, there is only one real problem with this; no matter how I change the oscillator, the clock changes everything. When I output the documentation of the test that I did, it says that I stayed at 100 Hz all the time. Now I have to convince the QC inspector that no, I

didn't really stay at 100 Hz all the time. I really changed between 90 and 110 Hz in some prescribed manner. But there is no way in the world that I can ever prove that. (without other independent instrumentation/analysis, Ed.) I made a remark previously on the problems of the sine wave test.

VIBRATION TEST METHODS

Let's go to good old Method 514, the standard vibration test method in Figure 7. The conditions that you have to create come out of Section One through a tailoring process. The material that you find in there for fallback positions, if you don't have field data or a good prediction method for aircraft stores, is unchanged. There are some completely new requirements for testing ground vehicle equipment which I hear were introduced at the 53rd Shock and Vibration Symposium one exciting evening about a year ago.

If you just stand back and look at Method 514, you will observe that most of the testing is done with random vibration; hardly any sine wave testing remains. That sine wave testing which is still there calls for source dwells, where we now have to dwell at or near excitation frequencies, e.g., the rotor blade passage frequency for helicopter vibration tests. There are no resonant dwells, which all of us have loved so dearly for all these years, because it tore up everybody's equipment so fast. Resonant dwells have disappeared and, as a taxpayer, I have to think that is a great step forward. So we have these new categories that I have mentioned, especially for ground equipment, and we also have a minimum integrity test.

Basically, Method 514 has four ways of testing (Fig. 8). You can either do it on the traditional electrodynamic shaker; if you have big equipment that is transported quite a lot, you have to test it on a Munson Course. If it is smaller equipment to be transported often, you can put it on a package tester. Then there is a method for response control which is a random excitation, but it is tailored for external stores. We also have Method 520 which is a combined temperature, humidity, vibration, altitude test. It is sometimes called a CERT test or a mission profile test. Then we have a combined vibroacoustic-temperature test (Method 523), which I have said is mainly attributable to the Pacific Missile Test Center at Point Mugu. Of course, all of these are very good tests, but they have specialized applications. They should only be used under the right circumstances. I am a little concerned that somebody will write a "spec" and they will say, "Let's see. What should we do?" They will then see Method 520 or 523 and say, "Gosh! That sounds pretty exciting! I think I will include that." They may not realize the very large impact on cost. Those are very specialized facilities, and you don't just call those out willy-nilly. But I suspect somebody will.

I tried to assess what was going to happen in the laboratory in terms of facilities. In the area of vibration exciters, not much will happen. Whatever you have will probably continue to be used. If one wanted to make a rather wild guess, on the average over the years, I think the size of shaker that we need will probably tend to be smaller merely because we will come up with better test levels that won't be quite so conservative, and therefore, we will get by with smaller shakers. But in the area of the control systems to drive those shakers, much work must be done. SMOP is an acronym that stands for "Small Matter of Programming." Those of us who have digital controllers will be faced with a lot of SMOP's. For instance, I happen to have the software package for the random vibration on random vibration which is this new swept narrow band random vibration test that the Aberdeen Proving Ground has come out with. But if you don't have it, there is no way that you can do that test. If you do have it, you are in good shape. Also, if you are going to run sine plus random vibration test, you must get that software productized so that those tests can be done repeatably, safely, and in a way that keeps everybody happy.

Some of the tolerances have been changed, and when you get into more complicated tests, the tolerances with which you do them also tend to get more complicated. Demonstrating that you met those tolerances is even more complicated. MIL-STD-810D is a little deficient in the discussion of tolerances when using digital controllers. However, Figure 9 shows one positive thing for random vibration testing; the tolerance allows us to be "out of spec" but not too often. We are allowed to exceed the dB limit in the negative direction provided we don't do it over too great a bandwidth. We can also exceed the +3dB limit, i.e., overtest. I think it says, "At your own risk," or "At the seller's risk" or something like that. That should make life easier in the laboratory on a number of occasions. We will have a problem with (tolerances for) random on random vibration tests since there is absolutely no software available, and I can't even conceive of how one could develop the software to put satisfactory post-test data around that nonstationary process.

I was a little disappointed at the sine wave test tolerance because we still have a plus or minus ten percent tolerance; that is all it says. Presumably, that tolerance applies when we are mixing sine and random vibrations. It also still applies when you have a very complex sweep. Where there is a ten to one ratio in the test level, that means that at the high frequency end the tolerance about that level is equal to the test level at the low frequency end. That leads to some problems in signal to noise ratio. We have one "spec" that instantaneously drops from 70 G's to 7 G's at 1300 Hz. There is no control system that will do that! You have to come down over some kind

of a bandwidth which is never addressed. Furthermore, suddenly the tolerance I had a millisecond or two ago is now equal to the test level that I am trying to control to. I think we ought to be able to do better.

We have one new tolerance which is time, and it says plus or minus one percent (Fig. 10). Personally, since frequency or time are the independent variables, I never was quite sure why we even put a tolerance on them in the first place. Plus or minus one per cent of time says that if I am running the typical one minute, flight acceptance test on space hardware, the tolerance on that duration is six tenths of a second. None of my controllers can measure time to closer than the nearest second. That may lead to some discussion with some of our inspectors on occasion.

Documenting what you did in the test is given very little consideration in the standard. When we have stationary test conditions, we may have to put out a lot of data, but at least it is relatively straightforward to do so. When we run some of the nonstationary tests, it is a little difficult to know quite what one should do to prove we have tested correctly. The old swept sine vibration test was pretty easy; you just put it on a "Visicorder," and you either filtered it or you didn't. To get a good time history and a record of the complete test wasn't too bad. But for any other nonstationary processes, it gets a little more complicated. It can certainly lead to rather voluminous test reports.

One of the spectra that one finds in MIL-STD-810D, courtesy of the Aberdeen Proving Ground, is shown in Figure 11. Obviously, that is a tailored spectrum, and I think it probably comes under the classification of very fine needlepoint. The break points and PSD values are tabulated over on the right hand side of the figure. At 15 Hz, the PSD value is .08838 G² per Hz, and at 16 Hz the PSD value is .32948 which is about 6dB difference. But at 19 Hz, the PSD value has dropped back in excess of 6dB. So I suggest that maybe we overdid it a little bit with some of these. Furthermore, it is clear that you must use an analyzer with no more than a 1 Hz frequency resolution. That is the absolute maximum that you could use, and you would like to use something like 1/2 Hz if you are to have any chance of controlling that spectrum within any type of tolerance. The control system will not be able to do it if a lively test item is on that shaker.

Figure 12 shows the spectrum for the random on random, or a swept random vibration test. Notice that there is a broad band random vibration level down at the bottom and five spikes. The spikes are ganged together, and they slowly sweep up the frequency band. The first sweep bandwidth is 2 Hz wide. The table of those bandwidths shows the range of frequencies over which they sweep, and these

bandwidths are proportional to each other. It takes you 15 minutes to sweep the 2 Hz bandwidth from 30 to 35 Hz. Recently, I heard Senator Dole on television describe a rather rapidly moving event as watching paint dry. If you run one of these tests and try to check it out, it is sort of like watching paint dry. You sit there, and you look at the scope, and you think, "Gosh, is that thing moving? Is it doing what it is supposed to do?" And you sort of sit there forever. It will be interesting to see how this goes on. The other problem is when you use the available software to do this, you have to fool the system. You have to make the tolerances about the base band level wide enough at the high frequency end to more than embrace the narrow band peaks. You also increase the tolerance on the overall RMS much more than we are accustomed to doing so the system will not shut itself down every time it moves a frequency. That means that we are, in effect, throwing away all of our safety features and just keeping our fingers crossed.

Figure 13 shows the suggested spectra for propeller aircraft and for equipment sitting on the engine. It shows the source dwells that were mentioned before which can be tested as narrow band random spikes. I think that achieving those spectra is rather straightforward because those spikes just stay at one place; it is a stationary process.

A new test in MIL-STD-810D is called the minimum integrity, and it is a pretty straightforward sine sweep, (Fig. 14). It just seemed to me that it was rather stringent because the recommended test duration is three hours per axis sweeping at five G's which I thought was a pretty high minimum. The random vibration requirement for this minimum integrity test for stores is about 8 G's RMS, (Fig. 15); it is somewhat similar to the spectrum in NAVMAT P 9492 except that it ($0.04 \text{ G}^2/\text{Hz}$) is extended down to the low frequency end.

The Standard mentions what they call common test techniques. It describes them a little and gives guidance on what they are for and how they should be done. Some of them are not quite as common as others, at least not yet. Response characterization is a transfer function measurement. The Standard says you can do that either sinusoidally or with a random excitation. I think in the era of FFT's, to do it with sine wave excitation is not quite keeping up with the technology, and you certainly don't get anywhere near the information that you can from random excitation.

The Standard also calls out several test types, such as engineering development and environmental worthiness. I think these tests will be a boon to the laboratory and to the development programs because they will help us to avoid the problem of running a full bore qualification test on a one-of-a-kind piece of equipment. This always tended to get people a little nervous especially when it forced us to

do things at extremes. Now these two earlier kinds of tests encourage us to do the testing at levels appropriate to how we are going to use the equipment.

CONCLUSION

Having gone through these two standards and mulled over these thoughts that I've just been sharing with you, what can one conclude about how MIL-STD-810D compares to MIL-STD-810C? Technically we've made some very significant advances in a number of areas. We should be able to make our testing and our efforts a lot more useful than they have sometimes been in the past. I don't really see any major facility impacts, unless you count software as facility. Tests will be somewhat more expensive, but hopefully there will be a decrease in the overall program costs. But I also think there will be a few little problems along the way as we learn to live with MIL-STD-810D. I have at least imagined in my own mind some really chaotic conditions that can probably arise as we try to use this. On the other hand, some of those are kind of fun to be in the middle of anyway and will pay off in the long run.

- PROGRAM SCALE
- ENVIRONMENTAL ENGINEERING
- TEST PLANS/PROCEDURES
- FACILITIES
- TEST COSTS
- COMPUTERIZATION
- SKILL LEVELS
- SAFETY
- WITNESSING

Fig. 1 — Overall Impacts of MIL-STD-810D on the Dynamic Testing Process

- "C" I. INTERNAL EQUIPMENT } QUALIFICATION
N. ASSEMBLED STORES }

- "D" I. ENVIRONMENTAL WORTHINESS
II. QUALIFICATION TEST
III. MISSION PROFILE
IV. CAVITY RESONANCE

Fig. 2 — Acoustic Test Methods (Methods 515 and 523)

- "D"
- I. BASIC DESIGN
 - II. EQUIPMENT TO BE PACKAGED
 - III. FRAGILITY
 - IV. TRANSIT DROP
 - V. CRASH HAZARD
 - VI. BENCH HANDLING
 - VII. PYRO
 - VIII. RAIL IMPACT
 - IX. CATAPULT LAUNCH/ARRESTED LANDING

Fig. 3 — Shock Test Methods (Method 516)

- "C"
- I. BASIC DESIGN
 - II. TRANSIT
 - V. BENCH HANDLING
 - IV. HIGH INTENSITY
 - VI. RAIL IMPACT

INNOVATIONS:

- SHOCK SPECTRUM SPECIFICATION & TEST
- SAWTOOTH FALLBACK (NO 1/2 SINE)
- DELETE WHEN VIBRATION DEMONSTRABLY MORE SEVERE
- TRAPEZOID FOR FRAGILITY!!!
- LAUNCH/LANDING WAVEFORM

Fig. 4 — Shock Test Requirements — Innovations

- ESSENTIALLY UNCHANGED, EXCEPT
- PULSE METHOD PREMIER METHOD
- ALL-RANDOM ALTERNATE
- MULTIPLE SUPERPOSED SINE WAVES

Fig. 5 — Gunfire Vibration Test Method (Method 519)

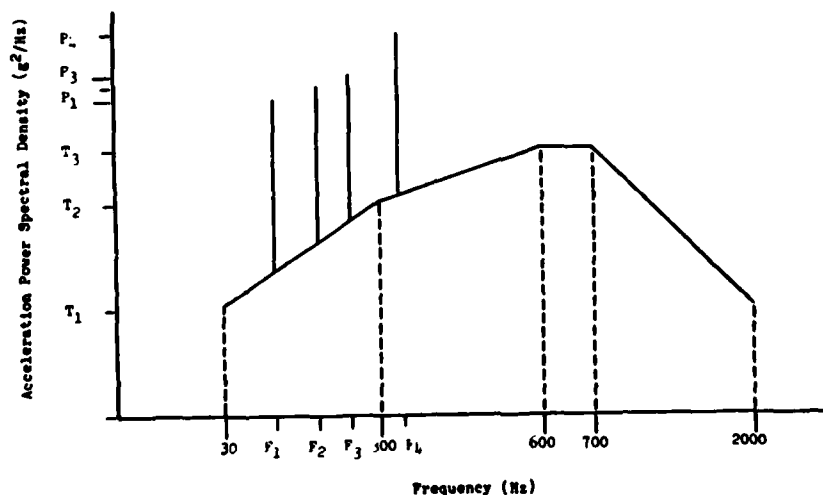


Fig. 6 — Generalized Gunfire Induced Vibration Spectrum Shape

CONDITIONS:

- TAILORING
- A/C STORES UNCHANGED
- GROUND VEHICLES COMPLETELY NEW
- MOSTLY RANDOM
- SOURCE DWELLS
- NO RESONANCE DWELLS
- NEW CATEGORIES — ESPECIALLY FOR GROUND
- MINIMUM INTEGRITY TEST

Fig. 7 — Vibration Test Methods (Method 514) — Changes

514: I. TRADITIONAL E.D. SHAKER TESTS

II. MUNSON COURSE

III. PACKAGE TESTER

IV. EXTERNAL STORES (RESP. CONTROL)

520: TEMPERATURE, HUMIDITY VIBRATION, ALTITUDE (CERT)

523: VIBRO-ACOUSTIC, TEMPERATURE (PT. MUGU)

Fig. 8 — Vibration Test Techniques (Methods 514, 520 and 523)

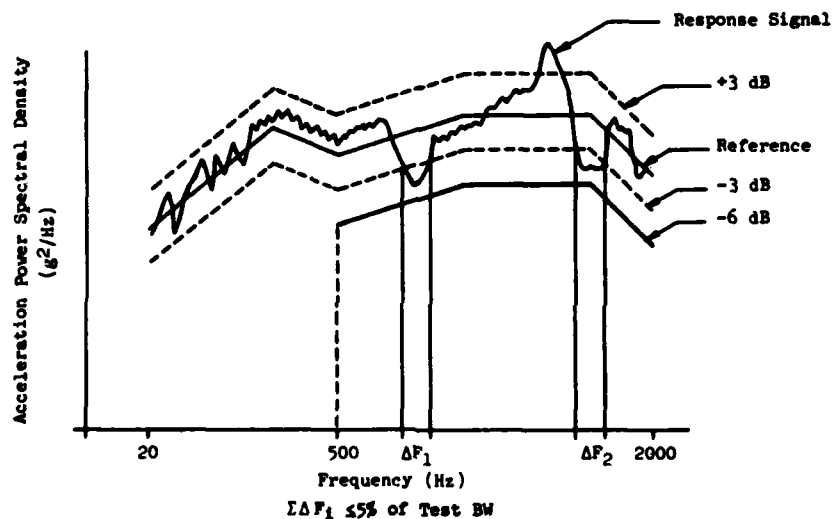


Fig. 9 — Example of Acceptable Performance Within Tolerance

- TOLERANCES
 - RANDOM
 - ROR
 - SINUSOIDAL
 - SINE PLUS RANDOM
 - TIME ($\pm 1\%$)
- DOCUMENTATION
 - STATIONARY
 - NON-STATIONARY

Fig. 10 — Vibration Testing Tolerances

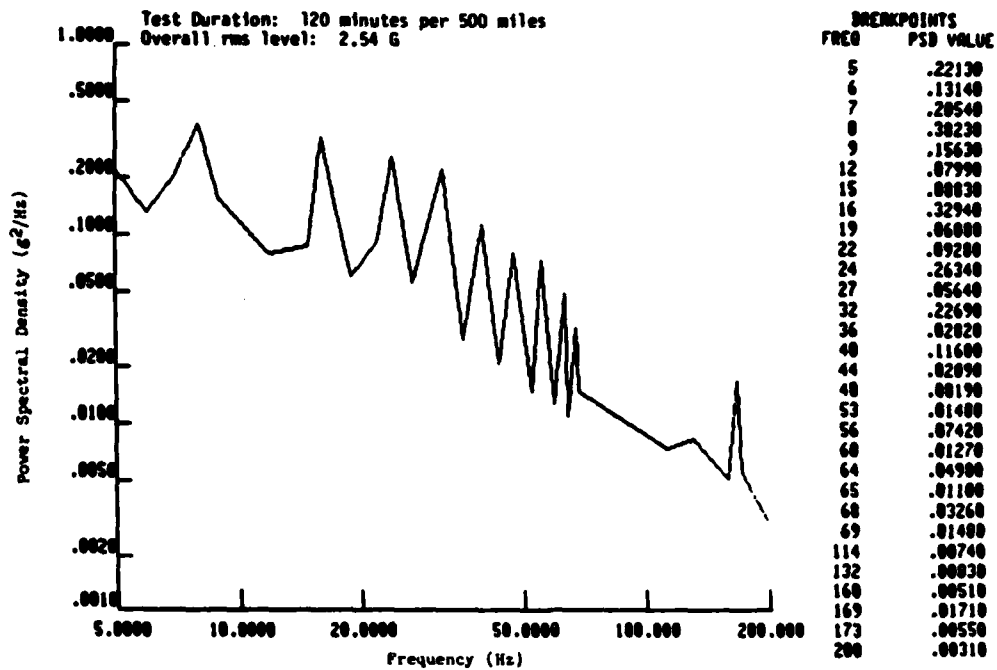


Fig. 11 — Basic Transportation, Composite Tactical Wheeled Vehicle Environment — Longitudinal

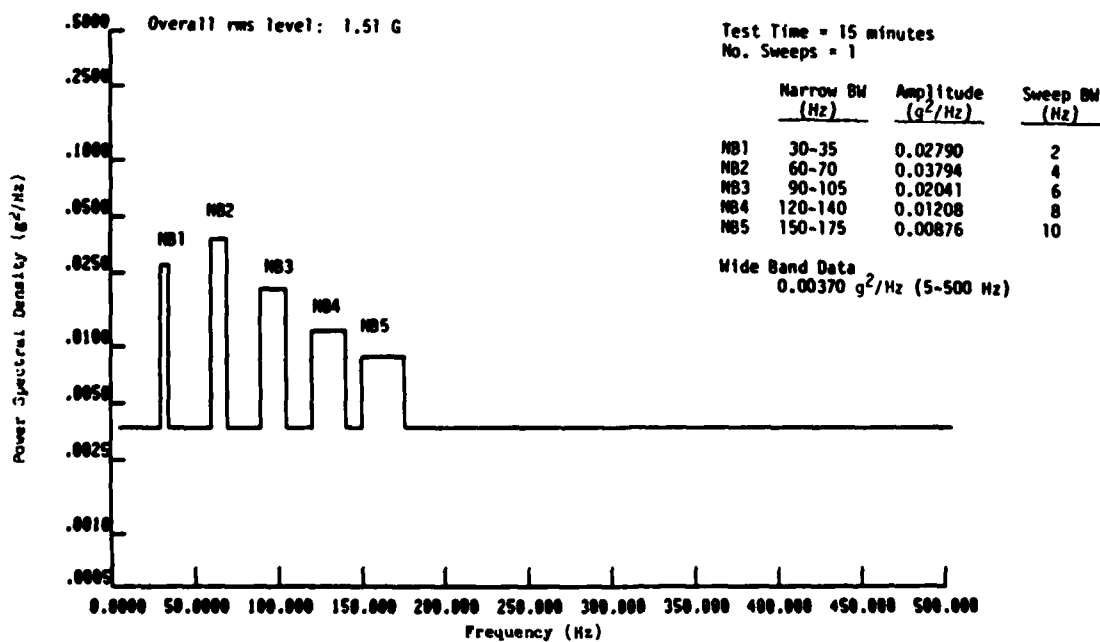


Fig. 12 — Example of Swept Random Vibration Test Spectrum

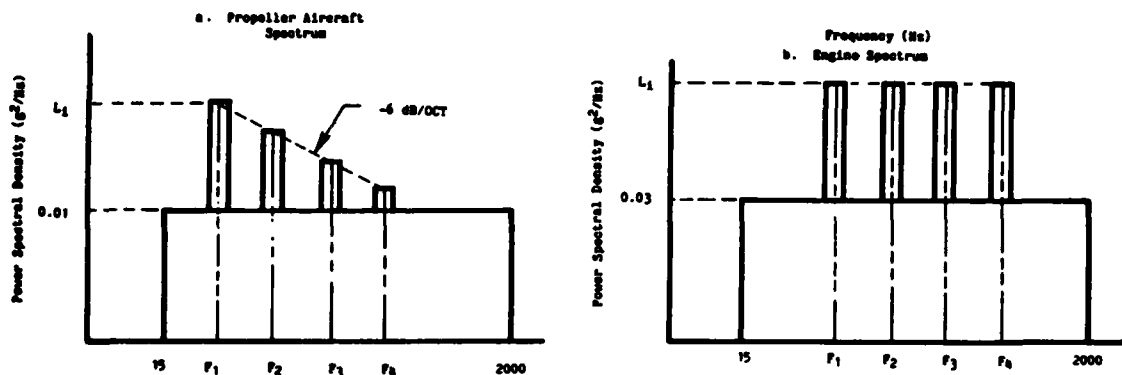


Fig. 13 — Suggested Vibration Spectra for Propeller Aircraft and Equipment on Engines

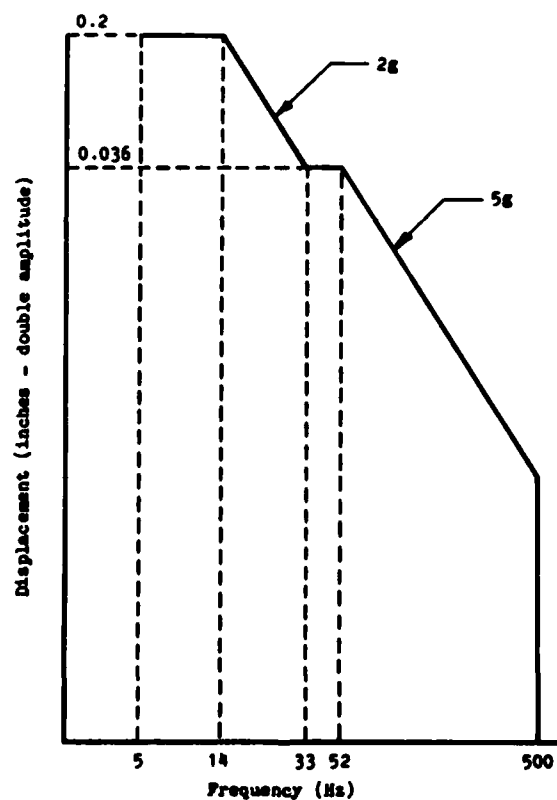


Fig. 14 — Minimum Integrity Test — Helicopter Equipment

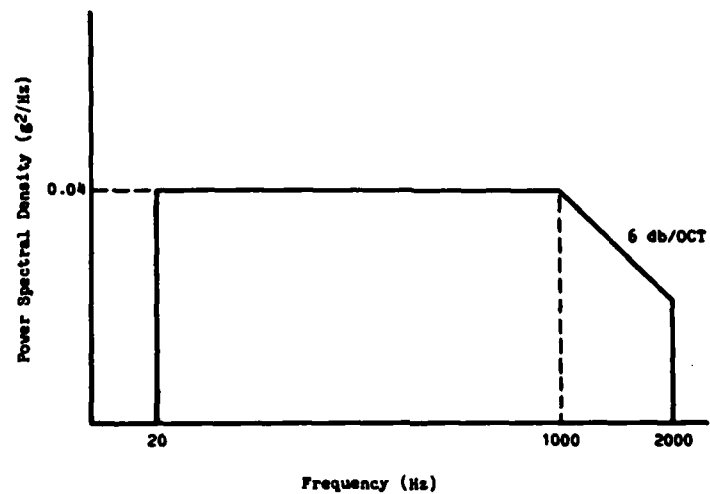


Fig. 15 — Minimum Integrity Test — Aircraft/External Store Equipment

THE CHANGING VIBRATION SIMULATION FOR MILITARY GROUND VEHICLES

Jack Robinson
Materials Testing Directorate
Aberdeen Proving Ground, MD

INTRODUCTION

Many changes in laboratory test schedules are found in MIL-STD-810D. Some changes are subtle while others are pronounced. One particular area of this MIL-STD which has undergone the pronounced change is the laboratory vibration simulation of secured cargo transport in military ground vehicles. We have suddenly gone away from the swept-sine schedules (Figure 1) that have existed in this internationally recognized testing document since its inception in June 1962. The A, B, and C versions, published in June 1964, June 1967, and March 1975, respectively, maintained the same type of schedules, but contained changes based either on a limited re-analysis of some existing data or on the analysis of some limited new data.

But now suddenly, and perhaps mysteriously to many, the swept-sine test schedules for military ground vehicles that we had all become accustomed to over the many years have vanished from the document; and random and complex random test schedules appear for the first time for these vehicles. To most, there was no warning that such a change was taking place; and to many, the driving forces behind these changes are not clear.

The intent of this paper is to examine why and how these changes in the vibration simulation for military ground vehicles have taken place, the impact of these changes, and what is needed to adequately accommodate them.

Why

There are many factors which have cumulatively influenced the decision to make the changes which have taken place in the military ground vehicle simulation schedules contained in Method 514.3 of MIL-STD-810D. Perhaps the one single factor which had the most influence was the ever increasing concern on the part of the designer/developer, as well as the tester, that the sinusoidal schedules were not representative of the real world environment. In that light, they were often considered an overtest which resulted in costly overdesign of equipment just to pass the test. Whether the latter is true

"across the board" is not the thrust of this paper and will not be further addressed. The fact remains, however, that the sinusoidal schedules were not representative of the real world. That fact has been recognized by more and more people and thus became one of the most significant factors influencing the change. Additional factors which complemented this lack of real-world simulation include: the increased and improved capabilities in data acquisition and data analysis in recent years; the decision in 1977 by the US Army Test and Evaluation Command's (TECOM's) Shock and Vibration Technical Committee to adopt and initiate a well-defined, multi-year, success-oriented plan for updating laboratory shock and vibration test schedules for equipment transported on and installed in military ground vehicles; the funding to support this plan; and the continuous Command and managerial emphasis placed on the resulting program to ensure its continuance and timely completion. When the efforts were initiated for the development of MIL-STD-810D, there was a sufficient portion of this TECOM program completed to allow incorporation of a limited number of the resultant test schedules into Method 514.3.

How

The topic of how these changes in Method 514.3 took place must include the overall process as well as the rationale used in the development of the new laboratory vibration simulation schedules.

TECOM's Shock and Vibration Technical Committee decided at its annual meeting in 1982 to launch a major effort to finalize, and propose for incorporation into the document, as many ground vehicle simulation schedules as possible within both the stringent time frame and existing data constraints. This effort was obviously successful since several new vibration simulation schedules have been included in the "D" version of MIL-STD-810.

The development of these new vibration schedules resulted in a change in technique for the simulation of this ground vehicle environment - that being the use of random and complex random schedules in lieu of the old,

familiar swept-sine schedules. The rationale used during the development of these new schedules needs to be understood by those who specify the use of these schedules in test plans and by the engineers in the test laboratories who use the schedules to evaluate materiel.

During the process of developing laboratory test schedules, there are several separate and distinct items that must be addressed and for which rationale must be developed. These include: which vehicles are used; how the vehicles are loaded; where the vibration data are measured; what terrain the vehicles traverse; how the data are reduced and translated into a laboratory test spectrum; the use of exaggeration factors for accelerated testing; development of corresponding laboratory test times; and the test of the new schedule. We will now turn our attention to each of these areas as related to the schedules in MIL-STD-810D.

Vehicles

The first portion of TECOM's Shock and Vibration Committee plan was to address cargo transport. Realizing that various sizes and shapes of vehicles are used by the Army for hauling cargo, a list of these vehicles was developed; and vibration data were taken on each of these vehicles as they were available for this investigative work. Trucks, trailers, semitrailers, and one tracklaying vehicle made up this somewhat extensive list of vehicles. As of this particular presentation, data have been generated on all of these vehicles. However, during the development of MIL-STD-810D, data were available only on the 5-ton M813 truck, the 12-ton M127 semitrailer, the $\frac{1}{2}$ -ton M105 two-wheeled trailer, and the one tracklayer, the M548 cargo carrier. As test schedules are developed on the remaining vehicles, the plans are to update the MIL-STD accordingly. As new cargo vehicles enter the Army's inventory, data will be collected on them; and the laboratory schedules will again be revised. This same process will also soon extend into the arena of installed equipment in military ground vehicles.

Loading

An investigation done under contract to the US Army Aberdeen Proving Ground addressed cargo loading and restraint in military ground vehicles. This investigation, which dealt primarily with ammunition and general equipment types of cargo, reached the following conclusions:

a. Most ammunition and general equipment are transported from the manufacturer to the forward supply point in either a palletized or containerized configuration; and from the forward supply point to the using unit it is either palletized or becomes individual items.

b. The cargo load as a percentage of rated capacity of the transport vehicle tends to be at

or near the capacity of that vehicle (either in weight or size).

c. Cargo restraint is extremely variable. In some instances cargo is secured in both the vertical and horizontal planes on the transport vehicle; there are likewise instances where no securing mechanism is used at all - which is what we call a "loose cargo" configuration. The majority of the time, however, the cargo is secured in the two horizontal planes but not in the vertical plane; and the horizontal restraints vary from rigid to loose. This configuration is what we call "restrained cargo".

d. Steel banding is used in the field as a securing mechanism and is also a representative means of securing the test load to the vibration exciter.

Although the investigation concluded that cargo is transported as secured cargo only a portion of the time, the decision was made to utilize the secured cargo configuration throughout this portion of the investigative work and then address loose/restrained cargo (which is a much more complex environment to measure and, probably, to simulate) in the future in accordance with the overall TECOM plan. Steel banding was used as the mechanism for tightly securing the load to the vehicle bed. Wooden blocking was used as required to prevent horizontal load shifting on the cargo bed.

During conduct of the actual field testing on the various vehicles, 105 mm ammunition boxes containing sand bags were utilized as the cargo load. This load was selected purely from a convenience standpoint. The boxes were upweighed to their normal shipping weight and center-of-gravity location. In an effort to approach realism in loading but yet establish some conservatism in this research effort, it was decided to load each vehicle to 75% by weight of its load capacity. A previous study had revealed that decreased load weight increased the dynamic response of the vehicle cargo bed - which is actually the input to the cargo load. By utilizing 75% loading, some conservatism is built into the schedules to account for variations in loading which will exist in the real world.

Data Points

Data were taken simultaneously at nine locations on the cargo bed of each vehicle that was used during the investigative work. Strain-gauge type accelerometers were used as the sensing instruments and were mounted triaxially at each of the nine measurement locations. These locations were on the structural members which go across the width of the cargo bed underneath the relatively thin-gauged steel cargo deck or floor. These structural members support the floor. The geometry of the instrumentation layout is depicted in Figure 2.

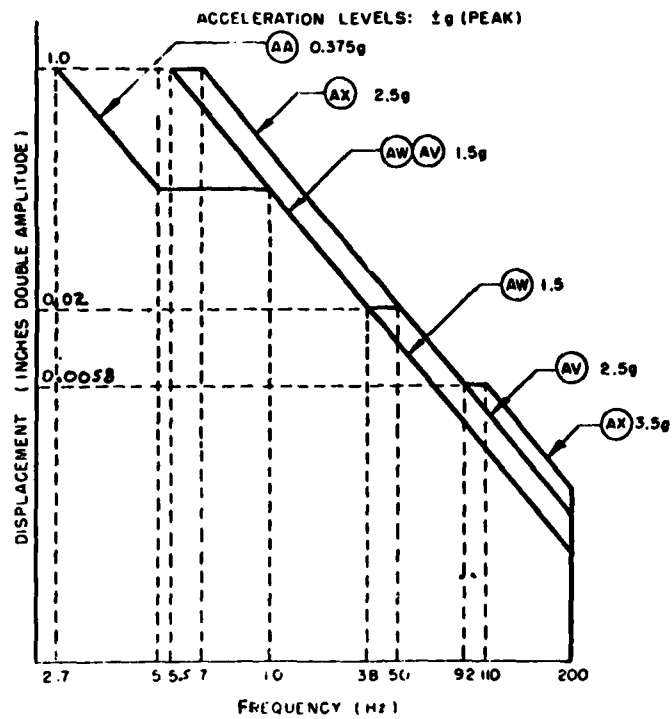
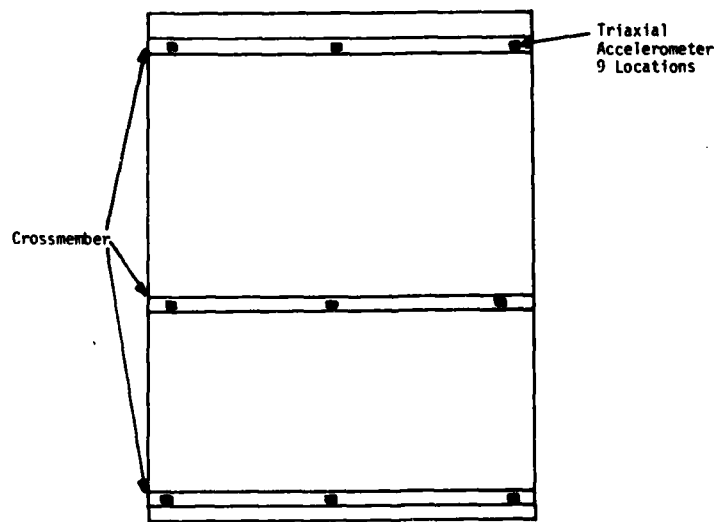


Fig. 1 — Vibration Test Curves for Equipment Transported as Secured Cargo, Equipment Category g



View from underside of vehicle cargo bed

Fig. 2 — Instrumentation Location

The instrumentation was mounted on the structural members instead of on the steel deck to avoid measuring the localized "oil-canning" response which usually occurs with a thin-gauged metal floor. Such a response is not indicative of the vehicle's input to the cargo due to the small amount of mass generating the force. The vibration environment which was measured was considered to be the vehicle's input to the cargo load.

Test Courses

The selection of the terrain over which to run the vehicles is as important as the selection of the instrumentation and load configuration. Running on a non-representative terrain ultimately produces a non-representative laboratory test schedule -- non-representative in the sense that it does not simulate real world conditions. An investigation was completed, under TECOM's overall plan, which addressed the establishment of the cargo transport scenario. This scenario identifies the various cargo transfer points (which segment the scenario), typical types of terrain in each segment, distance of travel expected on the various types of terrain, plus the types of vehicles that are generally used in the various segments. These results were verified with an Operational Mode Summary, which describes ammunition operational support concepts in foreign theaters. The investigation also established cargo transport distances within the Continental United States (CONUS) for transportation from the point of manufacture to a supply depot and from there to the port of embarkation. Likewise, the investigation established the intercontinental transport distances. The results of this work in established scenario distances are shown in Figure 3. As can be seen, CONUS transport is a maximum of 6,436 kilometers (4,000 miles); intercontinental transport (air or sea) is 8,045 kilometers (5,000 miles); and the foreign theater of operation mileage is a maximum of 856 kilometers (532 miles).

As the various segments of the total scenario from the point of manufacture to the using unit in the foreign theater were established, the same combined investigative work also defined the typical road profiles and transport vehicles within the segments. For CONUS it is generally the major highway system utilizing commercial tractor-trailers; and for intercontinental transport, it is by air or sea as previously mentioned. Once the cargo reaches the port of debarkation in the foreign theater, the type of transport vehicle becomes different, as does the road profile (or terrain as we previously referred to it). The various types of vehicles used to haul most cargo are depicted in Figure 4. This shows that the military trucks (typically the 6x6 or 8x8 running gear arrangement) and semitrailers are used from the point of debarkation to the forward supply point. The terrain is comprised of paved or improved gravel roads, along with a combination

of secondary roads, trails, and off-road conditions. The major portion of the terrain is the latter. From the forward supply point to the using unit, the transport vehicle is the 6x6 or 8x8 military truck towing the two-wheeled type trailer, with both the truck and trailer hauling cargo. As the figure shows, the two-wheeled trailer should be considered as the mechanism of transport within that segment since it produces a more severe vibration environment than the trucks. The physical size of the cargo items influences the selection of vehicle at this point. If an item is too big to physically fit in the two-wheeled trailer, then it would be transported by truck; and you would test to the appropriate laboratory schedule. The major portion of the road profile from the forward supply point to the using unit is the secondary roads, trails, and off-road -- all of which we will describe further on in this paper.

At the using unit, either the two-wheeled trailer (again, as limited by the physical size of the item) or the M548 tracked cargo carrier is used as the cargo carrier. The M548 is primarily an ammunition resupply vehicle, thus laboratory testing should be conducted accordingly. The major portion of the road profile of this final segment is also referred to as secondary roads, trails, and off-road.

The characterization of the road profile as secondary roads, trails, and off-road terrain implies that this is a rough terrain. A comparison of this terrain descriptor with the various test courses at the US Army Aberdeen Proving Ground led to the following correlation. Secondary roads can be depicted by the Cross-Country No. 1 course at Aberdeen which is made up primarily of gravel and has both sharp and sweeping curves. The road surface ranges from smooth to rough (roughness being due to potholes, washboard and rutting). The potholes and other sharp depressions are usually limited to a depth of 15 cm (6 in.). Depiction of off-road and trails was determined to be four of the test courses found at Aberdeen's Munson Test Area. These are the Belgian Block, Two-Inch Washboard, Radial Washboard and Spaced Bump courses.

The Belgian Block course is paved with unevenly laid granite blocks forming an undulating surface. It duplicates the rough, cobblestone road found in many parts of the world. The motion imparted by the course to a vehicle is a random combination of roll and pitch and high-frequency vibration.

The Two-Inch Washboard course is a 51 mm (2-in.) double amplitude sine wave course with the wavelength being 0.6 meter (2 feet). It depicts the washboard effect found on the dirt roads in many parts of the world and imparts a high-frequency chatter to the vehicle and thus a sustained high-frequency type of vibration environment.

The Radial Washboard course represents the

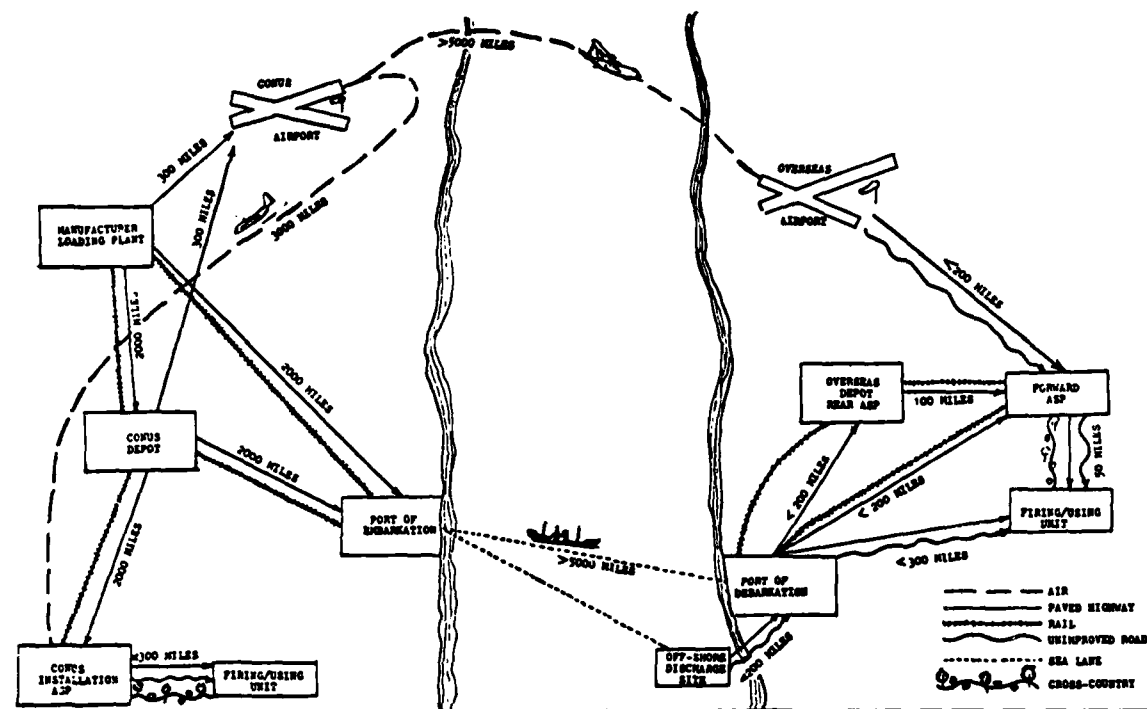


Fig. 3 — Movement of Ammunition from Manufacturer to User

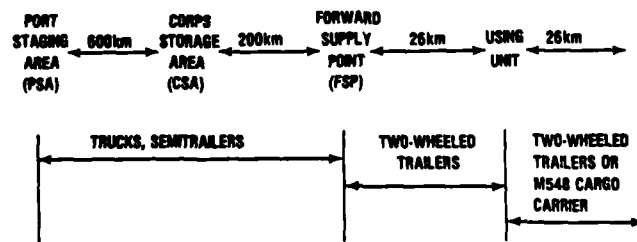


Fig. 4 — Foreign Theater Cargo Scenerio

washboarding found on curves of unimproved dirt roads. The wavelength varies from 0.3 to 1.5 meters (1 to 5 feet), and the double amplitude varies from 51 to 102 mm (2 to 4 in.). It produces various vibration frequencies on the vehicle.

The Spaced Bump course consists of 76 mm (3-in.) rounded concrete bumps that cross the concrete road surface at various angles. The spacing is designed to allow the vehicle's suspension system to "settle down" between bumps. This course imparts a combined vibration and mild shock environment to the vehicle's suspension system.

Before leaving this particular subject, we need to define the percentage of distance the various road profiles exist in the scenario. From the port staging area to the forward supply point, approximately 65% of the terrain consists of secondary roads, trails, and off-road; and 35% is primary road. From the forward supply point to the using unit, 70% of the terrain consists of secondary road, trails and off-road; and 30% is primary road. At the unit, only 10% is a primary road; and the remaining 90% consists of secondary road, trails, and off-road type terrain. These percentages were established by the Operational Mode Summary.

Data Reduction and Test Spectrum Development

A portion of the philosophy used in transforming field (or the real-world) vibration data taken on various transport vehicles into laboratory test schedules deals with using the proper data. It is known that the rougher types of courses produce a much more severe vibration environment on the cargo and thus are considered to be the predominant environment. This predominant environment becomes the one used in the development of the laboratory test schedules. The philosophy is that if the test item can withstand this severe vibration portion of the scenario, it can withstand the total scenario.

The differences in levels of severity for wheeled and tracked vehicles on various road profiles are shown in Figures 5 and 6. As can be seen in Figure 5, the actual PSD spectrum for an 8x8 truck on a paved road is approximately 90% less than the PSD spectrum for the same truck on the secondary road, trails, and off-road-type course. You will notice that the environment appears to be random, not sinusoidal.

The tracked vehicle, which is generally depicted in one's mind as running either on paved roads or on hilly cross country dirt courses, has different results in a comparison of peak PSD spectra as shown in Figure 6. Here the paved road surface is considered to produce the most severe environment. This rationale is based on the following: (1) The overall vibration environment of the cross-country terrain is no more severe than the paved road;

and (2) the paved road produces higher level and more frequency constant periodics. The latter indicates the driver can maintain a more constant speed on the paved road, and the resulting higher energy periodics will cause greater stresses on the material. You will notice this environment appears as a random signal with a fairly constant amplitude level (relatively low) and has superimposed on it higher level periodic amplitudes. These periodic amplitudes result from several things: the interaction of the vehicle track with the drive sprockets; the track contact with the road surface; the fact that by design the track on one side of the vehicle generally has one more track pad than the track on the other side; the inability of the driver to maintain a very precise speed; and the necessity to make minor steering corrections during road travel.

The periodic amplitudes are tested for randomness, and the results are shown in Figures 7 and 8. The test presently utilized is not a firm test for randomness. However, it serves as a reasonable indicator of data randomness. Figure 7 shows the results of this test as it was applied to a helicopter. If the data were a pure sinusoid, there will be no scatter in the data (either in terms of amplitude or frequency). Scatter exists when the average and peak data levels and frequency band are not the same. The blade passage frequencies (11, 22, 33 and 44 Hz in Figure 7) of helicopters were considered to be sinusoidal even prior to MIL-STD-810D. As Figure 7 shows, except for the 11 Hz frequency, there is no scatter in those frequency data. (The scatter in the 11 Hz data amplitude could possibly be noise on that data channel.) The same figure shows a difference, however, with the other vibration data in the total spectrum. There is data scatter which signifies it is not sinusoidal, but random. Figure 8, however, shows the high level periodics to have significant scatter both in amplitude and frequency as do the lower level data, which indicates that all of the data are random, some relatively narrow in frequency band while other data are of broadband frequency. Future plans are to develop a software program which will statistically determine the exact distribution of data.

This more severe vibration environment should be properly "weighted" according to the percentage of transport distance established for that terrain profile in the overall scenario. This "weighting" is done to preclude using the most severe environment for all of the mileage in the established test scenario. This will be addressed later in more detail. After the recorded data are verified as being valid (that is, not being one-sided, not having noise or frame errors, not being clipped, or not having a DC off-set), the reduction of data and the development of the laboratory spectrum begins. This is a computerized process which is somewhat dependent on whether the data are from a wheeled or tracked vehicle, and it proceeds as follows:

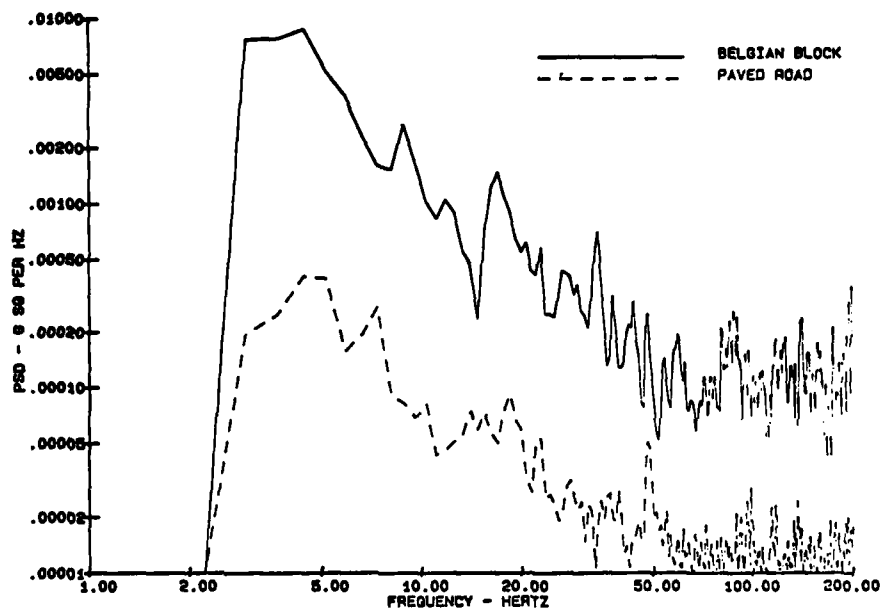


Fig. 5 — PSD of Truck on Belgian Block and Paved Road

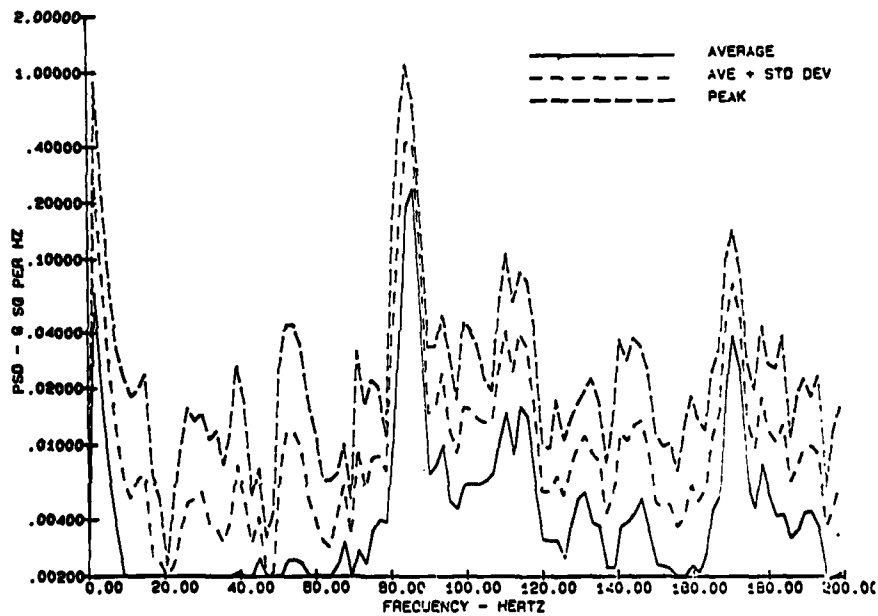


Fig. 6 — Typical Tracked Vehicle Data

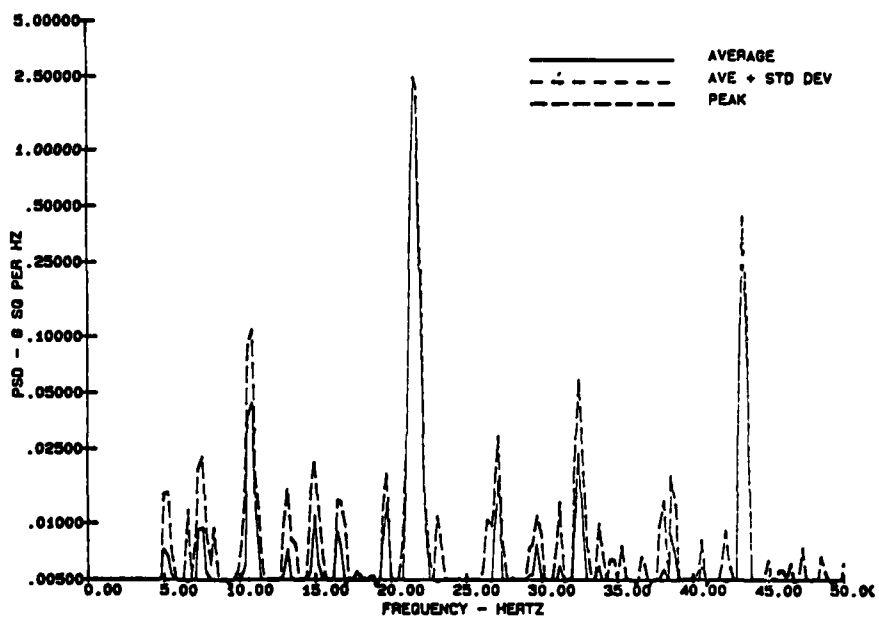


Fig. 7 — Typical Helicopter Vibration Data

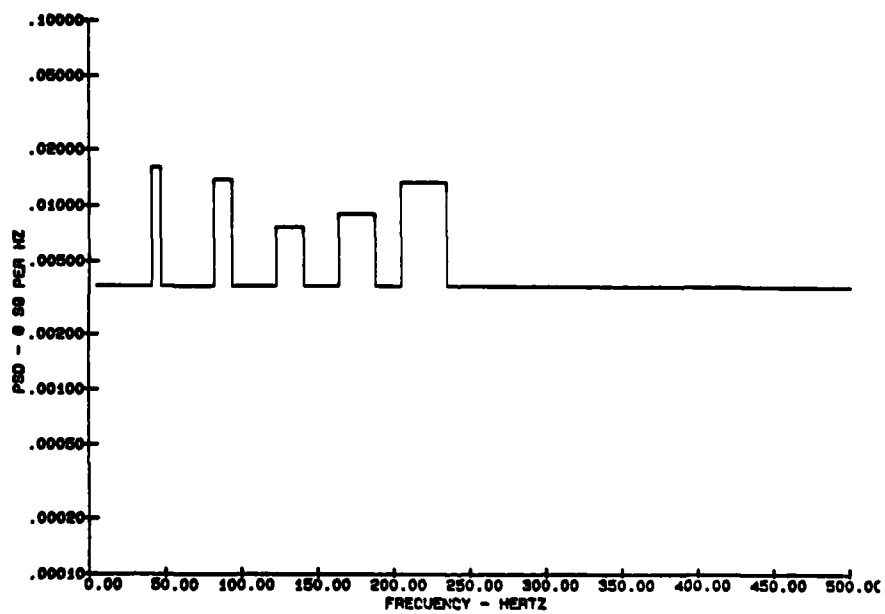


Fig. 8 — Typical Tracked Vehicle Test Spectrum

a. Wheeled Vehicles

For wheeled vehicles, the initial step of the computerized process is to analyze the raw (field) data from each channel on the vehicle during movement over the courses representing the secondary road, trails, and off-road conditions. This is done by making a PSD vs. frequency data file of each channel for each run. A run is defined as the vehicle traveling over one specific course at one specific speed (the courses are generally traversed at the maximum safe speed for the course, which experience has shown to usually produce the most severe vibration environment). The PSD's are computed by spectral line and 1,024 spectral lines are generally used. Two different PSD's are computed and saved; the peak and the average. In addition, the standard deviation per spectral line is computed and saved. The PSD used from this point on is the average plus one standard deviation. (This approach departs from past versions of MIL-STD-810 where the test spectra generally used the maximum values found in the data even if the value occurred only once.) Some conservatism is built in by adding a standard deviation to the average to account for the other terrain and vehicles for which no data were obtained, effects of tire pressure, vehicle wear, track tension, and so forth.

With all these PSD's computed and saved in files, the process now goes by data channel orientation, for instance vertical, as the schedule is developed for each orientation of the cargo bed. The selected PSD (peak, or average plus one standard deviation) for each location is overlaid with all other selected PSD's (from like orientation); and another analysis is made by spectral line to produce both the composite average and the composite peak PSD's. The standard deviation is again computed. The final PSD for the orientation becomes the average composite plus one standard deviation. This approach of combining the PSD's of like orientation for all of the data points on the cargo bed produces an averaging effect of the cargo bed input. Again, this tends to depart from the past maxi-max approach.

Once this final PSD has been developed, the final step begins in establishing the basic laboratory test spectrum. The PSD data are displayed on the terminal screen, and the engineer utilizes the cursor to produce a series of straight lines connecting the various breakpoints in the PSD (defined in terms of PSD amplitude (g^2/Hz) and frequency). As the real world data generally have many more changes in amplitude and frequency (see Figure 5) than laboratory digital controllers have the capability to control, the number of breakpoints is limited to 35 which is within the capabilities of digital controllers. By a careful selection process, the engineer can develop a meaningful test spectrum, one that encompasses almost all the real world data breakpoints in the frequency range below 100 Hz

and then carefully selected and characteristic breakpoints above 100 Hz. With most mechanical systems having natural frequencies in the range up to 100 Hz plus the demonstrated significant decrease in vehicle energy above 100 Hz, this approach provides the most meaningful laboratory test. As the number of breakpoints is limited, a smoothing effect takes place again providing a slight averaging effect. The final breakpoints are available on a printout to provide a more accurate spectrum definition for the laboratory engineer.

Concern has been expressed by a few that there are too many breakpoints in the wheeled vehicle spectrum, thus making it too specific. If one looks at enough wheeled vehicle data, it can be seen that this multipoint spectrum is always present, and thus the laboratory test spectrum philosophy of multiple breakpoints is realistic and is attainable. No data in the file look like the past "straight-line and less than six breakpoints" test spectra shown in Figure 1. Such straight line approach contributed significantly to the maxi-max philosophy of old.

At this point, the engineer addresses whether exaggeration factors are applicable in order to reduce test time. Before we look at that, let's turn to the data processing technique for the derivation of a laboratory test schedule for a tracked vehicle.

b. Tracked Vehicles

The initial process for tracked vehicles is the same as for wheeled vehicles -- verify that the data are valid, then compute the peak and average PSD's per spectral line per channel per vehicle speed plus the standard deviation. Again, the average plus one standard deviation is used in the further processing. At each road speed, all like orientation channels on the tracked vehicle's cargo bed are overlaid. The composite PSD's are computed along with the standard deviation, and the average plus one standard deviation is used. At this point there is a composite of all like orientation channels at each speed increment. The data appear as the relatively constant amplitude level superimposed with a higher level fundamental periodic and its harmonics. It is necessary to process the data in the aforementioned manner (or per speed increment) in order to maintain the relationship of the periodic frequency and its harmonics.

For each composite PSD, the periodic components are defined by an operator-controlled cursor from the data plotted on the terminal screen. The center frequency and associated acceleration amplitude (g^2/Hz) of each periodic and harmonic are computed along with PSD level of the relative constant amplitude level (with the periodic components removed).

Finally, the computed periodics and associated PSD levels are printed in groups of harmonics in ascending order. The average,

peak, and standard deviation for all the random values are computed; and the average plus one standard deviation is saved.

At this point, the current limitations of digital vibration controls dictates the method of development of the final test spectrum. As the current software can only accommodate five dynamically changing narrowband random spectra at various rates across a broadband spectrum, several test phases have to be developed to encompass all vehicle speeds. The printed information of periodics and harmonics are manually placed into narrowbands with the total width of each narrowband being chosen so that none of the succeeding narrowbands are overlaid in frequency. The periodic and harmonics for each vehicle speed must be present at the same time as this is what occurs in the real world. The final spectrum will look like that shown in Figure 9.

Exaggeration Factor

The use of exaggeration factors often becomes an area of real concern for many people as it implies an overtest. Exaggeration factors are used to reduce laboratory test time by increasing the amplitude of the input vibration spectrum. This is done to make the laboratory test time something that is manageable and to permit cost-effective testing in the laboratory. It is agreed that exaggeration can result in an overtest, but only if exaggeration factors are used incorrectly.

The most generally accepted theory underlying the use of exaggeration factors is Miner's theory; and it is well documented in various texts including the Shock and Vibration Monograph Series, SVM-8, entitled "Selection and Performance of Vibration Tests." Although several theories of cumulative damage exist, Miner's theory is probably the one that is the most universally applied because it is relatively simple and is as accurate as any. The rationale for using an exaggeration factor to reduce test time is based on the cumulative damage theory, assuming that the failure mechanism is fatigue.

Miner's theory takes into account both the endurance limit and damping characteristics of the material. The relationship of the real world and laboratory test spectra for a random test is:

$$\left(\frac{W_1}{W_2}\right)^{b/n} = \left(\frac{T_2}{T_1}\right) \quad (1)$$

where W_1 = amplitude of the real world environment

W_2 = amplitude of the laboratory test schedule

T_1 = time duration of the real world environment

T_2 = time duration of the laboratory

schedule

b = endurance curve constant which ranges from 3 to 25 with a representative value of 9 being used for many structural materials.

n = damping constant which ranges from 2 to 8 depending on the material and stress level. Generally for stress levels below 80% of the endurance limit of the material, $n = 2.4$.

$$\left(\frac{W_1}{W_2}\right) = \text{the exaggeration factor} \quad (2)$$

As the derivation of the vibration test schedules in MIL-STD-810D was intended to be for general use, the constant values of $b=9$ and $n=2.4$ were used.

The philosophy used in applying exaggeration factors in the ground vehicle simulation is that the exaggerated levels do not exceed the peak values which occurred in the field by more than 25%. (This is on a spectral line basis.) This value of 25% was established to provide a reasonable limit on the amount of exaggeration while providing for a manageable laboratory test time. The supporting rationale for this philosophy is that we ran only one vehicle, one time over the courses, with one driver. The peak data measured are for that one condition only; the 25% exaggeration permits a reasonable allowance for peak data which could be higher another time the vehicle traverses the courses.

The exaggeration factor is applied to each spectral line which increases the overall rms value of the spectrum by a factor of the square root of the exaggeration factor.

The exaggeration factor is applied to reduce test time in the laboratory and is used to determine the laboratory test time, to which we now turn our attention.

Test Times

Ideally, most test engineering personnel (and certainly all developers and program managers) would like to apply the field spectrum on a real-time, non-exaggerated basis to the test item in the laboratory. In some instances, such as missile flight, this is possible due to the short duration of the vibration environment. However, when we start to consider several hundreds or even thousands of kilometers of ground transportation, plus many test item samples on many programs, the use of real time is a practical and cost-effective approach. Thus we utilize the exaggeration factor that we have just addressed in order to reduce the laboratory test time to a manageable level. What is a manageable level? There is no definitive answer -- it really depends upon the philosophy of the particular test laboratory. Our philosophy in developing the test curves for MIL-STD-810D was to arrive at a test time of

approximately 2 hours which would permit vibration of three test loads in an 8-hour day.

The determination of a test time is more complex than it appears since three parameters must interplay to coincide with the philosophy of a manageable test time. These parameters are the exaggeration factor required, the desired test time, and the portion of the scenario used for the laboratory test (which is not necessarily the entire scenario). This last parameter can be the most confusing so we will direct our attention to it first.

As previously mentioned, the secondary road, trails, and off-road profiles produce the most severe vibration environment on wheeled vehicles. Similarly, paved road produces the most severe vibration environment on a tracked vehicle. The data from these terrain profiles become those utilized in developing the actual laboratory vibration test spectra. In developing test times, the percentage of the total distance the vehicle is expected to be on this most severe type of terrain must be known. These percentages were delineated earlier in this paper.

Utilizing the scenario given in Figure 4 as an example in developing the test time for the composite wheeled vehicle, we use the vibration spectrum from the secondary road, trails, and off-road terrain area which constitutes 65% of the mileage from the port staging area to the forward supply point. We arbitrarily selected 2 hours to be manageable laboratory test time. Additionally, we know from the data reduction/transformation process that the actual field vibration level for the composite wheeled vehicle was 1.55 g rms.

To arrive at the exaggeration factor, we use the relationship derived by Miner:

$$\frac{W_1}{W_2} = \frac{b/n}{T_1} \quad (1)$$

where W_1 = 1.55 g rms
 W_2 = unknown
 T_1 = 800 km divided by the average vehicle speed of 26.15 km/hr (as determined from average actual vehicle operation on those courses we've described)
 = 30.7 hrs. x 65% of time on that terrain = 19.89 hrs.
 T_2 = 2 hr
 b = 9
 n = 2.4

The computations produce a value of 2.87 gms for W_2 , or an exaggeration factor of 1.85. At this point, the actual environment is exaggerated by 1.85 (spectral line by spectral line), and the proposed laboratory spectrum is then checked to ensure it does not exceed the actual peak environment by more than 25%. As it was determined that the 25% limit was not exceeded, this exaggeration factor was

acceptable; and the test spectrum and test times were established.

Testing the Laboratory Spectrum

Our development procedure now utilizes a program which computes the required values of the various vibration exciter parameters (velocity, displacement, and acceleration) necessary to accommodate the laboratory spectrum. These values are then compared with the specifications of the exciter. If the exciter cannot meet any of these requirements, the exaggeration factor has to be reduced (increasing laboratory test time) in order for the parameters to fall into the range of exciter capabilities.

What Does All of This Mean?

Now we have derived the laboratory vibration test schedule which we will use to test items of cargo. But what does it really mean? Again, the philosophy is that if the cargo in a wheeled vehicle can withstand the more severe vibration environment of the secondary roads, trails, and off-road profiles in the foreign theater, it can withstand transportation over the entire scenario from the point of manufacture to the user in the foreign theater. This is based on the S/N curve which shows the higher the stress level, the lesser number of stress cycles required to reach the endurance limit; or simply, you use up the greatest portion of the endurance limit with the high stress test. In our example, the 2-hour laboratory test simulates the entire 800 km scenario.

What else does it mean? It means that those individuals involved in the design and development, and even testing, of equipment have to establish the scenario that the equipment must withstand. A lot more thought needs to go into test planning in order to subject the equipment to the most realistic test. The thrust of MIL-STD-810D is to identify the real scenario, measure the real world environment associated with that scenario, develop the laboratory test schedule simulating that scenario, and then -- and only then -- to conduct the test. We call it tailoring. It departs from the old traditional cookbook testing of previous versions of MIL-STD-810 where you merely "pulled out" a curve and used it. Tailoring provides for a better and much more realistic test. With that in mind, you may ask why any schedules appear in MIL-STD-810D. They are there for you to utilize if: first, you find these schedules match your scenario; or second, you do not have the resources to measure the environment of your scenario and are willing to take a chance on using these schedules. You will have to make this risk assessment.

Current Limitations

The current limitations on the use of the new approach are two-fold. One is that some

test laboratories are not properly equipped and must update their vibration controllers to the digital type in order to accommodate the new schedules. This involves funds and managerial direction, and it may or may not be serious depending on the size of the organization and the associated funding constraints or allocations. The other limitation is that of the current software for the digital controllers. The software currently does not adequately accommodate the swept narrowband-random-on-broadband-random schedules related to tracked vehicle vibration. This is a very serious problem since the only presently known software capable of doing this is available from only one source, and it can be used with only one particular digital controller. Although it works, it is operationally limited. It is known that several suppliers of controllers/software are currently working on this problem. It must be resolved, and completely adequate software must be developed to meet this need. Why then did we develop these schedules if we knew the controller industry couldn't properly handle it? The answer is simple. We felt it was time that the development of the laboratory vibration schedules required to realistically simulate the real world should drive the state-of-the-art in equipment/software design as necessary and not vice versa as has been the case in the past.

Summary

In summary, development of a laboratory vibration test schedule is not a simple task. To do it properly, the established philosophy must be understood and followed. The development involves identifying the real scenario (distance, terrain, and vehicles); selecting the data measuring points; acquiring the real world data while following the scenario; reducing and translating the data into the laboratory test spectrum; exaggerating (but not over exaggerating) the levels as necessary in order to provide for a suitable test time in the laboratory; and having the hardware/software to conduct the test. It is not an easy undertaking. It is one that has to be recognized as important by the managers and engineers at all levels of development, program management, and testing in order to provide the best and most cost-effective development process for the equipment we all expect our soldiers to use effectively in the field.

PANEL DISCUSSION

MIL-STD-810D

Moderator: Mr. Rudolph H. Volin, Shock and Vibration Information Center, Naval Research Laboratory, Washington, DC
Co-Moderator: Mr. Wallace W. Parmenter, Naval Weapons Center, China Lake, CA
Panelists: Dr. Sheldon Rubin, The Aerospace Corporation, Los Angeles, CA
Dr. Allen J. Curtis, Hughes Aircraft Company, El Segundo, CA
Mr. David L. Earls, Air Force Wright Aeronautical Laboratories, Wright-Patterson Air Force Base, OH
Mr. Howard C. Allen, Rockwell International Corporation, Downey, CA
Mr. John A. Robinson, U.S. Army Test and Evaluation Command, Aberdeen Proving Ground, MD

Mr. Mardis (General Dynamics): Dr. Rubin, regarding the tolerances on shock response spectra type tests, they seem a little tight, particularly the minus 0 dB tolerance bound. I have found in my past few years of doing electrodynamic shaker shock tests I have trouble matching the exact boundary conditions of the structure where the measurements were actually made. Therefore, I am not able to drive the structure that I am trying to test at all frequencies no matter how hard I try. I am using some proprietary equipment by a manufacturer to do this. They have some routines in their equipment to provide this. If I have to use a minus 0 dB tolerance, it looks like I will have to boost my whole spectrum up maybe 10 dB or so higher than I would wish. Before MIL-STD-810D came out, I did not have many guidelines to go by, so I established that I would permit myself a hole I could fall into over some portion of an octave across my frequency spectrum to fall out of tolerance on the lower end. Can you suggest how I can get around this tolerance problem? It is tough. I don't think I can meet the plus 6 dB, minus 0 dB tolerance for pyrotechnic shock at high levels on the shaker head much less on the unit under test.

Dr. Rubin: I agree. That is a tough problem. In my experience a lot of it has to do with the levels you are dealing with. If you can perform the test on an electromagnetic shaker, the levels are sufficiently low to do that. I have seen cases where the test control is certainly within that band. It is possible, but I am not saying it can always be done, and I'm sure it is specific to your particular set-up. So it is hard to make any general statements. In a

standard, however, we cannot foresee all of the exceptions and all of the situations that might arise where control may not be what one desires. In that case you have to work with your customer and convince him that you are doing the best you can to meet his requirements, and work out what deviations he will permit. I am afraid this is the way it really will work.

Dr. Curtis: First of all, I would like to comment with respect to the previous question. It seems to me it would have been better if the tolerance on the shock spectrum in some way permitted satisfying the requirement when a composite spectrum of the spectra existing at various mounting points met the requirement. This is somewhat analogous to power averaging several accelerometers in a vibration test, and for the same reason. We should be able to do the same thing in a shock test because if some attachment point is a node in vibration, it will still be a node in shock no matter how well you design your fixture. We need a general way of getting around that because you just cannot go back to the customer every time you have to run a shock test. I think we are kind of to that point. But Shel, as I listened to you this morning, I think what I heard you say was that every existing piece of software for performing shock tests on shakers is now obsolete because, to my knowledge, there is no software generally available which calculates both positive and negative spectra. There is no software which allows you to edit, to doctor, or to do whatever you have to do to your digitized time-history, to play around with or measure this equivalent duration that you described. I hope that is not what I heard.

Dr. Rubin: If the software does not permit looking at positive and negative levels, that does not make it obsolete. It means you will not be able to take advantage of the possibility of covering both directions in an axis with a single test. You will have to do what was done before, and test in each direction independently. So I believe, you have an incentive to get software with that capability as soon as possible because it will save you some testing. I cannot argue that the software does not have the capability now, but I think, as in many other areas, new capability will be required. I think this is a very modest capability to be able to deal with positive and negative responses in the same shock spectra analysis program. It does mean changing the program. In that sense it is not modest, but I believe, the degree of change we are talking about, although it is very significant, it is very modest. It can have a payoff in terms of the testing. So, it will not prevent you from continuing to use the old software; I think it will prevent you from taking advantage of one of the possible ways of reducing the number of shocks to meet the requirement.

Dr. Curtis: Either that or just switch the polarity on the accelerometer signal and recompute it.

Dr. Rubin: Well, if you can sell your customer on that, good luck.

Dr. Curtis: The reason for switching the polarity on the accelerometer signal and recomputing the shock spectrum is when you have the kind of wavelet type excitation which is in the present software, it is inherently symmetrical, but I cannot prove it. Furthermore, most of the software is based on about a 200 millisecond transient window. The transient will tend to occupy a good percentage of that window. Will that invalidate it? It will not look as pulsey as the picture you have.

Dr. Rubin: If the software has a fixed analysis window of 200 milliseconds, and there is no control over that, and if the standard says that there is some relationship between the analysis time and the effective duration of the field shocks, then that is an incompatibility. However, I would think you would be able to get around that, possibly on the basis of being able to demonstrate in some other way that you haven't introduced excitation over a longer period of time than the standard intended. But, the burden of proof will be on the tester to demonstrate that in some way. So, I think there is some way around it, but it will require some work to do it. The purpose of putting things in specifications or standards is to point out what you would really like to achieve and to avoid test conditions which are not intended. If the software or the hardware do not permit that, then there is a matter of convincing whoever has to write-off a deviation on the standard that you have met the intent. I think it is

important that you have those kinds of identification matters within the standard to point out what things you have to show. So yes, it could be a problem if you have software with a fixed capability. It will not meet the tailoring needs of the standard. What I hope will be achieved is that the software will be modified so it can handle the tailoring requirements in the standard. But I do not think that we could have in any way met the intent of a standard that had tailoring capabilities and stayed within the limitations of current software.

Mr. Smith (Hughes Aircraft Company): Dr. Rubin used a word that really pointed out the thing I am worried about. He used the word "deviate". According to the gist of the standard, if an environmental engineering specialist feels that this test with its vagaries, not deviations really, is something that meets the intention of the standard, and he really feels that the item was fairly and correctly tested to the best of his ability, then the question is what happens at that point? It is not really a deviation. Yet I know I will not feel comfortable at that point that the unit on the shaker really passed the test, and I can really take it off and not worry about it. What will happen? Is there an actual deviation, or opinion process, when these sorts of things happen? I just cannot believe that no matter how well I feel about it, that my say will really pass that part.

Dr. Rubin: I don't think I can answer your question. But if I understand the gist of it, when you run a test and something in the test has fallen outside of the limits in the standard, how can you know at that time whether it will be bought off, or whether you will be able to convince somebody that it is acceptable?

Mr. Smith: Even though you can rationalize it, and you feel reading Part III, or whatever explains the gist of the test, you have complied with that. What happens then? I am sure this will come up. In fact isn't that really the whole purpose of this standard to protect good hardware from these kinds of problems that really don't disqualify it?

Mr. Rubin: Yes, it is certainly part of it.

Mr. Smith: This is something that will have to be ironed out right away because that will be the first test of the standard when it happens. By the way, when will MIL-STD-810D go into effect? How soon will we be confronted with testing to MIL-STD-810D?

Mr. Earls (Wright Aeronautical Laboratories): When you get a contract that calls for the use of MIL-STD-810D after 19 July 1983.

Mr. Smith: So, has there been a test case yet?

Mr. Earls: It is out now, so it is effective on the date of the contract.

Mr. Smith: No, I mean what happens? Someone might say, "I passed the test even though I am out of specification. I have good reason to feel that is not important". Has that issue come up yet?

Mr. Earls: Is that any different than things that have been happening in the past?

Mr. Smith: Yes. Now it is blessed.

Mr. Earls: I don't think MIL-STD-810D has done anything new in that area. You might cite a problem with MIL-STD-810D that is in an area that you didn't see before; but there are individual contractual test problems that you have to solve at the time they come up, which has been prevalent all the time. It is between you and your customer.

Mr. Smith: It would seem from the wording that you have put that in a different area.

Dr. Rubin: In those cases when you are familiar with your test capability, and you know in advance that there will be some departure, I would think it would behoove you to discuss that in advance of the test and get some understanding of what will be allowed. That means the only thing that will come up during the test is something you didn't anticipate. Those are the problems that will have to be worked out at the time they arise by whatever mechanics that are set up for that purpose.

Mr. Galef (TRW): There are some things that have been in standards and specifications for so long that it is almost blasphemous to question them, but I will do it anyhow. One of them is the three shocks in every direction; why three? Why not one? Why not thirty? Another thing is the Q of 10. Why 10? It is not bad, but why?

Dr. Rubin: The three shocks in each direction has been in MIL-STD-810 historically, and I don't know where that particular number came from. I do know, in terms of the pyrotechnic shock testing that we have been involved with for aerospace applications, that the number of three was just an engineering judgement call because of the variability in creating that kind of environment with the way those tests are conducted. The scatter from test to test alone is such that you might even like to run more than three tests. You said that if you can't beat the tolerance, then it is a non-test anyway; so you will have to keep doing it over and over. The idea is to achieve the intended level of severity called out by the test three times, but there is no way to justify any one specific number across the board. It is a trade-off between at least some repetitions versus calling for so many tests that it becomes a very, very expensive item. The three just happens to be the number. I agree; I couldn't justify four versus three and so forth. With the matter of a "Q" of 10, again, there is no

way to come up with a single value of "Q" for damping that will apply to all hardware. The value of 10 is a reasonable number that lies somewhere in the range of what you typically find. I think it is good not to have the value of "Q" too high because then it makes the control and the generation of the test much more difficult. I am comfortable with a "Q" of 10 for general applications. On the other hand, if that seems inappropriate for some particular item of equipment which may be very lightly damped, again, that would be a departure that could be worked out. In other words, there is nothing that will prevent the use of a different value of "Q" if the proper arrangements are made. The important thing here is the value of "Q" used to analyze the field data should be the same as the value of "Q" that is used to control the test, so that they are compatible. It is not fair to change in midstream. That would be a possibility for change if an argument for doing something else can be made.

Mr. Rosenbaum (General Dynamics): I notice we finally have some fatigue relationships in the MIL-STD-810D using four for random vibration and six for sinusoidal vibration. I had a preliminary copy of MIL-STD-810D which used a material constant of eight-and-a-half on an overall g-RMS basis. Where did these numbers come from? Why were they used? Why was there a change?

Mr. Allen: I don't know how it got to eight-and-a-half.

Mr. Rosenbaum: How did it get to four and six?

Mr. Allen: We haven't nominally used that, however we have used the basic fatigue slope relationship, and we have nominally used four on the Apollo and the Shuttle programs. I don't know how this time equivalency relationship got to be eight-and-a-half. I would not attempt to justify that.

Mr. Galef (TRW): There has been fatigue in MIL-STD-810 since the C version came out over 10 years ago. As soon as they put the random vibration in, they also put in a method for accelerating the test. At that time they used a factor of four, that is an inverse slope of four for random vibration. This, by the way, is equivalent to a slope of eight for sinusoidal vibration, and I am very bothered by this inconsistency in the same paragraph of four for random vibration and six for sinusoidal vibration. Can you speak to that contradiction?

Mr. Earls: I don't have the answer to that, but somebody did go over that, and they said the eight was not equivalent to the four, and that is why it was changed to six. So, that is a mistake in the document if you are right.

Mr. Himelblau (Rockwell International): Let me shed a little bit of light. When we started the

process of deciding what slope to use, I got a hold of as much data as I could on random testing, mainly on 2024 aluminum. Some other metals were also tested, 7075 aluminum, etc. But the data that we found with the steepest slope came from some Langley tests by Clevenston and Steiner with a notch concentration factor of four. This is the data that we measured the steepest part of the slope on the S-N curve, and we came up with this as the slope coefficient. Since we were not involved with sinusoidal vibration, we didn't even look at any of the sinusoidal vibration data.

Dr. Curtis: I'd like to get back to shock for a moment. If I could hark back a few years, it seems to me that when we were first sort of adopting shock spectra and getting into that business, it was adopted on the basis that shock is sort of an ultimate strength kind of a damage producer, rather than a fatigue damage producer. Therefore, if you get up to that load once and nothing untoward happened, you would probably get up to that load many times and nothing untoward would happen. Therefore, if I do the proper test, I probably only have to do it once. Then we got into the discussion of the characteristics of whether the shock spectra were symmetrical or tended to be asymmetrical. If they were symmetrical, we would conclude that one test in one direction was adequate. But, if the shock spectrum was asymmetrical, and we wanted to make sure we loaded it in both directions equally severely, or severely enough, then we had to do it once in the positive direction and once in the negative direction. The other thing was that the shock spectrum was a measure of the damage potential, and we argued that if all the resonances were excited adequately, we could use the shock spectrum approach which discards any time characteristics like phase or the duration of the pulse, and so on. We even had equipment which would produce chirps which lasted a considerable length of time, but which could produce the shock spectrum. I thought everyone pretty well agreed to that, and we went down that road for awhile, but now I see that we are going back, and we are saying, "Yes, that is okay, but we want to put, if I may say so, unnecessarily severe restrictions on the time history which will cause a great deal of pain, grief and cost in implementing for a worth that I have difficulty appreciating.

Dr. Rubin: Let me go back to the sinusoidal versus random vibration. The argument for random testing is that you excite all of the resonances simultaneously, and any interactions or additions in stresses or loads that come about will thereby be more realistic. I think this same argument can be made here for shock. If you do fast sine sweeps, or chirps, and excite the resonances in some sequence of frequency order, and through the shock spectrum approach, basically excite each of those resonances to the maximum value, you are not getting the "possibility" of some damage

resulting from the superposition and simultaneity in the responses of your modes. I liken that to the same reason for doing random vibration testing versus sine sweep. I have never liked chirps, by the way.

Dr. Curtis: I don't either.

Dr. Rubin: It is a point. This was a definite concern. There were several ways of going about it. One of the ways of going about it is to say that you have to meet a shock spectrum requirement over a range of "Q". For example, you can specify three specific values of "Q", 5, 10, 50 or something like these. By saying that you have to meet a shock spectrum requirement for each of three values of "Q", means that you are quite limited in the kind of waveforms that can be used to meet the specification. That is a way of controlling it. It seemed to me that introduced an awful lot of complexity, and basically the same intention could be met much more simply if one constrained the time to be reasonable. So, with those two choices I selected what I thought was the simpler approach.

Dr. Curtis: If my test requirement is based on a single measured shock spectrum, admittedly I have thrown away phasing information in calculating that spectrum. Now, I appreciate your point very well, but my shock spectrum will most likely be some kind of envelope average percentile of a number of spectra, each of which will peak, and therefore the peak of the spectrum will not be reached simultaneously. In fact, it wasn't in the same place in each event on which I based this test. I would just like to suggest that maybe we are overdoing it a little bit for general application.

Dr. Rubin: Again, I see the same situation. If there is a doubt in terms of severity, one has to play it on the conservative side. I think the same situation applies to the random vibration test where specifications are generated on the basis of a multitude of spectra and one creates envelopes and percentile curves. The situation is the same. I don't know how to get around it. If you are going to perform a test to a single spectrum requirement, you have to cover all the possible bases. Again, I see a parallel between what is done in random vibration and in shock.

Mr. Hancock (Vought Corporation): One word particularly interests me, and that is the use of the word specification. I guess in the case that you are talking about, where this does become a specification, then what you are saying has total merit. Yes, we do have some sorts of problems. But, if I read paragraph four correctly, I, in my environmental engineering wisdom, knowing that this is not an adequate representation of the real world, am beholden to change the requirements prior to the time it becomes a specification that is written into the

contract. But the requirements are still based on the standard. Am I correct?

Dr. Curtis: You will have to remember when you do your environmental engineering, and come up with a good shock spectrum, which you hand off to the lab, you will have to remember to tell them that all this stuff about how spectrum package does not apply in this case. If you forget, they cannot do the test. So, you do have an out if you remember properly.

Mr. Hancock: A while ago, someone else brought up the question about how this gets changed. I believe the proper time, according to paragraph four, is when we submit the environmental test plan for approval.

Mr. Strauss (Rocketdyne): Dr. Curtis, you mentioned before that we throw out the time-history when we submit the shock response spectrum for testing, or that we don't use it. What we have is a specification which is only the shock response spectrum; one way to verify that we are doing a good job is to keep the time-history. When you write-up a requirement, or when you get field data, you keep the time-history. That way we know if it is really a damped sine-type shock wave or if it is a short duration pyrotechnic shock. I think we could keep the two together and use it as a basis for a better test. The other thing I wanted to mention is that there was a presentation this morning by someone from Sandia who said they are working on a way of coming up with superimposed damped sines to simulate the same type of thing; maybe they might come up with a different type of test requirement other than the shock spectrum. I think if the time-history were really based on the signal that looks like a damped sine wave, we could possibly simulate it in the lab by a series of damped sine waves. But many tests have been performed that show that there is a definite difference in damage potential using a shaker type damped sine shock from a pyrotechnic simulation. I think if we do not use the 20 millisecond duration on the shock spectrum analysis, then we have a large difference in damage potential in our tests if we are trying to simulate the same thing.

Dr. Rubin: I think what you are mentioning now is getting to the point that Aller made earlier in having to do with the correlation of the inputs to the test article. This has to do with sine waves and so forth. This is a problem, and I am not saying that the standard solves it. But again, this problem exists for all vibration testing, and we have not solved it there. It is an issue that is there, and we will have to make our best engineering judgement on it. There are no magic solutions to it. With regard to saving the time-histories of the field shocks, I certainly do not object to that. In fact, by asking that an effective duration be identified, we are asking for some very specific information other than the time-history. But at least we are asking for the duration of the shock so we

do not lose sight of the duration the significant portion of the shock; then we can try to do the same thing in the laboratory. I think saving the entire history as a data bank is always a good idea.

Mr. Volin (Shock & Vibration Information Center): I have a few observations that I noticed in the discussion this morning. The first one is we have a criteria of 100 inches per second for a waiver. I am wondering to what extent people try to get around doing the pyrotechnic shock test by using a waiver. Another problem I have noticed is the question of how you really know that you have valid pyrotechnic shock data, because I have heard too many times that it is a question of do we really know what we are measuring. As a matter of fact, at one meeting I heard somebody say, "What you are really measuring is the natural frequency of the accelerometer." Getting a little closer to earth, we have one test procedure in MIL-STD-810D that has been carried over from MIL-STD-810C and probably from some of its earlier versions, and it is the rail impact test. Is there any way to simulate that in the laboratory? From what I can see, there is a free-fall sled apparatus for testing packages. Of course, one could actually go to a railroad yard and conduct a test that way. But that gets to be expensive, and it is not a simulation; it is the real thing.

Dr. Rubin: The rail impact test? I really have not looked at that requirement in terms of other ways to perform it. The other point you brought up was the matter of waiving shock tests. I have run into it a couple of times; it doesn't come up very often. If you have a piece of equipment that is relatively far removed from a source of pyrotechnic shock, the levels can become relatively mild compared to other requirements. There has to be some way to cut off the need for testing, otherwise you are faced with, if you are on any kind of a vehicle or platform that has that kind of shock, a requirement to test for it regardless of what you can demonstrate in terms of severity. So, I personally feel there should be some way of waiving the requirement if you can demonstrate that you are covered by some other test. The question of the 100 inches per second - that is just based on the best information that I have been able to find, and it is based largely on some Navy experience. It is not a tremendous amount of experience, but we had to pick something. I picked that number on the basis of what I was able to find, and I put a factor of safety of two on it. I am hoping by identifying it in this way, that maybe some more information will come out of the woodwork, and we will be able to justify some better numbers in the future. But, it is a start, and it represents the best information, at least that I could find.

Mr. Davis (Ford Aerospace): I have a question on the application of the narrow band random

vibration on random vibration. The way random vibration is usually applied in the laboratories a roughly normal distribution of peaks is assumed, with a peak to RMS ratio of three. Are you assuming this same distribution for narrow band random vibration on random vibration?

Mr. Robinson (Aberdeen Proving Ground): I am not sure what the assumption is. I think it might be a trade-off between what the field data show and the capabilities of the existing test control software. I really don't have a firm answer for you on that.

Mr. Davis: I asked this question because, while I have not done a statistical evaluation of it, I ran some of this data through a very narrow band pass filter centered right at the track laying frequency. The resulting data look like an approximately constant amplitude sine wave. In that case, you have a very different peak to RMS ratio than you would have with a random signal at approximately the same frequency. So, I wonder if the appropriate peak to RMS value and the appropriate way of simulating random vibration on top of random vibration need to be defined.

Mr. Robinson: Yes. We are attempting to look at the data you are talking about now, at Aberdeen. This preliminary look at the data, to verify whether it is actually a random signal, or whether it is indeed sinusoidal, has revealed that if you look at tracked vehicle data over a long enough period of time, (I am not talking about minutes but a definite period) you will find the amplitude will vary because the vehicle cannot maintain a constant speed. There is no way a tracked vehicle will maintain a constant speed because there are just too many parameters; one parameter is the driver. He just cannot control the vehicle as closely as he can control a wheeled vehicle. That is why the tendency is to become a random type of environment, although very narrow band, rather than a sinusoidal type of environment.

Mr. Norris (Martin Marietta): I have two general questions for the panel on random vibration control strategy. Both questions pertain to digital control, multi-channel, extremal, or peak response strategy. First, what is the technical legality of weighting one or more of the control accelerometers in the driving direction? Second, what is the technical legality of incorporating cross-talk accelerometers in the control loop? There seems to be a trend toward doing that. If you put cross-talk accelerometers in the control loop, when the cross-talk accelerometer takes control, the input in the intended test direction is down. You will not meet the test specification. It is happening; people are doing it. There is nothing in MIL-STD-810D or C that puts a limit on what you can do. So people just take liberty with whatever they have to do to get through a test. I think MIL-STD-810D

should have had some kind of limits on it, or some kind of guidance on it.

Dr. Curtis: Since the beginning of time specifications cannot ensure integrity.

Mr. Norris: How about weighting of one or more of the accelerometers?

Dr. Curtis: Are you talking about a response control test?

Mr. Norris: I will give you an example if you consider testing captively carried stores.

Dr. Curtis: That is a response control test.

Mr. Norris: It tells you to put your input in six dB down, measure your responses, and where the responses are greater than six dB above the input, notch the input. However, if you run extremal control, you really don't have to go through that. The way it is being interpreted is there is an accelerometer on the captively carried store for response and an accelerometer on the fixture input. They weight the fixture input accelerometer by six dB and then run extremal control. Why do you have to do that if you are running extremal control? In other words, every time the input accelerometer has control you are six dB down. If it has control and it is weighted, it means your response is at least six dB down, or the input would not have control.

Dr. Curtis: I am familiar with response control testing. When we do it, we have a specification which has two spectra; one is a maximum input spectrum, the other is a maximum response spectrum. Now the response control test in MIL-STD-810D, in effect, has those same two spectra. It has the input spectrum, and then one could imagine another one six dB higher which is the maximum response spectrum. It just doesn't appear there in the standard; but in effect, it is there. Given that that is so, then assuming your extremal control software works properly, then I think, what you are doing is quite legal, and it meets the requirements of the standard.

Mr. Norris: Don't you think you would have a six dB undertest at those frequencies?

Dr. Curtis: The whole point of a response control test is to empirically put notches in the input at those points where, in the field, the impedance match between the support structure (which you are replacing by a shaker), and the test item would mandate a notch at that frequency. That is the whole philosophy of response control testing.

Mr. Norris: Except that if you do not weight that input accelerometer, then you still do not ever overtest because whenever the response tries to go over the input, it takes control. But, by weighting the drive of the input

accelerometer, it will not take control until the response is six dB lower than your specification.

Dr. Curtis: Well, I would envision kind of a double system if I wanted to do that. I guess I would power average my input to three or four accelerometers.

Mr. Norris: That is an average now, not extremal control.

Dr. Curtis: But, I have that power average input. Then I notch that by having each individual accelerometer also go back into the loop with a modified sensitivity to create notches in the right place. So, there is no reason why I can't notch the input as well as the response; that is, limit the input points as well as the response points. Who says what is an input point? Why is it different than any other point?

Mr. Norris: Well, you are given a profile to meet with a maximum response. That is response testing. Is that true?

Dr. Curtis: No. At least I see it a little differently than that. I say I am given an input to meet except at those frequencies where I must decrease the input to limit the response to some other requirement.

Mr. Norris: Very good. I just did not meet the specification. Because at those points where my response was down and my input had control, I was six dB down from my specification.

Dr. Curtis: But, if some point was only down because some other point meets the maximum response, I haven't violated the requirement.

Mr. Norris: No. That is not what I am saying.

Dr. Curtis: If it is notched at the input and no other point met the maximum response, then I do think you bend the rules a little.

Mr. Norris: That is exactly what happens when you weight that input accelerometer. It doesn't take control until the response is at least six dB below your specification. That is an example of control weighting that I was asking about.

Mr. Caruso (Westinghouse Corporation): We have run several of those tests in the lab. First of all, the curve that is shown in MIL-STD-810D, if we are talking about external stores, is a response control curve, it is not an input curve per se. It is a threshold. If the standard tells you to lower those levels six dB and you do that, you can't be violating it if you do what it told you to do. So, you are not undertesting. You are doing exactly what the standard told you to do. Second, you only need one control accelerometer; you don't have this conflict of accelerometers. The way we have traditionally done the test in the past has been

to use an accelerometer at the input to the store, at the mounting lugs essentially. Then we put response control accelerometers, not really response control but response measurement accelerometers, at the forward and aft ends of the store. Then, using the procedure outlined in the Standard, we would start the test at a six dB level below the threshold curve indicated. But first of all, we would do some sort of modal survey of the store to determine where our resonances were. Then at those frequencies we would add additional energy to that input curve to get the responses of the pod up to, or in excess of, the threshold that was given. But, there is only one control accelerometer. We are not overriding one accelerometer with another. The input is still the input. We have determined empirically what the input has to be to get the responses of the ends that are needed. We can monitor that throughout the test, but the input is only from one point. It is a very simple procedure to implement, so I am not quite sure what the confusion is. Also, again, if the procedure has told you to drop six dB to do it right, then by dropping six dB, you are certainly not undertesting. You are following the procedure.

Dr. Curtis: Henry, I think the gentleman's question was the question of trying to do a response control test with the notching actively on line by taking advantage of the extremal control option which is built into the digital software. This is a little different than what we have done and the method you just described where one apriori calculates where the notches should be and how deep they should be.

Mr. Caruso: We are not doing that per se. I think one of the subtle differences is the difference between our response control test and our response definition test. We are not controlling the responses per se; we are still controlling the input, but we are controlling it in such a fashion that responses come out the way we want.

Mr. Silver (Westinghouse Electric): We do use extremal control. If the off-axis exceeds a threshold, I think, again this is a tailoring concept, but it is also built into that particular specified concept. If the off-axis exceeds as a threshold, you hold it until it is no greater than the threshold, and that is the way the standard is written. That is what you do to meet the standard as it is defined. As I see it, the problem with that response is that we have a bad concept; it seems to me when we originally wrote this document it wasn't done the way it was intended. It was changed to put an input acceleration in which, to my mind, is incorrect because you don't consider the impedance of the pod. If it is a big pod, there are very large differences in apparent mass of that pod at the point of input. If you require the acceleration to be some flat amount, you fail to recognize that. In the process, you generate apparent response peaks that are not

the real resonant peaks of that pod. If you checked it as a free, free body on a force basis, those resonances wouldn't be there. So, that is what I think we are doing wrong. We should put in flat armature current or a flat force function into the drive rod to the store, and we should not require a 6 dB down acceleration at that point. I think that is a much more important aspect of what we are doing wrong.

Mr. Smith (Hughes Aircraft Company): They talk about the aliasing filters on shock tests. This is the first time I have thought about aliasing with shock tests. But, just going along with the thought - you usually analyze shock at a greater sampling rate, more than twice the maximum frequency. For example, in a particular application I happen to be sampling at eight times the maximum frequency. Do you then interpret the aliasing filter as being four times the maximum frequency? Would that be a legitimate interpretation of that requirement? It isn't made clear. It definitely would not be the maximum frequency. That wouldn't make any sense at all.

Voice: This question is on shock data analysis. The aliasing frequency is half of the sampling frequency. When I analyze shock data, I use a sampling rate that is much greater than twice the maximum analysis frequency. There are really no rules that say how many samples I can use; the more, the better. The bigger a buffer and the more time I have, the better an answer I can have. But, that would imply that every lab has a different requirement on what aliasing filter they use. If I sample at 100,000 samples a second, should I use a fifty thousand Hz aliasing filter?

Mr. Galef (TRW): I think the question is due to confusion between the maximum analysis frequency and the maximum frequency that is present in the data. It is the latter of course that is important in aliasing. If we have the resonant frequency of the accelerometer, which may be several hundred thousand Hz, then we would have to sample enough faster than that to avoid it, or else use an anti-aliasing filter to keep it out.

Dr. Curtis: I think I recognize sort of a general rule that says when one is analyzing transients, that you better use a bigger guard band. No one knows how big, but as you say, the bigger, the better. Probably a factor of five is the number I hear bandied about.

Mr. Andress (Spectral Dynamics): I translated a specification that I read from an accuracy requirement point of view. I forget what the number is, but I translated that to a sampling rate of ten times. That is the way we approached it. The aliasing filter falls right along with it. But it seemed to me that that specification set the sampling rates that we are going to have to use.

Mr. Parmenter (Co-Moderator): Thanks, Rudy. Rudy initially asked me to summarize the tone of the discussion. Actually, as you heard there were many tones. This is kind of random here. We started off with the person in the space program being locked in on specifications and having to work his way out. Now we ended up with some topics here that were extremely debatable. What is the best approach, even at this date? I think that would be about it really. There were several tones rather than a tone.