

IDA

INSTITUTE FOR DEFENSE ANALYSES

**Defense Science Study Group IV:
Study Reports
1994-1995**

Volume I

W. J. Hurley
N. P. Licato

19980324 041

DTIC QUALITY INSPECTED

February 1996

Approved for public release;
distribution unlimited.

IDA Paper P-3296

Log: H 97-000050

This work was conducted under contract DASW01 94 C 0054, Task A-103, for the Defense Advanced Research Projects Agency. The publication of this IDA document does not indicate endorsement by the Department of Energy, nor should the contents be construed as reflecting the official position of that Agency.

© 1997, 1998 Institute for Defense Analyses, 1801 N. Beauregard Street, Alexandria, Virginia 22311-1772 • (703) 845-2000.

Provided that copies will be given to authorized persons only, this material may be reproduced by or for the U.S. Government pursuant to the copyright license under the clause at DFARS 252.227-7013 (10/88).

INSTITUTE FOR DEFENSE ANALYSES

IDA Paper P-3296

**Defense Science Study Group IV:
Study Reports
1994-1995**

Volume I

W. J. Hurley
N. P. Licato

PREFACE

The Defense Science Study Group (DSSG) is a 2-year educational program, sponsored by the Defense Advanced Research Projects Agency (DARPA), which introduces outstanding young professors of science and engineering to the defense community and to current national security issues. The program has two basic components. The first focuses on group activities and provides a broad introduction to the defense community. This is achieved through a series of briefings by senior military and civilian officials and through visits to Joint Commands, industrial facilities, and military installations.

The second component of the program provides each member with the opportunity to personalize the DSSG experience by selecting a specific area of interest and spending about 2 weeks reviewing the activities of DoD in that area. This is done during the June and August sessions of the program's second year. The June session is held in Washington, DC, and the members interact with IDA and DARPA staff as well as with military and civilians from throughout DoD. The August session is held at Los Alamos National Laboratory where the members interact with Laboratory staff and prepare brief reviews of their subject areas. In November, at the final session of the program, the members brief the results of their studies to the other members and mentors of the program.

This volume contains the reports of the Fourth DSSG class which met during 1994-95. Again, the reports are brief, informal reviews of areas of specific interest to the members. Their primary purpose was to enable the authors to determine the interests of the defense community in the areas and to begin to make contact with appropriate individuals within that community.

One of the reports was done at the classified level and is available in a separate volume.

The reports were reviewed and organized by William J. Hurley and Nancy P. Licato who are, respectively, the DSSG Program Director and Administrator.

CONTENTS

Volume I

STUDIES AND ANALYSES

A. Solid State Alternatives for Image Intensification	1
<i>Kevin F. Brennan, School of Electrical and Computer Engineering, Georgia Institute of Technology</i>	
B. An Investment Strategy for the Air Force in Space Technology for the Next Twenty Years	25
<i>Daniel E. Hastings, Department of Aeronautics and Astronautics, Massachusetts Institute of Technology</i>	
C. Gathering Position Information Using Radio Tags and GPS	53
<i>Gaetano Borriello, Department of Computer Science and Engineering, University of Washington</i>	
D. Comparative Advantages of Titanium Alloy and High-Strength Steel Submarine Pressure Hulls	71
<i>William C. Johnson, Department of Materials Science and Engineering, University of Virginia</i>	
E. Land-Based Acoustic Sensor Arrays	105
<i>Christopher S. Kochanek, Department of Astronomy, Harvard University</i>	
F. Technical Issues Pertaining to Environmentally Sound Ships	123
<i>Ann R. Karagozian, Mechanical, Aerospace and Nuclear Engineering Department, UCLA</i>	
G. Nanocrystalline Materials for Absorbing Microwave and Infrared Radiation	147
<i>Brent T. Fultz, Department of Materials Science and Engineering, California Institute of Technology, and S. Lance Cooper, Department of Physics, University of Illinois</i>	

H. From Chips to Ships: Applying VLSI CAD Techniques to Naval Vessel Design	191
<i>Gabriel Robins, Department of Computer Science, University of Virginia</i>	
I. Laser-Assisted Friend or Foe Identification	219
<i>Clifford Pollock, School of Electrical Engineering, Cornell University</i>	
J. Military Simulations Using New Tools Involving Complex, Adaptive Systems	239
<i>Jean M. Carlson, Department of Physics, University of California, Santa Barbara</i>	
K. Exploiting and Controlling Cavitation	255
<i>Michael J. Shelley, Courant Institute of Mathematical Science, New York University</i>	
L. Mathematics and Theory in Virtual Engineering	267
<i>John C. Doyle, Electrical Engineering, California Institute of Technology</i>	
M. Military Applications of Diamond	289
<i>Richard B. Kaner, Department of Chemistry, University of California, Los Angeles</i>	

Appendixes

- A. Glossary
- B. Distribution List for IDA Paper P-3296, Volume I

Volume II -- Classified

Detection of Mines in Marine Environments	1
<i>A. Paul Alivisatos, Department of Chemistry, University of California, Berkeley</i>	
<i>Stephen D. Kevan, Physics Department, University of Oregon</i>	

**A. SOLID-STATE ALTERNATIVES FOR IMAGE
INTENSIFICATION**

**Kevin F. Brennan
School of Electrical and Computer Engineering
Georgia Institute of Technology
Atlanta, Georgia**

ABSTRACT

Image intensification is important in many military operations. It enhances the military's capabilities to perform in low light levels particularly in maneuvering, piloting, and reconnaissance operations. Currently image intensification is accomplished using vacuum-tube-based technology. Futuristic military programs, such as the 21st Century Land Warrior, demand full integration of direct view imagery including image intensification, GPS data, stored map data, infrared signal, and transmissions from others for all dismounted warriors. To accomplish this task, digitization of data and their subsequent fusion will be necessary. Though tube-based image intensifiers provide very high performance, as measured in terms of light level amplification, their inherent limitations render some uncertainty as to their usefulness in future, integrated, sensing applications. For these reasons, it is important to consider solid-state alternatives to image intensifier technology. This paper addresses the limitations of tube-based image intensifier and explores some of the possible solid state detector alternatives.

SOLID-STATE ALTERNATIVES FOR IMAGE INTENSIFICATION

1. INTRODUCTION

One of the key developments in modern military operations is the capability to operate under limited visual conditions such as darkness, fog, rain, smoke, etc. Two basic, different approaches have been developed for sensing in the dark. The first method, which is the principal subject of this paper, operates by amplifying the ambient optical radiation to produce an image that can be seen with the unaided eye. This technique is called image intensification and the devices that are used to perform this function are called image intensifiers. The second method relies on thermal imaging. In this technique the heat emitted by an object is collected and converted into an image (thermal contrast is converted into visual contrast on a display). Thermal imaging relies on the collection of far longer wavelength radiation than image intensification and can be used during both day and night.

Thermal imagers provide far better imaging through smoke, fog and rain than do image intensifiers. This is due to the fact that thermal radiation is emitted at much longer wavelengths, ~3-14 microns, longer typically than the dimensions of smoke, fog, or rain particles within the air, and is consequently not strongly scattered by the particles. Additionally, thermal imagers can detect objects at far greater distances than image intensifiers [1]. Subsequently, for long range, full day/night/foul weather applications a thermal-imaging system is superior to an image intensifier. However, image intensification is of primary importance in piloting, maneuvering, and reconnaissance where the operator needs to be able to distinguish between objects with the same temperature signature.

Image intensifiers amplify the ambient background light. Table 1 illustrates the range of light levels that the human eye responds to and representative sources and illuminance. As can be seen from Table 1, an image intensifier would need to operate over many orders of magnitude if it responded only to optical radiation originating either from the moon or starlight. From the chart, mean starlight provides between 10^{-3} and 10^{-4} lux illuminance. The response range of the image intensifier can be reduced if it takes advantage of near infrared radiation as well as optical radiation. Even though in darkness with no moon where little ambient optical radiation is present, there is still near visible,

near-infrared radiation. The sources of this radiation are zodiacal light and airglow. Zodiacal light originates from solar radiation scattered by residual interplanetary dust within the orbital plane. Airglow is caused by excitation of atoms and molecules in the upper atmosphere that are heated by the sun during the day and then release this energy at night. While the radiation produced by airglow and zodiacal light cannot be seen by the unaided eye, this radiation can be used by image intensifiers as a source of illumination. Typically, airglow and zodiacal radiation provide an illuminance between 10^{-2} and 10^{-3} lux, an order of magnitude higher than starlight alone [2]. As will be discussed below, existing GEN III image intensifiers have high response to these wavelengths, typically on the order of 0.4 to 0.9 microns.

Table 1. Low Light Sources

Illuminance (lux)	Source	Typical Seeing Conditions
10^6		Too Bright
10^5	Glare	Full Sun
10^4	Hazy Sun	
10^3	Partial Cloud	Optimum Performance
10^2	Room Lighting	
10^1	Candle Light	Reduced color and texture
1	Full Moon	Inaccurate color and texture
10^{-1}	Half Moon	Discrete objects discernible
10^{-2}	Thin Moon	Limit of color perception
10^{-3}	Thin Moon	Some outline perception
10^{-4}	No Moon; Starlight	Some contrast perception
10^{-5}	Overcast; No Moon	Limit of light perception

In general, any low light level detector must be able to distinguish the input signal of interest from noise. If the noise level of the detector and following electronics is extremely low, then low-light-level detection can proceed without front-end amplification. For example, at high input signal levels, front-end gain is unnecessary in a well designed detector since the converted input signal is sufficiently above the noise floor so that the signal can be readily discriminated from the noise. As the signal level diminishes in magnitude, it ultimately approaches the noise floor making signal discrimination very difficult. At low input signal levels then, improved detection can be achieved either by lowering the noise of the system or by increasing the signal strength. It

is difficult to lower the detector noise significantly without cooling well below room temperature, which is undesirable because of the concomitant cost, weight, and size constraints. Alternatively, by providing gain, the detector boosts the signal level well above the noise floor of the detector. Therefore, for image intensification the device must provide front end gain.

Image intensifiers were first introduced in the 1940s. All of these designs used tubes as do most of the current image intensifiers. The primary advantage of tubes, as we will see, is their ability to yield very high gain at very low additional noise, thus greatly improving the signal-to-noise ratio of the detector. The first tubes were called Generation 0 or GEN 0 and relied on infrared illuminators to bathe the surroundings. These devices have peak response in the blue-green region and are characterized by the presence of geometric distortion. Immediate successors to the GEN 0 tubes were the GEN I tubes, which were passive and had greater sensitivity. In these tubes, a photocathode is used along with electron acceleration through vacuum to achieve gain. These tubes were also subject to geometric distortion. A major advancement was achieved with the development of the GEN II tubes which incorporate a microchannel plate, MCP, to provide high amplification. The basic design of the GEN II tubes is still used today as well as the tubes themselves. In the next section, the basic operating principles of the GEN II and GEN III image intensifier tubes will be reviewed.

2. REVIEW OF EXISTING IMAGE INTENSIFIER TUBES

Existing image intensifier tubes consist of three main components, a photocathode, a microchannel plate, and an output light phosphor as shown in Figure 1. The tube is of course in high vacuum and requires high voltage for operation. Photons are incident onto the semitransparent photocathode where they are converted into electrons. The electrons are then emitted into the vacuum by the photocathode and are accelerated by an electric field into the MCP channels. The electrons accelerate to high energies through the MCP channels. During the course of the electrons' flights through the MCP channels, they suffer substantial collisions with the sides resulting in the production of numerous secondary electrons providing signal gain. Depending upon the applied voltage, efficiency and size of the MCP, one incident electron can lead to the ultimate generation of ~1000-3000 descendants [3]. Upon exiting the MCP, the electrons are accelerated again through vacuum under the influence of another electric field towards a phosphor screen. After striking the phosphor screen, the electrons are

converted back into photons resulting in a greatly amplified optical image. In most designs, the gap between each component is very small such that there is little lateral drift of the electrons. Consequently, an image focused onto the photocathode is correctly reproduced on the phosphor screen. This arrangement is called proximity focus. Finally, the output image must be inverted since the input objective lens produces an upside down image at the tube input. Image inversion is typically obtained using a fiber-optic inverter [3] known as a twister. Alternative designs utilize electrostatic inverting electron lenses between the photocathode and the MCP for image inversion [4]. Sketches of the these two GEN II designs are shown in Figure 2.

It is important to recognize that the GEN II design provides for very high signal gain through the use of the MCP. High gain is critical in these devices since there is a fair amount of inefficiency within the tube. There is substantial signal loss in the photocathode as well as the phosphor screen due to their respective inefficiencies. Subsequently, high gain is necessary to overcome these losses.

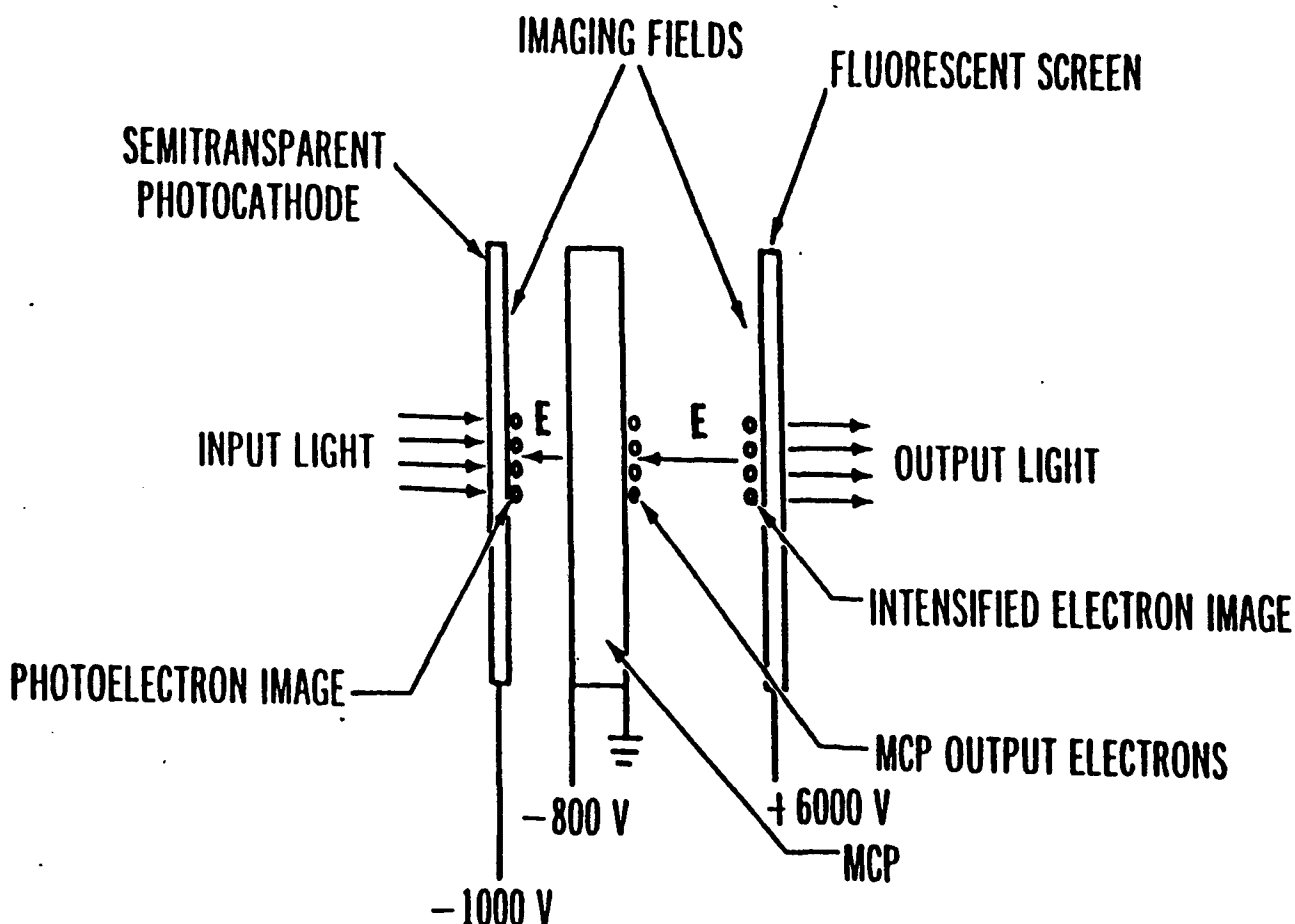
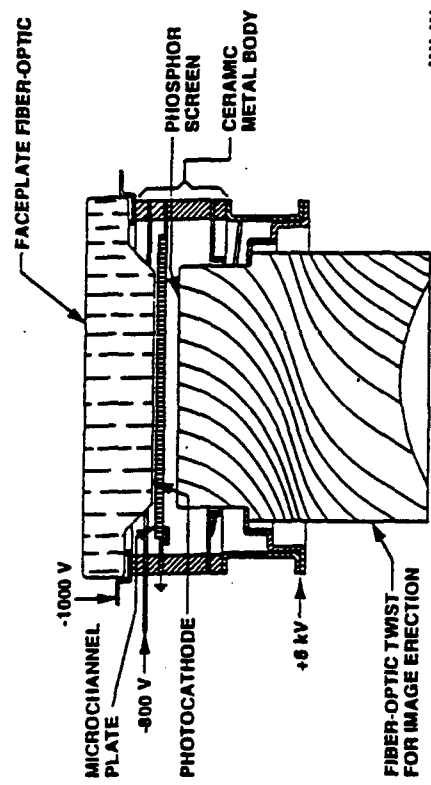


Figure 1. Basic Tube Anatomy for GEN II and GEN III Image Intensifier Tubes

Wafer Tube



ESI Tube

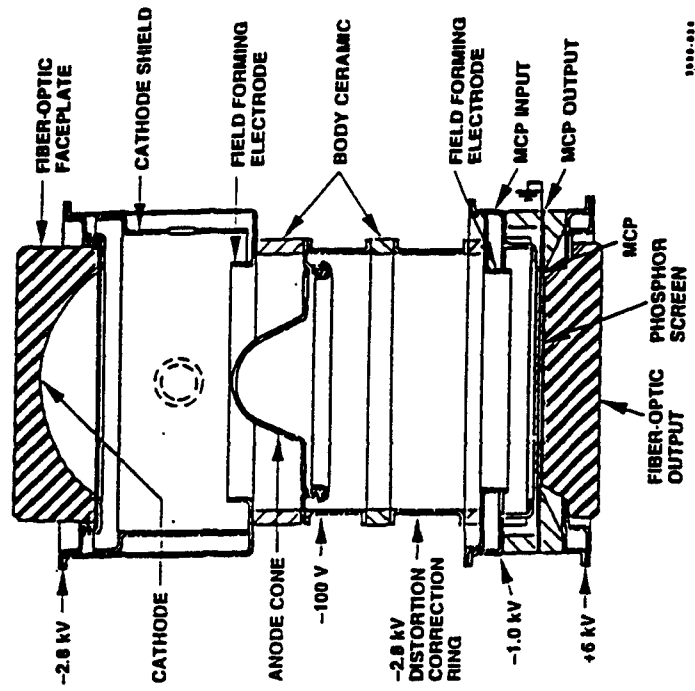


Figure 2. Sketches of Fiber Optic and Electrostatic Focusing GEN II Tubes

Current generation image intensifier tubes are known as GEN III devices. The GEN III tubes have basically the same design as the GEN II devices incorporating similar components. The primary advancement of the GEN III tubes is the use of a much more efficient GaAs photocathode, which results in much higher sensitivity. The GaAs photocathodes though are highly sensitive to contamination and thus require a very high vacuum environment. Nevertheless, the improved sensitivity gained through use of a GaAs photocathode warrants any complications that arise from the production and maintenance of high vacuum. Sensitivity is measured in output current vs. input light intensity. For the GEN III tube, the original specification was for 1000 microamps of current per lumen of light input. Through continuous process improvements, average sensitivity of current GEN III tubes is 1500 $\mu\text{A}/\text{lm}$. A sensitivity of 2000 $\mu\text{A}/\text{lm}$ is achievable, twice that of the original specification. Increased sensitivity broadens the useful spectral range of the intensifier. As shown in Figure 3, the advanced GEN III designs have extended the sensitivity of the original GEN III devices down to 400 nm from ~ 550 nm [4]. This more closely matches the spectral range of the GEN II devices but at much higher sensitivity at all of these wavelengths.

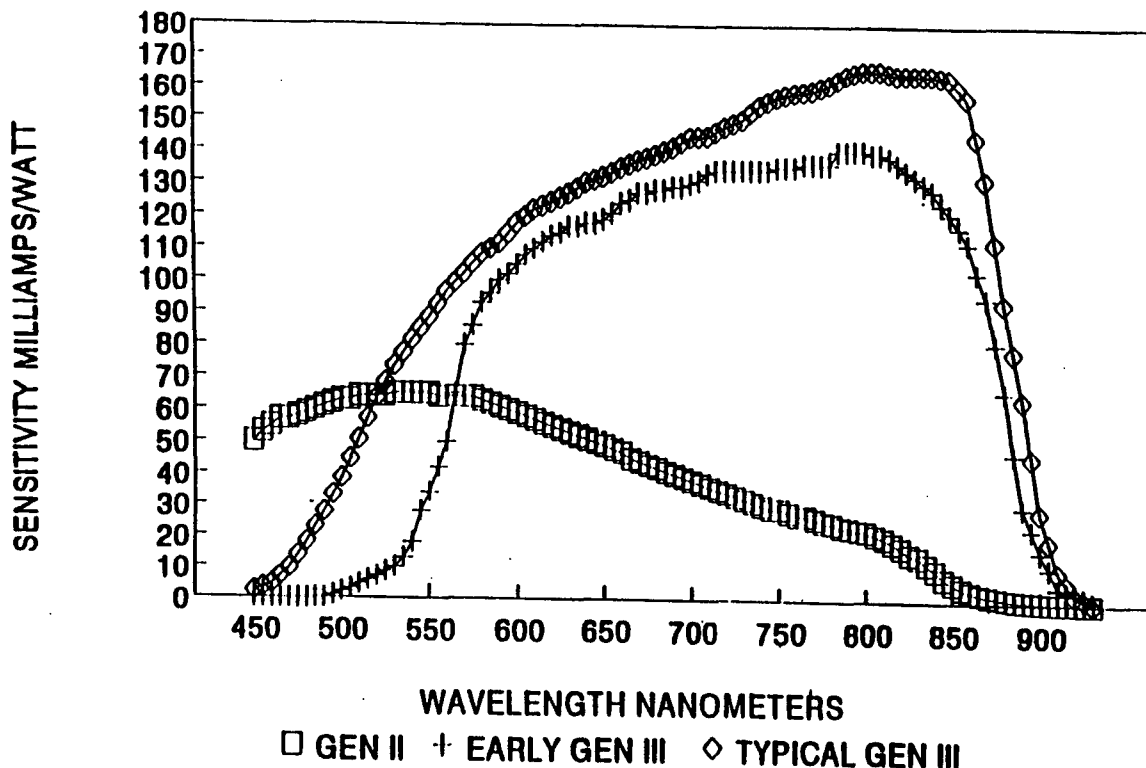


Figure 3. Photocathode Sensitivity; GEN II, Early GEN III, and Current GEN III

The GEN III devices also offer higher resolution than the GEN II tubes. The original specification for the GEN III tubes was 36 line pairs per mm (lp/mm) [4]. The major factors that govern tube resolution are the MCP channel separations, the spacings between the MCP and the photocathode and phosphor screen. Ultimately, reduction in the MCP channel spacing is limited by the Nyquist resolution limit. However, in practice, tube resolution is generally poorer than the Nyquist limit due to the resolution limits arising from the component spacings. The component spacings affect the lateral spreading of the electrons and hence the resolution of the imager. The closer the components are spaced the less lateral spreading. This results in improved proximity focusing. However, there is a limit to how closely the components can be spaced due to the appearance of local bright spots caused by emission points. Processing procedures can correct for these bright spots to some degree. With these procedures, GEN III tubes have been built to over 50 lp/mm resolution. The improved resolution enables both increased range at which fine detail can be seen as well as improved contrast.

As mentioned above, perhaps the most important issue in the operation of image intensifiers is the signal-to-noise ratio, SNR. Obviously, a high signal-to-noise ratio implies that the signal can be readily discriminated from the noise. At low SNR, it is typically difficult to resolve the signal of interest from the surrounding noise. At very low light levels the ability to see fine detail is less affected by the resolution of the tube than the SNR [4]. The range to which objects can be seen under very low light levels is directly proportional to the SNR [4]. The SNR can be improved in two different ways, through increased photocathode sensitivity and by reducing the MCP noise figure. The SNR is directly proportional to the square root of the photocathode sensitivity. This implies that a tube with twice the sensitivity of another tube will provide 40 percent higher SNR and hence 40 percent higher low-light-level resolution. The SNR is inversely proportional to the MCP noise figure. As discussed above, the MCP must provide high gain in order to amplify the input optical signal. However, if the amplification process is noisy, then the tube will deliver very bright noise rather than a clear, amplified signal. Consequently, the key function of the MCP is to provide high gain at very low additional noise.

Tube lifetime is nearly five times longer under identical operating conditions for the GEN III devices than the GEN II devices [4]. Extended tube lifetime has the obvious advantages of longer service lifetime of a fielded system as well as longer performance before degradation begins.

3. MILITARY REQUIREMENTS FOR IMAGE INTENSIFIERS

The primary military usage of image intensifiers is for piloting, maneuvering, driving, walking, observing, and targeting. All branches of the military are customers for this technology. However, this report will focus mostly on the Army's needs and requirements for image intensifiers since it is believed that the Army constitutes the largest users of image intensifiers.

The 21st Century Land Warrior program has identified cameras and more specifically image intensifiers as important ingredients in outfitting dismounted and mounted warriors. The Dismounted Battle Laboratory, located at Fort Benning, GA, [5], sees the principal applications of cameras for the infantry as follows. Within each platoon there will be roughly three soldiers equipped with cameras, with infrared, image intensification, and optical capabilities. If possible, it is desirable to merge all capabilities within one camera for each soldier. In a typical battlefield scenario, these soldiers will provide advanced reconnaissance and will be equipped to transmit images back to a platoon leader. Though unmanned air vehicles, UAV, will be used as well for reconnaissance, Fort Benning believes that reconnaissance teams will still be primarily used [5] in actual battlefield situations.

Being that the Army envisions dismounted infantry equipped with image intensification capability, weight, size, and power consumption become major requirements. Current helmet mounted units are in the neighborhood of 10 pounds. The objective is to reduce the entire system weight anywhere from 2-4.5 pounds depending upon the user [4]. For Air Force piloting operations, the proposed weight limit is at 4.5 pounds. When the combined weight of the helmet, oxygen mask, and eye protection visor are subtracted, only 600 grams remain for night vision capability [4]. State-of-the-art binocular intensifier tubes currently weigh 530 grams. Therefore, the weight of the tubes presents an important problem. Battery lifetimes for the system are currently around 12 hours so an extended battery lifetime or lower power consumption is also highly desired. Field of view is of greater importance in piloting than to the dismounted warrior, though in both applications improved field of view is of value. Current systems supply only a 40° field of view while pilots prefer at least a 60° field of view [6]. An entire helmet mounted system currently costs roughly \$10,000 [6] severely limiting the number of units the Army can deploy. Obviously, at this cost it is not possible to equip all dismounted infantry with helmet mounted night vision capability. Though the GEN III tubes do not represent the primary expense, ITT, one of the two major manufacturers

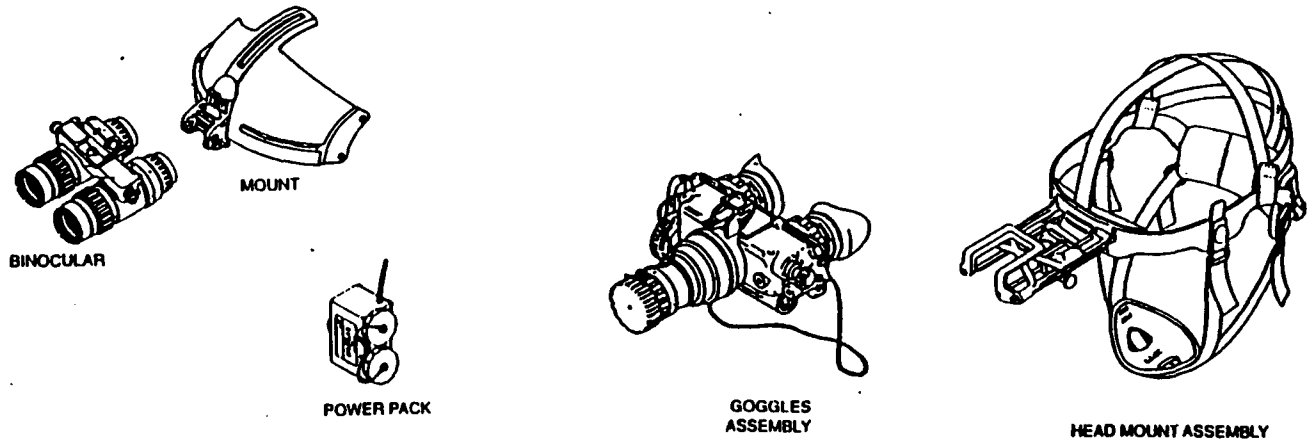
of the GEN III image tubes, expects that the tube cost cannot be lowered below a few hundred dollars [7] each. Therefore, the cost of two tubes alone is anywhere from \$500-\$1000. Consequently, even if the power supply, mounting unit, etc., were reduced substantially in cost, it is not currently believed that the cost per unit could possibly drop much below \$1000, owing in part to the cost of the image intensifier tubes themselves.

Shared situational awareness is one of the key aspects of the 21st Century Land Warrior program. To accomplish this mission, it is necessary to integrate information from multiple sources, i.e., direct view including image intensification, GPS data, stored map data, infrared signal, and transmissions from others. In order to integrate these data together it is important that it be encoded digitally. Digital storage of imagery enables fusion of multiple wavelength imagery as well as incorporation of map and GPS data onto one image presented to the soldier on a head-mounted visor display. Current image intensifier tubes do not provide for digital output. In fact, the output of these devices is directly viewed, providing a very poor interface for image fusion. A very different approach will need to be adopted in order to provide fusion of the imagery derived from a image intensification tube.

Currently, there is a significant effort to improve the GEN III image intensifier tubes. There is a program called the Advanced Image Intensifier Advanced Technology Demonstrator (AI²ATD) coordinated through the Army. The goal of this program is to increase the field of view to 60° and increase the resolution by 25 percent from currently fielded units. New tubes, Omnibus III, have been developed and deployed in systems, particularly the ANVIS and PVS-7. These systems are sketched in Figure 4. The refined Omnibus III tube offers 25 percent better range and operates in light levels 40 percent lower than the GEN III tubes. These refinements have been achieved through improved processing, better photocathode materials, and phosphor screens. Though the newest tubes have made significant strides, further improvement is necessary. Specifically, the weight, size, cost, and digital data output requirements present serious challenges to tube based designs.

ANVIS and PVS-7

- Omnibus III tubes give 25% better range
- Contrast over 100% better on larger objects
- Effective in light levels 40% lower



ANVIS

PVS-7B

Figure 4. Current State-of-the-Art Helmet Mounted Intensifier Systems

4. ALTERNATIVES TO IMAGE INTENSIFIER TUBES

As was discussed in the previous section, image intensifier tubes have several limitations that may reduce their usefulness in future military image intensifier systems. Perhaps the most important limitation of tube based intensifiers is the fact that they do not readily provide for electronic output which can be easily digitized thereby enabling fusion of data. The most obvious means of circumventing this problem is to develop a solid-state image intensifier in which electronic output is naturally obtained. A solid-state detector would provide direct electronic output that could be readily digitized for use in future integrated systems. Other presumed advantages of solid-state detectors are lower cost, lower power requirements, lower weight and size. In this section, two different solid-state based approaches will be examined. The first is based on charge coupled devices, CCDs, while the second uses CMOS based active pixel sensors.

a. CCD-Based Devices

One of the most important CCD designs uses silicon metal-oxide-semiconductor, MOS, structures, known as surface channel CCDs. The principal advantages of silicon-based CCDs are their very low cost and high technological maturity. A CCD is composed of a two-dimensional matrix of MOS structures each called a pixel. Charge is collected underneath each pixel and is clocked out sequentially across the array by applying a sequence of voltages to each pixel. Nearly perfect charge transfer between pixels is required for accurate image reconstruction, which in turn limits the readout speed of the device.

The CCD can be directly applied to low-light-level imaging, but this typically requires cryogenic cooling, slow scan readout rates, and extended integration times [8]. The main limitation in using CCDs directly in low light levels is that in their simplest implementation, they do not provide any front-end gain. Subsequently, in order to collect sufficient carriers to overcome readout noise, etc., the integration time of the CCD must be extended resulting in slower readout rates. In fact, for very low light levels, insufficient carriers are collected to overcome the noise floor in a typical room-temperature operated CCD. For these reasons, the direct usefulness of CCDs is restricted in low-light-level applications. Alternatively, CCDs are often combined with image intensifier tubes to provide electronic output.

Two of the most promising means of using CCDs for image intensification are either optically coupling the CCD to the phosphor screen of an existing tube using an objective or a fiber optic plate or inserting the CCD in place of the phosphor screen and directly bombarding it with electrons [8-10]. Usage of CCDs in these modes does not circumvent the limitations of the tubes raised above in terms of cost, size, and weight since these devices are again vacuum devices which incorporate all or some of the tube components. Nevertheless, these CCD designs provide an electronic output which can be digitized to provide data integration.

The use of CCD designs in the electron-bombarded mode eliminates the microchannel plate, phosphor screen, and output optics of the tube designs. Electrons emitted from the photocathode and accelerated by an electric field applied in vacuum impinge directly onto the CCD. Gain is achieved from the multiplication process which occurs when very high energy electrons (several keV) are stopped within a semiconductor material. On average, an electron-hole pair is produced for every ~ 3.6 eV of energy in silicon [9,10]. Consequently, a gain of several thousand is achievable. The relatively low

noise level of the CCD readout at room temperature, (typically 10-100 electrons/pixel frame [8]) coupled with this high gain ensures that the CCD can detect individual photoelectrons. The CCD must be backside bombarded to avoid damaging the front side connections as well as to optimize carrier production. Backside bombardment requires in turn wafer thinning down to about 10 μm [9]. Electron bombarded CCDs require a radiation-hardened device. Damage to gate structures from keV x-rays and ballistic electrons can result in increased dark current as well as a reduction in pixel well capacity.

Optical coupling into a CCD is typically accomplished using a fiber optic coupler. This can be accomplished without harming the array using a nonpermanent oil coupling or direct bonding [10]. The particular advantage of this approach is that it can be operated with existing image intensifiers. In addition, the optically coupled CCDs have good resolution performance.

CCDs can also be directly used if they are modified to provide front-end gain. By inserting an avalanche photodiode, APD, array on top, a CCD can be effective in low-light-level imaging. An APD is a photodiode sufficiently reverse biased such that carrier multiplication via impact ionization occurs. Incident light is absorbed within the photodiode and the subsequently photogenerated carriers are multiplied through impact ionization events. The carrier multiplication provides gain. The multiplied carriers are collected within the depletion layer of the APD and are transferred into the CCD wells for readout after a suitable integration time. By combining the amplification of the APD array with the CCD readout, low light level imaging is possible.

Though CCDs are relatively mature, are used in many camera products, and can be adapted for low light level applications, they have some important limitations. The most serious limitation of CCDs is their need for nearly perfect charge transfer efficiency [11]. Table 2 shows the importance of charge transfer efficiency. As can be seen from Table 2, only ~13 percent of the collected electrons are available at the output of the array if the charge transfer efficiency, CTE, is 0.999 for a 1024 x 1024 array [11]. This means that only 1 electron in 1000 can be lost. Since a typical charge packet in a CCD contains about 1000 electrons, then only 1 electron can be lost per transfer. At a CTE of 0.999999, where a far better fraction of electrons are available at the output, 98 percent, this translates into the loss of only one electron after 100 transfers! Because a single broken bond in the semiconductor crystal can capture a signal electron, then nearly perfect semiconductor crystalline material is required to maintain the required CTE. This is the primary problem with CCDs, and the principal reason why extension of CCDs to other

material systems is technologically difficult. The requirement of nearly perfect charge transfer efficiency makes CCDs (1) radiation "soft," (2) limited in their readout rate, (3) difficult to reproducibly manufacture in large array sizes, and (4) difficult to extend to alternative materials to enhance the spectral responsivity [11]. Additionally, CCDs are difficult to integrate with CMOS circuitry and require complex clocking circuitry with a large power dissipation. In an attempt to overcome the limitations of CCDs, an alternative technology, active pixel sensors, APS, has been developed [11-12].

Table 2. CCD Performance vs. CTE [11]

Array Size	Charge Transfer Efficiency	Fraction of Electrons at Output
1024 x 1024	0.999	0.128
	0.9999	0.815
	0.99999	0.980
2048 x 2048	0.99999	0.960

b. Active Pixel Sensors

Active pixel sensors, APS, preserve many of the advantages of CCDs, i.e., high sensitivity, etc., but overcome many of its disadvantages. Table 3 summarizes the advantages of the APS technology over CCDs. The major advantage of the APS technology is that charge transfer across an entire row or column of the array, which is basic to CCDs, is avoided. Instead charge transfer occurs only between the collection gate and the sensing gate. Subsequently, though high quality material is of course optimal, the performance of the sensor is less sensitive to material quality overall than a comparable CCD. A further important advantage of APS electronics over CCDs is the much lower power requirement of APS devices. The basic APS structure is based on CMOS circuitry and as such operates at very low power and voltage. In comparison, CCDs, use much higher power.

Table 3. Comparison of CCD and APS Technology

CCD	APS
Difficult to integrate with CMOS circuitry	Easy incorporation of on-chip signal processing, either analog or digital
Complex clocking circuitry requires	Only one row active at a time; transfer between pixels does not occur
Large power dissipation	Lower power dissipation
Difficult to scale to large array sizes	No pixel charge transfer; large arrays far less complicated to build
Near-perfect material requirement complicates application to other materials	Can be extended readily to other material systems since crystalline purity is greatly relaxed
Readout rate limited by sequential readout	Readout by pixel only; random access possible

Operation of an APS structure can be understood using Figure 5 [11]. Light is incident onto the photogate, marked as PG in the figure. Photogenerated electrons are collected and integrated underneath the photogate. After signal integration, the pixel can be read out by the following process. Row selection transistor S is biased high (+5 V). This activates the source follower amplifier. The reset gate R is briefly pulsed to +5 V to reset FD, the floating diffusion output node. The output of the source follower is then sampled onto the holding capacitor, CR, using transistor SHR. Next, the integrated charge under PG is transferred in FD using the transfer gate, TX, and by pulsing PG low. The new source follower output voltage is sampled by the holding capacitor, CS, using transistor SHS. The advantage of storing the signal charge on SHS and the reset charge (before signal transfer) on SHR, enables correlated double sampling of the pixel. In other words, the baseline charge stored in the diffusion node can be subtracted from the collected signal charge, eliminating the thermal noise from the pixel.

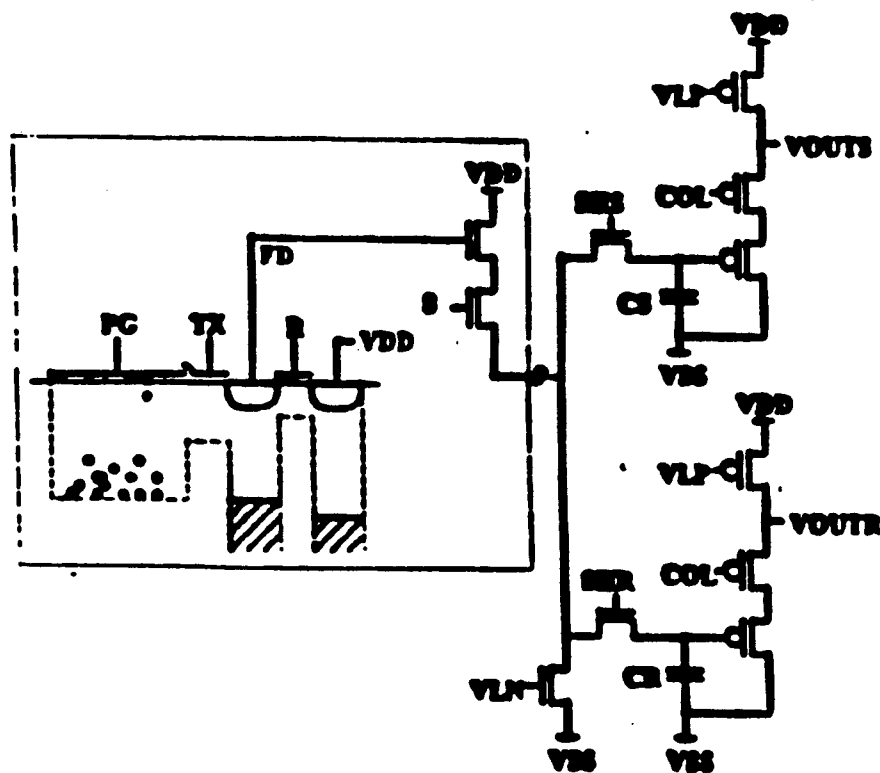


Figure 5. CMOS APS Unit Cell, Enclosed by the Dashed Line, and Related Circuit

128 x 128 two-dimensional CMOS APS optical imaging arrays have been developed [13]. These chips have been incorporated into simple cameras and excellent performance has been obtained [14]. Aside from the technical advantages APS devices offer, it should be noted that silicon based APS devices use standard CMOS processing, thus leveraging the CMOS industry. This has tremendous cost advantages since an APS imaging chip can be made from CMOS foundry services eliminating the need to construct a separate manufacturing line for chip development as is the case for CCDs.

5. NOVEL INTEGRATED APD/APS IMAGE INTENSIFIER

From the above discussion, it is clear that the development of a completely solid-state alternative to existing tube-based image intensifiers offers many advantages. Though CCDs are a relatively mature technology, their direct use in low-light-level imaging is questionable. In this section a novel approach to low-light-level imaging using APS technology is proposed. It is possible that by combining APDs with APS that

a new completely solid state image intensifier can be developed that will overcome many of the limitations of tube and CCD based devices.

Figure 6 shows a simplified sketch of the proposed APD/APS pixel. The proposed device will be made from the GaAs materials system since the absorption length in GaAs in the near infrared is relatively small, $\sim 0.5 \mu\text{m}$. In silicon, the absorption length at these wavelengths is very much longer, $\sim 20 \mu\text{m}$, which would be far less attractive due to recombination losses. Initially, the APD is biased well above the point where avalanche breakdown occurs. The applied voltage is then removed. The diode stays in the nonequilibrium, depleted mode, recovering slowly only through dark current generation/recombination events. If the dark current generation/recombination current is small, which is easily achievable in well made GaAs diodes, then the diode returns to equilibrium very slowly. During this recovery time, signal charge can be collected from incident light. Incident light is absorbed in the top, p+ region of the APD, producing electron-hole pairs. Since the diode is reverse biased, the photogenerated holes are collected in the p+ region while the electrons diffuse to the depletion region. Once within the depletion region, the electrons are accelerated by the electric field to high energy where they can impact ionize. Gain is provided by the impact ionization events. The initial carriers and all of their subsequent electronic descendants are collected at the edge of the depletion region and stored by the junction capacitance [15,16]. At the end of the integration time, a bias is reapplied to the diode and the stored charge is read out into the APS device. The operation of the APS device is similar to that described in the previous section.

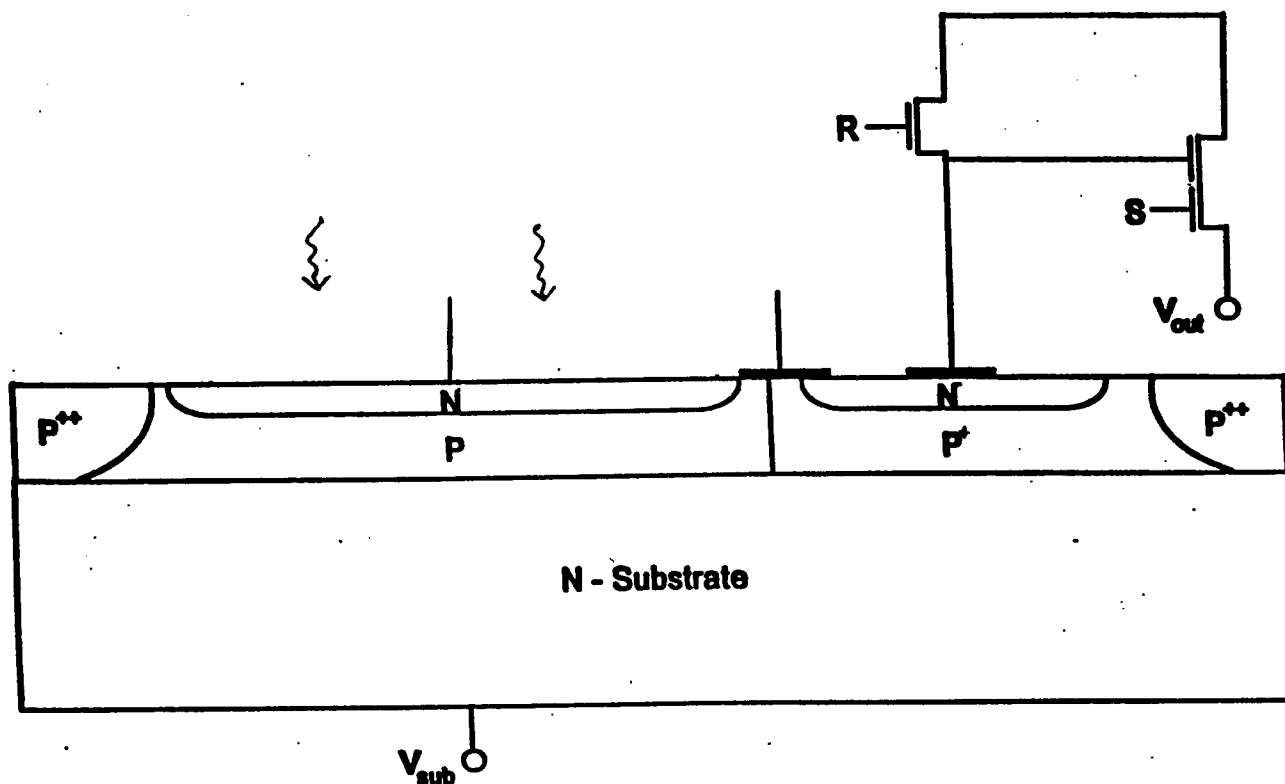


Figure 6. Simplified sketch of the proposed APD/APS image intensifier

The APD must be designed such that relatively high gain can be achieved at little excess noise. Very low noise operation can best be achieved in a photodiode if the carrier ionization rates are very dissimilar [17]. Unfortunately, avalanche photodiodes made from bulk GaAs or AlGaAs exhibit very poor noise performance owing to the nearly equal electron and hole ionization rates in these materials [18]. However, new APD designs [19] have been developed using the GaAs/AlGaAs materials system which exhibit very low noise. In these APD structures, gain is achieved using a series of built-in mini-junctions within the depletion region. Since the ionization occurs principally from the action of the electric field within the mini-junctions, the voltage requirement of the diode is 5-10 times lower than in a standard APD structure. This results in a significantly lower power device. Integrating these devices with APS structures can potentially lead to very low light level, completely solid-state detectors. Since both the superlattice APD and APS electronics are relatively low voltage, low power devices, it is expected that the APD/APS array will operate at very much lower voltages and power than competing CCD arrays, and orders of magnitude lower than tube based image intensifiers. In summary, the integrated APD/APS devices can be readily configured to produce a

digitized output, are separately addressable, operate at low voltages, have very low readout noise, can be extended to many different materials systems providing for multi-spectral capability, are compact and light weight. Therefore, it is conceivable, that the proposed APD/APS array can operate as a very effective image intensifier.

6. SUMMARY AND CONCLUSIONS

Futuristic military programs, such as the 21st Century Land Warrior, demand full integration of direct-view imagery including image intensification, GPS data, stored map data, infrared imagery, and transmissions from others for all dismounted warriors. To accomplish this task, digitization of data and their subsequent fusion will be necessary. Though tube-based image intensifiers provide very high performance, as measured in terms of light-level amplification, their inherent limitations make their use questionable in future digitized image intensifier systems. For these reasons, it is important to consider solid-state alternatives to image intensifier technology. In this paper, the limitations of tube-based image intensifiers, future military requirements of image intensifiers, and present and potential solid-state alternatives were addressed.

Though CCD technology is presently very mature, the inherent weakness of CCDs, their requirement for nearly perfect charge transfer efficiency and its concomitant requirement of nearly perfect crystalline quality, make CCDs less attractive than active pixel sensors, APS, for many imaging applications. APS detectors eliminate across the chip charge transfer, greatly relaxing noise and processing constraints. In addition, APS devices are separately addressable and can readily provide a digitized output. Integration of GaAs based APS electronics with GaAs/AlGaAs superlattice APD structures can potentially yield a light weight, compact, low voltage and power, completely solid state image intensifier.

REFERENCES

1. E.N. Philips, *Electro-Optics—Image Intensifiers*, ITT Internal Report, pp. 1-19, Doc: JSI/100, ITT Defense, Roanoke, VA.
2. E.N. Philips, *Range Performance Analysis of Image Intensifier Systems*, ITT Internal Report, ITT Defense, Roanoke, VA.
3. C.F. Freeman, "Image Intensifier Tubes," *Applications of Electronic Imaging Systems*, *SPIE*, Vol. 143, pp. 3-13, 1978.
4. Roy H. Holmes, *Night Vision 94 Image Intensifiers*, ITT Internal Report, ITT Defense, Roanoke, VA, 1994.
5. C.H. Thornton, Major P.S. Hilton, and F. Mazzocchi, private communication, Dismounted Battle Laboratory, Fort Benning, GA, June 20, 1995.
6. Ray Stefanik, private communication, Night Vision Laboratory, Fort Belvoir, VA, June 6, 1995.
7. Chip Hambro, private communication, ITT Defense, Roanoke, VA, January 12, 1995.
8. W. Enloe, R. Sheldon, L. Reed, and A. Amith, "An Electron Bombarded CCD Image Intensifier with a GaAs Photocathode," *SPIE*, Vol. 1655, *Electron Tubes and Image Intensifiers*, pp. 41-49, 1992.
9. J.C. Richard and R. Lemonier, "Intensified CCD's for Low Light Level Imaging," *International Electronic Imaging Exposition and Conference*, pp. 13-16, 1986.
10. T.F. Lynch, "Development of Intensified Charge-Coupled Devices (CCDs) and Solid State Arrays," *SPIE*, Vol. 143, *Applications of Electronic Imaging Systems*, pp. 36-41, 1978.
11. E.R. Fossum, "Active Pixel Sensors: Are CCDs Dinosaurs?" *SPIE*, Vol. 1900, *Charge Coupled Devices and Solid State Optical Sensors III*, pp. 2-14, February 1993.
12. S. , S.E. Kemeny, and E.R. Fossum, "CMOS Active Pixel Image Sensor," *IEEE Trans. Electron Dev.*, Vol. 41, pp. 452-453, March 1994.
13. E.R. Fossum, private communication, Jet Propulsion Laboratory, Pasadena, CA, January 1995.
14. S.K. Mendis, S.E. Kemeny, and E.R. Fossum, "A 128 x 128 CMOS Active Pixel Sensor for Highly Integrated Imaging Systems," *IEDM Tech. Digest*, pp. 583-586, 1993.

15. H. Komobuchi, M. Morimoto, and T. Ando, "Operation and Properties of a p-n Avalanche Photodiode in a Charge Integrating Mode," *IEEE Electron Dev. Lett.*, Vol. 10, pp. 189-191, May 1989.
16. H. Komobuchi and T. Ando, "A Novel High-Gain Image Sensor Cell Based on Si p-n APD in Charge Storage Mode Operation," *IEEE Trans. Electron Dev.*, Vol. 37, pp. 1861-1868, August 1990.
17. R.J. McIntyre, "Multiplication Noise in Uniform Avalanche Junctions," *IEEE Trans. Electron Dev.*, Vol. ED-13, p. 164, 1966.
18. P.P. Webb, R.J. McIntyre, and J. Conradi, "Properties of Avalanche Photodiodes," *RCA Review*, Vol. 35, pp. 234-277, June 1974.
19. K. Brennan, "Theory of the Doped Quantum Well Superlattice APD: A New Solid State Photomultiplier," *IEEE J. Quantum Electron.*, pp. 1999-2016, October 1986.

**B. AN INVESTMENT STRATEGY FOR THE AIR FORCE IN
SPACE TECHNOLOGY FOR THE NEXT 20 YEARS**

**Daniel E. Hastings
Department of Aeronautics and Astronautics
Massachusetts Institute of Technology
Cambridge, MA**

AN INVESTMENT STRATEGY FOR THE AIR FORCE IN SPACE TECHNOLOGY FOR THE NEXT 20 YEARS

VISION

The recommendations for technology investments derive from a vision of the Air Force in space in the 21st century, in which the Air Force has achieved survivable, on demand, real time, global presence that is affordable. This vision represents a revolutionary increase in capabilities for the Air Force and is achievable with targeted Air Force technology investments and adaptation of commercial developments. These technology investments will enable the United States to maintain military superiority by the exploitation of space through four themes:

- Global Awareness
- Knowledge on Demand
- Space Control
- Force Application.

THE CURRENT SITUATION IN SPACE

Space will continue to be the proverbial high ground for the foreseeable future. Desert Storm showed that space assets integrated with air, ground, and sea assets can play a critical role as force enhancers in fighting and winning conflict. Senior Air Force leadership¹ asserts that "space systems signal America's stature as a world power and aerospace nation. Control of space and access to it are fundamental to economic and military security. Ask the 20 foreign countries who will have space capabilities by the year 2000: a presence in space implies influence, power and security."

The space enterprise can be divided roughly into four areas:

- Launch systems
- Spacecraft bus systems

¹ Hon. Sheila Widnall, Secretary of the Air Force, September 1993.

- Spacecraft payload systems
- Spacecraft operations.

In the area of launch systems, the United States has an old, unresponsive, and relatively expensive set of launchers. Foreign launch systems have taken a substantial fraction of the world market and the number of countries able and willing to launch payloads is continuing to increase. In the area of spacecraft bus technology, the United States is in a leading position; the one major area where other countries have taken the lead is in spacecraft propulsion where U.S. technology is behind what has been accomplished in the former Soviet Union. The United States is still leading in the integration and operation of spacecraft payload systems both from the component level into spacecraft and in the development of constellations of spacecraft. This leading position is due to previous Government investments in research and development in space technology; maintaining this lead in the future depends on the technology investments that the U.S. Government and private sector make today.

The reduction of resources available to the DoD in the post-Cold War era means that DoD investment in space technologies and space systems must be firmly rooted in the goal of affordable systems. To this end, the DoD must plan its technology investment with a clear view of technological advances in the commercial world. It is undesirable and unnecessary for the DoD to develop every technology for its space systems on its own. The commercial sector will develop many technologies that the military can adapt for its own use with minimal investment. On the other hand, there will always be unique requirements for military systems that necessitate the use of technologies that have no commercial application, that push the performance limits of dual-use technologies, or whose timescale and risk are not attractive to the commercial sector. The DoD should carefully target its investments in technology to achieve the highest possible return. Technologies that are candidates for DoD investments fall into one of three possible categories:

- Revolutionary technologies in which the DoD must invest vigorously, because they are critical to the military mission and have little or no application in the commercial sector—Without DoD investment, these technologies will not advance. These technologies will enable a substantial increase in the exploitation of space by the DoD.
- Evolutionary technologies in which the DoD should invest, because they are similarly critical to the military mission and have little or no commercial application—These technologies will enable gradual improvements that over

time can significantly improve the performance or reduce the life-cycle costs of military systems.

- Technologies in which little DoD investment is required, because they will be led by the commercial sector—In these areas, the DoD should carefully monitor the progress that industry is making and invest only to the level necessary to adapt commercial technologies to the military mission.

The DoD should not underestimate the benefits of a healthy synergism between military and commercial research and development.

The Air Force's investment in space technology has fallen in recent years both as a fraction of TOA and as a fraction of spending on R&D. The inception of SDIO in 1983 subsumed the primary Air Force space technology programs under the SDIO umbrella. During the heyday of SDIO, total DoD investment in space technology programs was in excess of \$500 million per year. The primary emphasis during this time was highly survivable space technology development and demonstration. In addition, major programs were initiated under the SDIO umbrella in active and passive space sensors, radiation hardened electronics, advanced high data rate communications, high efficiency and high power solar arrays, and high density power storage technologies along with cryocoolers and other related structural technologies. Great emphasis was placed on technologies and systems that were highly survivable to a variety of threats including laser, nuclear, and microwave effects. The evolution of SDIO into BMDO and the current program direction to address theater missile defense had a negative impact on the space technology development budget and, as a result, space technology investment has decreased from \$500 million per year down to \$200 million per year. This dramatic decrease in space technology development investments will have a serious impact on the dominant position of the United States in space systems development, both commercial and military, for years to come. Current Air Force investment in space technology is grossly inadequate. The exploitation of space for military advantage requires aggressive and continuous investment.

SPACE TECHNOLOGY AND SPACE SYSTEMS IN THE FUTURE

In the next 10-20 years, the commercial world will see the development of four types of space-based systems that will be available to both friendly and unfriendly nations, corporations, and individuals on a worldwide basis. These systems will provide commercial services but will also be militarily useful. In addition, these systems will either involve other countries that build or purchase them or will involve international

consortia of investors. These systems will lead to the growth of new service industries based around their use that will be economically powerful. (A current example is the GPS; the growth of civilian users, including the FAA, is now creating a dilemma about the compatibility of unfettered access to precision GPS with the DoD's need to maintain a competitive military advantage.)

Four types of commercial services that will be available include the following:

- Global positioning and navigation services—While the DoD already has GPS, other countries are developing equivalent systems or augmenting the existing one; similar capabilities will be available through the development of personal communication systems. They will enable navigation with an accuracy at least several tens of meters.
- Global communication services—Several systems have already been proposed such as Iridium, Globalstar and Inmarsat-P. These systems will provide universal communications services between mobile individuals to almost anyplace on the surface of the earth. These systems will work transparently with local cellular systems and will enable rapid telecommunications development in underdeveloped parts of the world.
- Information transfer services—These services will enable data transfer between any two points on the surface of the earth at rates ranging from a few kb/s to Gb/s. Proposed systems include Orbcomm, Spaceway, Cyberstar and Teledesic. Individual users will be able to access large amounts of data on demand. Direct TV from direct broadcast satellites is a harbinger of what will be possible.
- Global reconnaissance services—These services will provide commercial users multispectral data on almost any point on the surface of the earth on a meter scale resolution. This data will span the range from the radio frequencies to the IR to the visible. This information will be available within hours of a viewing opportunity and on the order of a day from the time of a request. Proposed systems include improvements to the French SPOT as well as Orbimage, World View and various types of radarsats.

Each of these services will be part of the global infosphere. It will be possible for a person of means to locate themselves on any point on the earth, communicate both by voice and computer to other points on the earth and have a good picture of the local environment. Both the services and the technologies that enable them will be commercially available all over the world. Given the enormous magnitude of the commercial market, military and NASA communications will have to be fully integrated

with, and technologically dependent on, the exploding market-driven communications technologies.

Nevertheless, there will be military-specific needs that are not encompassed by these four types of commercial services:

- Geographically selected denial of high-precision global positioning (sufficient for weapons delivery) to an opponent, and assured friendly access to those same services
- Assured access to communications that are robust against jamming and tampering and/or covert, including local surge capacity to deployed forces
- Assured relay of very high data rate intelligence information from geosynchronous distances
- Day/night all weather reconnaissance of low contrast stationary and moving targets with hyperspectral resolution and in the least possible time.

The proliferation of space applications at affordable prices will tend to offset the current U.S. military edge. Capabilities in these and other areas will enable the U.S. military to have the advantage over an opponent who is also exploiting the infosphere.

SPACE TECHNOLOGY DEVELOPMENTS IN THE COMMERCIAL WORLD IN THE NEXT 20 YEARS

To deliver the services described above in a competitive environment, the commercial world will invest in bringing many technologies relevant to space to commercial viability. The technologies that the commercial world will develop include the following:

- Technologies for manufacturing many identical spacecraft
- Technologies for efficient spacecraft operations
- Low-cost high-performance electronics and computers
- Technologies for commercial global communications
- Small expendable space launch systems
- Systems-level simulation-based design
- Technologies for automated spacecraft checkout.

These technologies will result in standardized, modular bus designs that can be launched on any compatible launch vehicle, simplified payload designs, commoditized payload elements and efficient (e.g., autonomous) operations. In addition, the

commercial world will develop management techniques to reduce system cost and delivery time as well as refining techniques for cost estimating and scheduling. Relying on the commercial world to develop these technologies, the Air Force will need to invest only where it is necessary to adapt these technologies to meet specific military requirements. However, not all of the functions needed by the Air Force will be achievable solely with commercial developments.

A fundamental question for the Air Force in space in the future will be to determine which technologies will allow the United States to retain advantage in space.

IMPLICATIONS OF THE VISION FOR THE AIR FORCE IN SPACE

The vision for the Air Force in space requires increased capability over projected commercial systems, yet these increases must come in a time of decreasing budgets. Therefore, the cost of space systems must be reduced to make these capabilities achievable. The costs of space systems are dominated by the costs of the individual elements, the costs to launch the space elements, and costs to operate them. Historically, the cost of space hardware has scaled directly with mass. In order to break the current cost paradigm in each of these areas it is necessary to invest in or to invest to adapt several key technologies from the commercial world. The relevant technologies are those that reduce the satellite mass for the same or increased functionality, technologies for launch vehicle cost reductions and performance improvement and technologies for spacecraft automation.

FOUR ACHIEVABLE THEMES THAT WILL CONSTITUTE THE VISION FOR SPACE

With the attribute of affordable systems, the space technology investments can be grouped under the four themes of Global Awareness, Knowledge on Demand, Space Control, and Force Application. These themes will be enabled by the targeted investments by the DoD as well as the related investments in the commercial sector.

Global Awareness

Global Awareness is the idea that space technology will enable the ability to see in near real time everywhere on the surface of the earth or in the air or near space, under all weather conditions at any time. The integration of this ability with the command and control system for a military operation will enable the U.S. military to respond and

outthink any potential adversary in a context where space based information will be available on a worldwide basis. The timely acquisition and use of information inside a decision loop of any enemy will confer a tremendous advantage on U.S. forces. Global Awareness also has enormous deterrent value. Any adversary will know that he is under continuous surveillance by active and passive means at all times, under all conditions.

Global Awareness will be powerfully enabled by Air Force investment in technologies that will enable large sparse apertures, evolving in the direction of clusters of cooperating satellites. Such clusters will enable aperture sizes that are bigger than those now only available with large satellites. In addition, constellations of large numbers of smaller satellites will allow economy of scale in production and have reduced vulnerability relative to single satellites. Also important to Global Awareness are the technologies for space-based active probing such as synthetic aperture radar (for day/night all-weather coverage) as well as technologies for passive probing through multispectral sensors. These capabilities will enable any point on the surface of the earth or the air to be scanned in a wide range of electromagnetic bands

Knowledge on Demand

Knowledge on demand is the idea that at an individual warfighter could request knowledge about some area of operations. The warfighter has always benefited from having strong situational awareness of all situations in which he or she is called upon to fight. The human mind is very capable of assessing patterns in information and using those patterns to make decisions. As the infosphere envisioned by the commercial world develops, there will be a plethora of information available at many levels to a warfighter and to the commanders of U.S. forces. Indeed, there will be so much information to collect, analyze, assess, synthesize, and disseminate that the quantity will be overwhelming. What the warfighter needs is not information, but knowledge. Knowledge will come from a fusion of information from all types of sensor sources (air, ground, and sea as well as space) together with communications to deliver knowledge to the warfighter.

The warfighter could request to see all the new threats in an area or an update on old threats or new targets. The request would be entered into a global integrated information system and, if appropriate, a space-based set of sensors would provide the knowledge. The communication would be direct to the system, the request would be processed by the system, the data would be collected by the system, the knowledge

extracted from the information gathered by the system and that knowledge would be sent to the warfighter. This use of space- and ground-based assets combined with Global Awareness will enable direct and timely readout to tactical users. This integrated use of space based assets is one of the aspects of information dominance and information warfare. The technologies that will enable Knowledge on Demand are the technologies of image processing, high data rate antijam communications, data fusion, artificial intelligence, neural networks, and distributed processing.

Space Control

The value of space systems and the advantages that they will give to the United States will be so large that an adversary would be foolish not to target the space assets. In the next several years, the technology to selectively target an individual satellite will have proliferated all over the world and will be available to anyone at a relatively modest cost. Space systems can be targeted by a determined adversary with electronic warfare, high power microwaves, lasers, and with collateral nuclear weapons and kinetic energy kill vehicles as ballistic missile technology proliferates. The use of these degradation mechanisms can be made sufficiently precise so as to allow a whole range of options ranging from temporary blinding of a sensor to permanent destruction of a sensor to physical destruction of a satellite. With this range of technology available in the world, it is important that the Air Force invest in the technologies for space control in a hostile environment. These technologies will allow U.S. systems to survive and function in the kind of hostile environment that almost any adversary will be able to create in the future. The distributed satellite systems necessary for Global Awareness and Knowledge on Demand will be inherently survivable since functions will be spread among many satellites.

Space control can be divided into three technology areas: space asset surveillance, space asset negation, and space asset protection. It is important to know what assets are in space, to determine what capabilities they have, and to be able to distinguish them from background chaff and debris. The technologies that will enable effective space asset surveillance are sensor technologies. Once assets are identified, it may be necessary to undertake negation of an adversary's assets using directed energy, kinetic kill vehicles, or information warfare. The technologies that will enable space-based asset degradation are autonomy technologies will enable a smart interceptor to be released from carrier spacecraft and then accomplish a mission to degrade a specific

satellite without requiring ground control. However, it is also important to increase the survivability of friendly satellites. Space asset protection is necessary against both natural threats such as orbital debris and radiation as well as human-generated threats. Threats can be handled by making satellites difficult to find, track, and then damage or kill. The technologies to substantially enhance survivability are low observable technologies and maneuvering technologies (which require high power). Space-based directed-energy weapons for the protection of space assets will be enabled by the technologies of high power generation in space.

Force Application

The application of force from space to ground or air will be feasible and affordable in the next 20-30 years. Force application by kinetic kill weapons will enable pinpoint strikes on targets anywhere in the world. Such force projection will enable Global Awareness to extend to global presence. The current Air Force mission area of force application includes both nuclear and conventional deterrents to place adversary terrestrial targets at risk. The technology for precision kinetic energy strike of fixed terrestrial targets from space-based or ballistic missile platforms is available to the United States now. Technologies such as Micro Electro-Mechanical Systems (MEMS) could substantially improve the affordability of such systems. Technologies for similar conventional strike of mobile targets is possible given the appropriate targeting and command and control. Discussion of this kind of capability so far has focused on very limited capacity for a narrow range of targets. However, the technology suggests the possibility of a dramatic change in the means available for global power projection, making logistic delay negligible and recovering the investment in energy for logistic deployment directly as destructive energy on targets. Force application by means of directed energy weapons will be feasible if the Air Force invests in the power technologies for large amounts of power generation and energy storage. The technology to enable this application will not be developed by the commercial sector and must be developed by the Air Force.

The equivalent of the Desert Storm strategic air campaign against Iraqi infrastructure would be possible to complete in minutes to hours on more or less immediate notice. U.S. perspectives on this kind of capability are colored by past investment in conventional force projection and by Cold War attitudes about deterrence. The use of ballistic missile platforms for conventional strike raises an ambiguity in

nuclear deterrence that would have been destabilizing in the bipolar Cold War context. Use of orbital platforms for conventional strike raises a similar ambiguity with respect to treaty verification for the treaty banning weapons of mass destruction in space. While these ambiguities may give U.S. decision makers pause, it should be realized that the opportunity for others to exploit this avenue to global power will be readily accessible to the large community of nations achieving access to space. Awareness of this opportunity should help motivate the importance of space control and missile defense.

RECOMMENDATIONS FOR INVESTMENT IN SPACE TECHNOLOGY

A combination of targeted Air Force investment and adaptation of commercial development will enable a revolutionary change in Air Force space capability. Such change in the next 20 years will be affordable and based around Global Awareness, Knowledge on Demand, Space Control, and Force Application. Air Force technology investment must be carefully directed to provide the greatest return.

Revolutionary Technologies in Which the Air Force Must Invest

Several key technologies offer the possibility of a substantial increase in the exploitation of space by the Air Force, the potential impact of which is so great that the Air Force must invest now. These technologies are as follows:

- High energy density chemical propellants to enable spacelift with high payload mass fractions. Specific impulses of 1,000 seconds or greater (in high-thrust systems) are desirable.
- Lightweight integrated structures combining reusable cryogenic storage, thermal protection, and self diagnostics to enable a responsive reusable launch capability.
- High temperature materials for engines, rugged thermal protection systems, and power beaming applications.

These three technologies will enable much larger payload fractions to be lifted to orbit by factors of four or more and combined with affordable operations will enable much cheaper access to orbit. Therefore they have the potential to revolutionize the launch equation and remove the significant barrier that high launch costs impose.

- High performance maneuvering technologies such as electric propulsion (thrusters greater than tens of Newtons at specific impulses of thousands of seconds at near 100 percent efficiency are desirable) and tethers for momentum exchange.

- Technologies for high power generation (greater than 100 kW) such as nuclear, laser power beaming, and electrodynamic tethers.

These two technologies will enable space-based weapons such as high power lasers, space-based radars with wide search areas, and satellites that can maneuver almost at will. They have the potential to substantially remove orbital dynamics as a barrier to where satellites can go.

- Technologies for clusters of cooperating satellites (e.g., high-precision stationkeeping, autonomous satellite operations, and signal processing for sparse apertures).

This set of technologies will enable a new vision for space applications where functionality is spread over many satellites rather than only in a single satellite. They have the potential to enable new applications from space (such as Global Awareness) at affordable cost.

Evolutionary Technologies in Which the Air Force Should Invest

The Air Force should invest not only for revolutionary change, but also for evolutionary improvements in performance or reduced life-cycle costs to its systems. The technologies that offer such benefits include the following:

- Launch vehicle technologies
 - Engines, upper stages, and solar thermal propulsion
 - Vehicle structures (e.g., Al/Li or advanced composite tankage).
- Satellite bus technologies
 - Structure technologies (e.g., lightweight structures, active vibration suppression, precision deployable structures, and software-controlled multifunction surfaces)
 - Energy storage technologies (e.g., Electromagnetic Flywheel Battery)
 - Attitude control technologies, including attitude sensors and ACS algorithms
 - Radiation hardening technologies for spacecraft electronics
 - Low-observable technologies
 - Micro Electro-Mechanical Systems (MEMS).

- Sensor technologies
 - Large, sensitive focal plane arrays and associated readout and cooler technologies for ultraspectral sensing of small low-contrast targets and long wavelength detection against the cold background of space
 - Active sensor technologies (e.g., large lightweight antennas, high-efficiency RF sources, and high-energy semiconductor lasers for SAR, MTI radar, and LIDAR)
 - MEMS (including on-chip optics).
- Communications technologies
 - Very high rate, long-distance optical communications
 - Multi-beam adaptive nulling antennas.
- Data fusion technologies, including automatic target recognition
- Space-based weapons technologies
 - Laser weapons technologies (e.g., large lightweight optics)
 - Smart interceptors (e.g., autonomous guidance)
 - RF weapons technologies (e.g., lightweight energy storage) for EMP and jamming.

Commercially-Led Technologies

Another set of technologies that will allow for evolutionary change in Air Force space operations will be driven by the commercial sector. These technologies merit minimal investment by the Air Force, yet the Air Force should do what is necessary to adapt the following technologies to its needs:

- Small launch vehicles
- High-efficiency energy conversion and storage
- High-data-rate RF communications
- Technologies for debris reduction
- Information storage, retrieval, and processing technologies and protocols
- Image processing, coding, compression, and VLSI architectures
- Neural networks and artificial intelligence
- Technologies for spacecraft manufacturing
- Technologies for vehicle and spacecraft operations.

CONCLUSIONS

In this report, the following have been shown:

- The international exploitation of space services will grow.
- The Air Force will be able to take advantage of complementary commercial investment.
- There are revolutionary technologies that will enable a new vision for the Air Force in space.
- To effectively support the warfighter from space, active, sustained investment in these technologies is essential.

Annex A

REVOLUTIONARY TECHNOLOGIES

Annex A

REVOLUTIONARY TECHNOLOGIES

REVOLUTIONARY LAUNCH VEHICLE TECHNOLOGIES

The highest-leverage technology area impacting launch vehicles is the development of high-energy-density materials (HEDM) for use as propellants, which offer the promise of significantly reducing both booster and upper stage weight and hence lowering launch costs per pound of payload delivered. While the existing USAF program of this name is oriented primarily towards rather modest near-term improvements in Isp, (in the range of 7-20 seconds), much more effort should be devoted to revolutionary advances in this field. If Isp could be increased from today's limit of about 450 seconds to nearer the theoretical limit of 1500 seconds, payload mass fractions would increase (all other things being equal) by approximately a factor of about six. While it is still not clear how this might be done with propellants capable of being contained in a reasonable structure, such an achievement would fundamentally change the nature of spacelift. HEDM is already examining one advanced concept to use metals as fuel additives. Metal atoms or small molecules stored in liquid hydrogen, for example, could yield Isp gains of 50 or 100 seconds, resulting in a 25 percent increase in performance in a system such as the STS. The long-term goal for energetic propellants should be an Isp increase of much greater than 100 seconds. Material systems such as metallic hydrogen would produce such increases in Isp. The use of computational chemistry techniques to enable the design of even more energetic propellants needs to receive increased emphasis, because such technologies could enable missions thought to be impossible by today's standards. The potential benefit of HEDM technology justifies a significant investment by the Air Force.

Another high leverage technology area within launch is concerned with materials which can withstand extremely high heat loadings without failing. Very high heat loadings are generated in rocket engines. Engine specific impulse (which scales as the square root of the maximum temperature allowed in the combustion chamber) and thus engine thrust-to-weight ratios could dramatically increase if materials that could withstand extremely high heat loadings were available. In addition to increasing

maximum specific impulse, these materials would enable critical portions of the engine to be operated without the added complexity and mass of an active cooling system. Such dramatic improvements may enable such revolutionary concepts as the scramjet engine to become a reality; a scramjet would have the distinct advantage of taking its oxidizer from the atmosphere rather than carrying it along and hence could devote more of the vehicle mass to payload. The potential benefits of the development of such materials justifies substantial sustained development.

In addition to the need in rocket engines, high temperature materials that do not require active cooling are a vital area of technology investment for advanced thermal protection systems. These are essential for military TAVs and RLVs. Advanced thermal protection systems become particularly important if air-breathing (e.g., Scramjet) propulsion concepts are used for a large portion of the TAV flight profile within the atmosphere. For advanced vehicle concepts, material systems capable of 4000K on one surface and cryogenic temperatures of 4K on the other side will be required for use in engines and structures.

Vehicles may be designed with the cryogenic tanks integrated for vehicle load structure, such as wing elements, while still requiring a thermal protection surface. For RLVs, the high temperature materials will be used for the thermal protection system will also have a requirement to be operable and more importantly, field repairable. Integrated structures combining reusable cryogenic storage with a thermal protection system would reduce the overall dry weight of the vehicle. In addition, the integrated structure would include self-diagnosing sensors, enabling the vehicle to report on its own condition, which is critical to short turnaround time. Thus lightweight integrated structures combining reusable cryogenic storage, thermal protection, and self-diagnostics to enable a responsive reusable launch capability are revolutionary technologies in which sustained investment must be made.

REVOLUTIONARY TECHNOLOGIES FOR SATELLITE BUSES

Electric propulsion is a revolutionary technology that can enable moving spacecraft to different orbits, executing orbital plane changes, and performing routine spacecraft attitude changes. Electric propulsion has a tremendous potential for reducing spacecraft weight, and that allows the use of smaller launch vehicles with dramatic cost savings. There are three types of electric propulsion: electrothermal (e.g., arcjets), electromagnetic (e.g., plasma engines), and electrostatic (e.g., ion engines). A typical

specific impulse for arcjets is 450-1000 sec, 1500-2500 sec for plasma engines, and 2000-3500 sec for ion engines. The three categories of thruster technologies are shown in Table 1. Thrust levels are currently very low (fractions of a Newton) and need to be improved for many applications. The power required for an electric propulsion system is proportional to the specific impulse and could require tens of kilowatts of power.

Table 1. Classes of Electric Propulsion Systems

Thruster	Specific Impulse (s)	Thrust (N)	Propellant
Electrothermal			
Resistojet	300-850	0.125-0.5	N ₂ H ₄ , H ₂
Arcjet			
1-10kWe	450-850	0.17-0.23	N ₂ H ₄ , NH ₃ , H ₂
10-30 kWe	700-1400	1.0-2.2	
Electrostatic			
Ion Thruster			
1-5 kWe	2000-4000	0.04-0.2	Xe, Kr, Ar
5-20 kWe	2500-6000	0.2-0.6	
Stationary Plasma Thruster (SPT)	800-2500	0.02-0.08	Se, Kr, Ar
Electromagnetic			
Pulsed Plasma Thruster (PPT)	200-1750	17.8 μ -3.5	Teflon
Magnetoplasma-dynamic (MPD)	2000-6000	20-100	Ar, H ₂

Electric propulsion (EP) has nearly 30 years of space flight experience, during which time thruster designs have matured as improvements based on flight tests and on new technology have been incorporated into operational systems. Nevertheless, EP has so far played only a limited role in military space systems. Technical concerns have included thruster performance, power availability, guidance, navigation, and control (GN&C), and spacecraft interactions. Non-technical issues have included development costs, scheduling, mission constraints at block changes, and lack of familiarity with the strengths and limitations of EP. Nevertheless, EP makes increasing sense as the size of satellites decrease and technology continues to advance.

Advances in solar-electric power, autonomous GN&C, and electric thruster technology can support an expanded role for EP that will help meet the challenge of new mission applications, including advanced space control techniques. The ability to reposition or reconstitute satellites (without a significant penalty to operational life) is

needed by military commanders during quick-response deployments. Past EP flights have focused on low-power thrusters for small velocity-change maneuvers. Today, high-power thrusters and solar arrays offer the enabling technology for large velocity-change maneuvers and orbit raising without the time delays characteristic of past EP systems.

Some initial applications of the electric propulsion concept have been demonstrated in geostationary orbit, where some spacecraft use kilowatt-class arcjets to perform station keeping. This initial application is likely to be replaced by ion engines in the near future. The next payoff would be obtained by using electric propulsion for low altitude station keeping and attitude control, then extending the technology for LEO-to-GEO transfers where order of magnitude weight savings can be achieved. Again, these orbit transfers would require engines capable of handling tens of kilowatts. Electric propulsion can be used for stationkeeping in a distributed spacecraft systems or to make small continuous random changes in the spacecraft orbit to make spacecraft more difficult to track. Finally, only a small amount of propellant need be reserved to deorbit spacecraft to avoid debris.

The Air Force must fund an aggressive program to develop and demonstrate electric propulsion engines with specific impulse between 2,000 and 2,500 seconds and power handling capability of greater than 10 kW, as well as basic research into the physics of electric propulsion. This class of engine (coupled with an efficient power generation system), could enable the orbit transfer of a several thousand pound spacecraft with significant cost reduction. The coupling of high specific impulse and small size makes electric propulsion an ideal technology for small spacecraft.

Two candidate engine types that should be developed are plasma engines and ion engines. Plasma engines similar to the stationary plasma thrusters (SPT) or the anode layer thrusters (ALT) developed by the Russians could represent the first stage of development. The most important research to develop these engines will be the material selection. Any material must enable lifetimes of 8,000 hours in components such as high-temperature high-energy-density magnets and cathodes carrying over 100 amperes of current. Development efforts should include ground testing of these engines (to prove lifetimes of up to 8,000 hours) followed by space testing.

Another unique maneuvering technology that the Air Force should investigate is the use of tethers for momentum transfer. Two satellites tied together with a nonconducting tether compose a system whose center-of-mass motion is straightforward

to predict, but predicting the motion of the individual satellites becomes much more difficult.

The availability of high power on a spacecraft is an enabling capability. All current spacecraft are either power limited or restricted in some measure by inadequate electrical power. Power limitations impose restrictions on the communications subsystems and the propulsion subsystems and currently make large space-based radars and space-based weapons relatively unfeasible. A revolutionary change in capabilities will result from power technologies that can provide large amounts of power onboard satellites. Large amounts of power will be enabling on spacecraft in the same sense that large amounts of random access memory have been enabling in personal computers. When one is limited to very small amounts of power (or memory), the number of applications is limited and there are applications that are not even conceivable. If power is not an issue, then previously hard applications become easy and new applications become possible. Evolutionary development of solar array based power technologies will see improvements to tens of kilowatts on satellites over the next decades. However, all solar collection systems in earth orbit are limited by the solar constant of 1.4 kW per square meter. Large powers from solar collectors require large collection areas. For substantially larger powers (> 100 kW), several different types of technologies will have to be explored. Powers of this level will make large space-based radars, space-based directed-energy weapons and the use of high-performance electrically driven maneuvering technologies possible. A natural technology to enable high power is nuclear power in space; however, this technology has been seen to be unacceptable because of political and environmental limitations. Thus it is desirable to develop other technologies that may provide large power levels in space. In addition to continued development of safe nuclear systems, two other sources of continuous power in space that should be explored are the concepts of electrodynamic power generating tethers and power beaming from one location to another (e.g., from space to space). The development of these and other technologies for high continuous power will have a revolutionary effect and the Air Force should invest in these as well as continue to invest in solar collection technologies.

Over the years, there have been several programs in nuclear-powered spacecraft. NASA has been using Radioactive Thermal Generators (RTGs) for the interplanetary missions which generate few tens of watts of power. Russia has flown nuclear reactors in space, and BMDO has a joint program with the Russians (TOPAZ), where the Defense

department bought three of the reactors to do some laboratory experiments. DoE had a program, SP 100, to use nuclear power in space along with an Air Force program for nuclear propulsion. These programs have been canceled. However, nuclear power remains as one of the attractive alternatives in generating large amounts of power in space. To build a reactor for space applications has many challenging technical aspects including development of high-temperature light-weight materials, active cooling technologies, extremely radiation hard and high temperature electronics, and fail safe system architectures. Setting the emotional issues of nuclear power aside, this technology offers a viable alternative for large amounts of power in space. The Air Force should continue the efforts towards making a safe nuclear reactor in space a viable option. Existing joint programs with Russia offer a low cost alternative and should be pursued.

Electrodynamic tethers are basically long wires that are drawn across the earth's magnetic field. Just as in a electrical generator, the motion of a conductor across a magnetic field will cause a voltage to be generated. If a current can be made to flow from the ionosphere through the tether and close back in the ionosphere then power can be generated. This power comes at the expense of orbital energy since the tether feels a drag force which acts to slow it down. Thus the tether effectively changes orbital kinetic energy to electrical energy. Hence a continuous power system would be composed of a tether and a thruster to reboost the orbit. Alternatively, a system can be designed that uses a tether to extract energy during part of an orbit and then reboosts during another part of the orbit. Electrodynamic tethers can also be used as thrusters by reversing the current flow through the tether with an onboard power supply. In addition, since all electrodynamic tethers are tethers, they can also be used for momentum exchange between two tethered spacecraft. Since electrodynamic tethers work by using the voltage drop that comes from moving across the earth's magnetic field, they are limited for effective use to orbits where the field is strong enough to give reasonable voltage drops. This limits tethers to orbits below a thousand kilometers from the earth's surface. There are many technical issues to be resolved with high power electrodynamic tethers. These include the extraction of large currents from the ionosphere (tens of amperes), the emission of such large currents back into the ionosphere and the dynamic stability of such large unidimensional conductors in orbit. This technology offers one high risk, high payoff way to achieve high powers in space and should be pursued.

Power beaming to a spacecraft using high power lasers offers another option for obtaining large quantities of power in space. In one concept, a high-power ground-based

laser is used to form a collimated beam onto the spacecraft. Multiband gap solar arrays on the spacecraft convert the laser power into onboard electricity for the spacecraft. In another concept, a space-based laser driven by a large solar array or a nuclear reactor could be used to beam power to another spacecraft. In both concepts, in order for the receiving spacecraft to have small arrays, the arrays must be capable of processing equivalent power densities greater than 100 suns (140 kW per square meter). This would enable hundreds of kilowatts to be received on an array of the order of a few square meters. The limitation on this is the semiconductor materials which can convert such large power densities to electricity without large heat losses or without suffering permanent damage. The Air Force should invest in the basic research necessary to develop such materials as well as in the pointing, tracking, and continuous high power generation in a laser device. As this technology matures, this concept may offer promise for transmitting high powers to spacecraft.

REVOLUTIONARY TECHNOLOGIES FOR SATELLITE CLUSTER

The development of low cost, single function satellites offers new horizons for space applications when the satellites operate cooperatively either in clusters (local formations of satellites) or in constellations (satellites distributed both within an orbital plane and over a set of orbital planes). The vision of what can be achieved from space is no longer bound by what an individual satellite can accomplish. Rather the function is spread over a number of cooperating satellites. Further, these distributed systems of satellites allow the possibility of selective upgrading as new capabilities become available in satellite technology.

A detailed analysis shows that while distributed systems do not provide a cost-effective answer for all applications, there are many applications for which small sensors on many satellites scale very well and give cost-effective solutions. As an example, passive scanning imagers on dedicated satellites or communication constellations scale very well indeed. In addition, distributed systems have distinct advantages in survivability. This results from the distribution of capability over all the components. Individual satellites, once found and tracked, can be easily destroyed from the ground by high energy lasers or by kinetic kill vehicles. Distributed systems will degrade in proportion to the number of satellites lost. The flexible and proper interconnection of the rest will make the overall system intrinsically survivable. The development of Global Awareness will require an array of collectors with all-weather sensing. For example,

frequent revisit SAR of mid to low latitudes with one meter resolution could be achieved by a small constellation of low inclination, low altitude small satellites. These would provide all-weather, day/night observation capability. In addition, 1-meter mid-wave infrared, 2-meter long-wave infrared and 2-meter multispectral data could be provided by a constellation of single-purpose small satellites. Another possible use is of bistatic SAR in which a microwave illuminator is placed in a synchronous orbit with lower orbiting receivers or airborne receivers. This concept could be implemented as a low-cost space-based surveillance concept involving the use of approximately 10-20 LEO satellites for SAR coverage of a theater.

One of the revolutionary effects of the technologies which enable clusters of cooperating satellites will be the ability to flexibly form extremely large (in wavelengths) coherent apertures in space for sensing, communications, and weapons. The development path to clusters begins with systems of interconnected, cooperating satellites such as Iridium or Teledesic whose constellations will distribute function across their orbiting networks to provide global communications. The applications path to coherent clusters of satellites goes through sparse distributed aperture sensing satellites. The mission need driving the technology is the need to sense continuously the target and background environment in an area of interest. To provide continuous viewing opportunity over arbitrary spots on the globe requires constellations on the order of 100-1,000 satellites (depending on the viewing angle constraints) at altitudes on the order of 1,000 km. At altitudes on the order of 10,000 km, the number of apertures shrinks to the order of 10-20. At geosynchronous altitude the number of apertures reduces to the order of 3-10 depending on the need for high latitude coverage. For imaging applications, the aperture dimension required to maintain resolution scales directly with the distance to the target. However, the aperture may need to be only sparsely filled where the energy received is not the limit, e.g., with illumination from the sun. At low altitudes, a monolithic aperture may be reasonable. At moderate altitudes, a sparse, distributed aperture on deployable structure may provide equivalent performance. At higher altitudes, a cluster of cooperating satellites flying in formation can form the aperture dimensions required without the weight and cost penalty of a satellite subtending the entire aperture. The requirement on the cluster elements is to maintain autonomously relative positioning, attitude, and communication among the elements well enough to allow the aggregate to maintain phase coherence over the aperture. The distributed system then becomes a constellation of clusters.

Having achieved the technology for clusters of satellites cooperating for passive sensing, the same ability enables revolutionary change in active systems for active sensing, communications, and weapons. For active apertures for sensors, communications, or weapons, the aperture may be thinned but not sparse to the degree that the power (and waste heat) radiated per element is too high. Instead of a relatively small number of cooperating elements in a cluster, these applications will drive towards large numbers of identical cooperating (perhaps docked) elements which permit significant economies of scale in manufacturing and flexibility in launch. An example application of this approach is an alternative path to the frugal Global Precision Optical Weapon (GPOW) space-based laser. Rather than large monolithic flexible optics directing the beam from a single large laser powered by 10 percent efficient solar to chemical energy collection/storage and 25 percent efficient laser conversion, a clustered approach would employ phased diode lasers (like the Fotofighter concept) with 50-90 percent efficient laser conversion and 20-30 percent solar electric energy collection. This approach can also be applied to the generation of very intense RF beams from a set of separate elements on different satellites with the precision stationkeeping to enable all elements to radiate coherently. Such an intense RF beam could be used to overcome local jamming or to burn out sensitive electronics.

Thus, revolutionary capabilities will be enabled by the use of distributed systems and the Air Force must invest in the technologies for clusters and constellations of cooperating satellites (e.g., high-precision stationkeeping, autonomous satellite operations, very high performance communication links, distributed processing, and signal processing for sparse apertures).

**C. GATHERING POSITION INFORMATION USING RADIO
TAGS AND GPS**

**Gaetano Borriello
Department of Computer Science and Engineering
University of Washington
Seattle, Washington**

SUMMARY

The global positioning system (GPS) has developed into a major tool for the military. It permits the gathering of accurate position information for military vehicles and personnel on the battlefield. Knowledge of precise and even approximate position of personnel can aid immensely in command decisions and in reducing fratricide. Operation Desert Storm (ODS) provided first hand proof of the great utility of position information garnered from GPS. The venue of ODS, an open battlefield with clear line of sight to the GPS satellites, was an ideal proving ground.

Today's U.S. Armed Forces are being asked to perform more and more duties that are in venues with less than ideal conditions for satellite communication. Specifically, activities in urban environments involving peacekeeping duties (e.g., the former Yugoslavia) and pacification of cities (e.g., Haiti and Somalia) provide additional challenges to the gathering of position information. These difficulties stem from the need to monitor the position of individual soldiers rather than armored vehicles and the fact that GPS has some inherent limitations. Among these are that it can be jammed easily and cheaply over a very wide area and that building interiors do not permit unobstructed lines of sight to a minimum of three satellites (for 2-D positioning) [2].

These problems make it virtually impossible to rely on GPS for position information in potentially hostile urban (or even in the simple case of heavily foliated) environments. The lack of situational awareness (SA) in these operations has proven to be an important source of fratricide and a major factor in political incidents (e.g., harassment of locals, firing on crowds, etc.) involving U.S. forces. Accurate position information on each warfighter could prove to be invaluable information to both individual soldiers requesting reinforcement, to their commanders in making command decisions relying on rapid redirection of personnel, and in providing an auditing trail for search missions as well as post-incident investigations to other authorities.

Radio tags currently being used by several of the services for inventory control could be an interesting supplement to GPS in this arena [5]. This off-the-shelf technology is based on an interrogate/reply system that uses signal strength to determine approximate distance [6]. Organized in a cellular structure, interrogator units could be

used to provide the raw data needed to construct a reliable, albeit approximate, location map for an entire platoon or even larger contingent. Radio tags are primarily passive devices that respond only when interrogated thereby not divulging a soldier's position except when explicitly requested to do so. Most importantly, radio tags work well indoors as they were designed for use inside packages and containers stored in warehouses and trucks [7].

This combined system – radio tags with GPS – is the subject of this investigation. The result of this very preliminary study is that such a prototype system should be constructed and evaluated. The tag technology is already in the field. Only minor modifications and enhancements are needed to ready the equipment for a prototype system. A trial would involve the development of new software to deal with the data gathered by the interrogators. Verification can be accomplished by using a MOUT operation to simulate an urban battle environment and compare the system's performance against actual battle video footage. If it proves robust, it could provide an important augmentation to the already elevated functions of wireless communication that will be available to the 21st century soldier [3].

I. INTRODUCTION

A. SITUATIONAL AWARENESS

In recent years, one of the major themes in U.S. military operations has been enhanced situational awareness [3]. Field commanders are demanding the ability to know the position of all their assets including vehicles and materiel. The focus has been primarily on battlefield location with an emphasis on large-scale engagements involving hundreds if not thousands of locatable units.

The global positioning system was developed specifically for the purpose of providing accurate location information to aircraft, naval, and ground vehicles. The accuracy of GPS is to 30m as this was and is sufficient for large objects. The system was heavily used in ODS to aid in reducing fratricide and to permit the rapid redirection of forces on the battlefield.

Not only is the need for location information supported by battlefield experiences but also from problems in completely different venues. For example, incidents in Somalia and in Haiti have highlighted at least two more reasons why it will be more and more important to have position information for individual soldiers. The first concern is for incidents of fratricide in urban environments where soldiers are not aware of each other's position to the extent required for safety. Even though fratricide in ODS primarily involved vehicles, data show that dismounted infantry fratricide was the predominant problem in Vietnam and WWII [1, 4]. The second concern is the involvement of soldiers in political incidents where accountability or even criminal prosecution may be necessary. As U.S. forces are called upon to operate in urban environments and perform more and more peacekeeping duties, it is reasonable to expect that both of these concerns will only be more pronounced. Position information with a granularity of 10m will be required.

B. GPS

Unfortunately, GPS will not be able to help in these situations, at least not by itself. It will need to be combined with another more localized positioning system. The reasons for this are twofold.

First, a GPS receiver requires a clear line of sight to a minimum of three satellites in order to determine 2-D position information (point on earth's surface). Four satellites are required for a position in three dimensions. The GPS satellite constellation was designed so that seven to eight satellites are in fact visible at any point on an idealized earth's surface. The receiver then chooses to use the three to four with the strongest signal.

This model is perfectly appropriate for what GPS was originally intended, namely, to aid in navigation for aircraft and naval vessels. However, land masses pose some difficulties. Terrain such as canyons can obstruct enough satellites to make the system ineffective. Heavily foliated environments are also a problem as the GPS satellites' signals are only 10-20W at a frequency of 1.5GHz and as such are not able to penetrate dense or even moderate foliage. The system is rendered almost completely ineffective in dense urban environments. Not only is it unusable in the interior of buildings or underground but even buildings that are close together (as in many cities around the world) can obstruct enough satellites. In addition to electrical degradation of the signal, there is also geometrical degradation due to multi-path effects created by the many surfaces against which the signal bounces in an urban settings. In combination, these two effects can degrade GPS positioning by an order of magnitude [2]. Therefore, even if readings are possible, the accuracy may be reduced to as little as 300m. Clearly, this is not accurate enough for a platoon of soldiers operating in a dense urban area.

Second, GPS can be very easily jammed over a wide area. Since the signal from the satellites is of such low power, it takes a correspondingly modestly powered jamming device to completely swamp the signals. Moreover, the jamming is effective up to the horizon thus covering a very large area. Such a jammer would have to operate continuously as any suspension of jamming enables GPS receivers to reacquire their positions. Such a requirement on the jammer might appear to make it an easy target for elimination. This is certainly the case for air or sea operations where objects are relatively far apart and often over the horizon prior to engagement.

Again the situation is different on land. The low power requirements for the jammers makes it possible to build them inexpensively (approximately \$75) from parts obtainable from any radio hobby shop and be of a very small size (the size of a hockey puck). A popular hobbyist's magazine has even provided the detailed schematics and parts list for building such a device [2]. On land, a large number of jammers can be easily scattered over a wide area thereby effectively disabling GPS where it is most

needed – on the battlefield. In cities, the problem is even worse, as they can be virtually impossible to locate and at the very least require a time-intensive search.

C. RADIO TAGS

Clearly, GPS alone cannot solve the problem of mapping individual soldiers in an urban environment in real-time. Fortunately, another wireless technology can be combined with GPS and help mitigate many of the problems discussed above.

Logistics including shipment and indexing of inventory are the bane of many military operations. In recent years, the Army and Marine Corps have applied radio tags, a commercial off-the-shelf technology, to address this problem [5]. Radio tags are small (the size of a deck of cards), battery-powered, quiescent devices that are attached to crates, vehicles, or even individual items in a ship container [6]. The tags lie quiescent until they receive a radio signal from an interrogator unit. When this occurs, the tag awakens and sends back a reply to the interrogator containing its ID as well as any other information that may have been placed in its memory (e.g., contents of the crate, vehicle maintenance record, etc.). Interrogators are typically placed at the entrances of warehouses or trucks to keep track of what has entered a particular location. The interrogator periodically awakens the tags so that objects passing by are logged into a central database.

The tags can also be used in locating a particular object in a large storage area—an enhancement to the original tags requested by the U.S. Department of Transportation [5]. This is possible because the interrogators also measure the strength of the signal returning from the tag. By using a hand-held interrogator, the relative strength of the signal can be used to provide a hot-cold indication that instructs the person searching for the object that they are moving closer or further away.

The radio tags used in the U.S. Armed Forces are manufactured by Savi, Inc., of Mountain View, California [6]. They use a frequency of 433MHz (near the middle of European and U.S. standards for low-power operation) and are FSK modulated. At this frequency, the signal travels through virtually all building materials [7]. This was a design requirement due to its use in warehouses, container yards, and transportation vehicles. Multi-path effects actually help the tags as they help extend the range of the transmission. The tags are very low-power and designed to last for up to 5 years on their internal batteries. The power of the transmitted signal is approximately 10mW with a corresponding range of up to 100m. The small size of the cell is needed in order to

provide fairly accurate position information with evenly spaced interrogators networked in a cellular arrangement.

II. A LOW-COST URBAN POSITIONING SYSTEM

A. SYSTEM GOALS

The system will provide the capability to determine the position of approximately 100 soldiers in an area of one square kilometer within an accuracy of 10-20m. It must be possible to maintain up-to-date information within a few seconds of real time. There shall be no restrictions on the type of obstructions in the area including, but not limited to, buildings (of various materials) and foliage. The additional weight incurred by each soldier must be under two pounds and the incremental cost under \$200. A soldier's position should not be compromised due to RF emissions.

B. PROPOSED SYSTEM OVERVIEW

The system consists of two principal components: the hardware and the software. The hardware essentially exists today in the form of the Savi Tags and Interrogators but may need to be modified to increase its operating range by a factor of 5 [6]. Another issue is the data transfer rate which also needs to increase by at least a factor of two. These modifications should not pose significant problems and are to be expected to have a minimal impact on size and cost. The software for this application does not yet exist but should pose no significant problems that cannot be handled easily by today's embedded microprocessors.

There are three types of hardware: a fixed interrogator that can be placed within or on the perimeter of the area of operation, a portable interrogator to be worn by some subset of the soldiers involved in the operation, and a portable radio tag to be worn by all personnel. GPS units will be attached to all the interrogators (and in all likelihood to each soldier) [3].

The interrogators will periodically query all tags within their receiving range. The Savi Batch Collection algorithm allows one interrogator to collect the IDs of hundreds of tags virtually simultaneously [6]. Once the interrogator has the IDs of the tags in its range it can then interrogate each tag individually to measure the signal strength and thus approximate distance (within some interval of uncertainty that needs to be experimentally

determined [7]). The fixed interrogators provide anchor points for the mapping process and require either an accurate position either acquired via GPS or entered manually. The fixed interrogators should be small enough devices so that they can be easily hidden from view (once they acquire their location).

At this point, the network of interrogators will coordinate to distribute the distance measurements each interrogator has collected so that each interrogator has a complete copy of the data (time-stamped to avoid stale data). Each interrogator will forward its data to all its neighbors so that, incrementally, each will build up a more and more complete set of data. The software, residing in each and every interrogator, will then begin processing the measurements to determine the position of each soldier's tag. These final positions (or rather, most likely positions) will be re-broadcast by each interrogator so that each tag will get information about every other tag's position. A small LCD screen can be easily connected to an individual soldier's tag to graphically display the positions of other members of the platoon superimposed on a map of the area (if available and ideally including buildings and other landmarks).

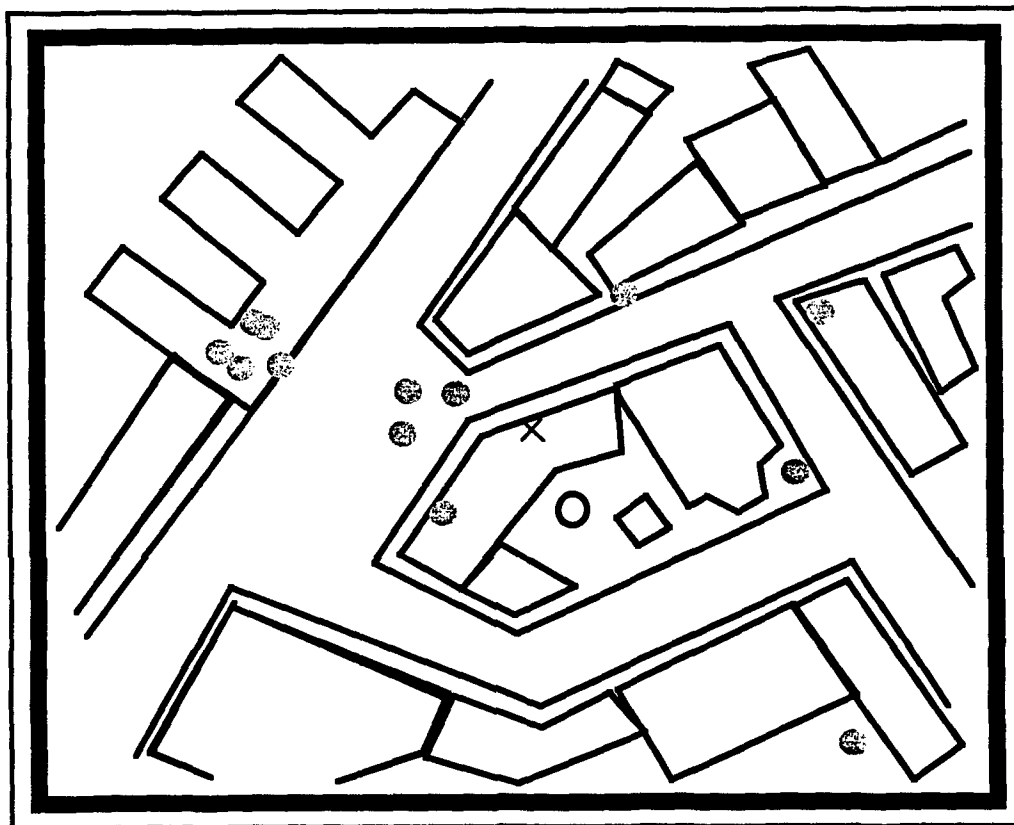


Figure 1. Soldier's Lcd Screen Showing Map of Area of Operation and Position of Comrades (shown as light gray regions on map) and Self (cross at center).

The system's software will be running in parallel in all the interrogators and performing exactly the same operations. The results should be identical, and a quick verification during broadcast can provide information as to the integrity of an interrogator's results. Rather than operating directly on this large amount of data, the interrogators first select the three best strength readings for each tag. This is done by choosing those nearest to being half the maximum range as these are likely to be the most accurate.

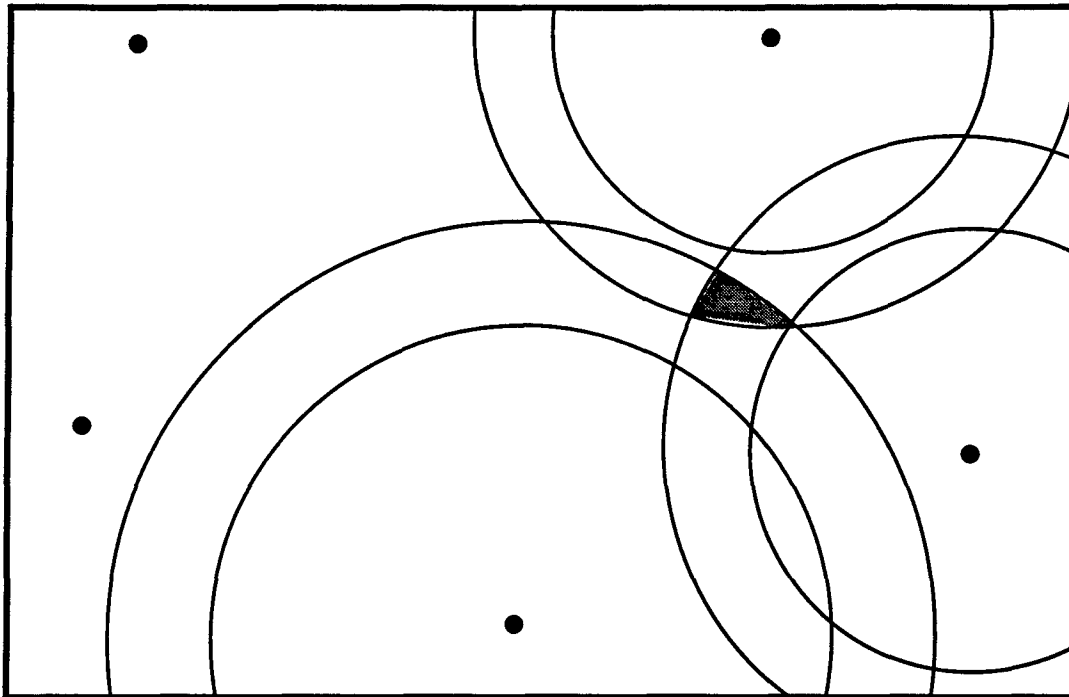


Figure 2. First Level of the Hierarchical Triangulation.

[Position of fixed interrogators (black dots) is used to triangulate the position of the mobile interrogators (with some uncertainty shown as shaded region)].

Triangulation is now performed in a hierarchical manner. The position of the mobile interrogator units is determined first. This is likely to be more accurate as these units will have better distance measurements due to their higher signal strength. Furthermore, ambiguities in the triangulation due to missing data (e.g., only two distance readings to a tag) can be resolved using some auxiliary information about the physical site. For example, knowing that two interrogators are at the perimeter of a region of operation can disambiguate between the two possible positions that would be obtained

from only two distance readings. Similarly, limits on the distance above and below ground can help disambiguate points in three-dimensional space.

Once the position of the mobile interrogators is determined, the distances of the tags to these units can be used to determine the positions of individual tags. Without this hierarchical organization, the data would be more difficult to use and require much more computation. Simple geometrical triangulation and data filters are all that is needed for each interrogator to determine the position of every tag, which it will then broadcast. Memories in the fixed and/or mobile interrogators can be used to keep a log of all the calculated positions for later playback.

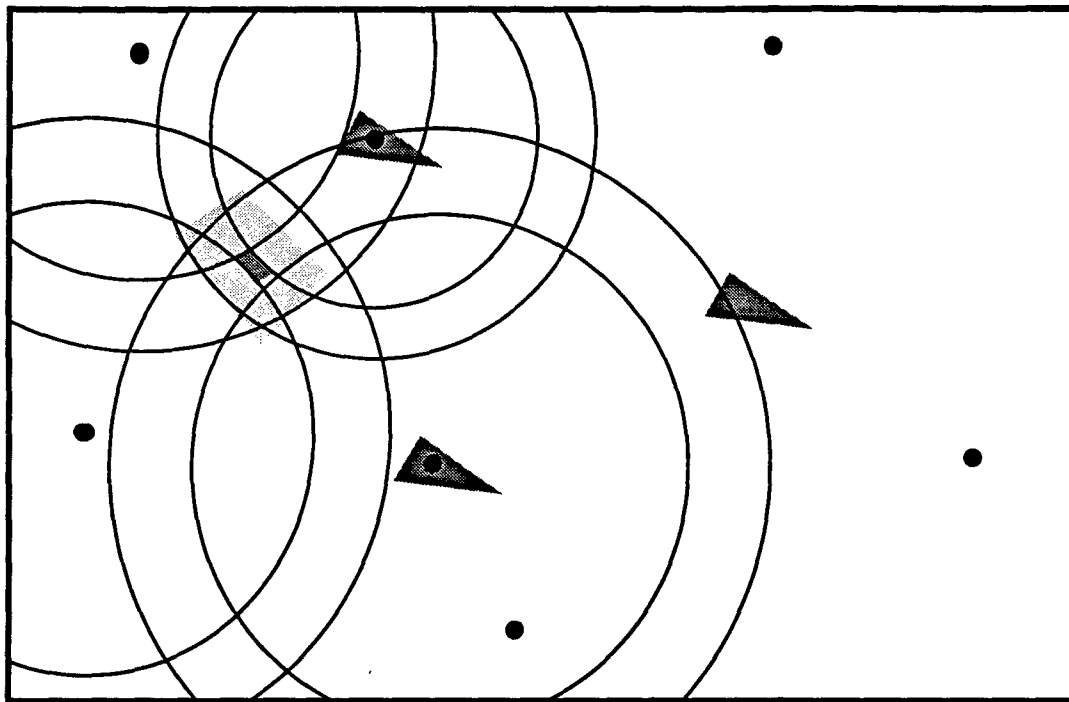


Figure 3. Second Level of the Hierarchical Triangulation

[Position of fixed interrogators (black dots) and mobile triangulators (darker shaded regions) is used to triangulate the position of individual tags with higher uncertainty (lighter shaded region)].

C. DESIGN AND IMPLEMENTATION ISSUES

Many issues arise in the development and deployment of such a system. These range from adequate data rates to deployment of the fixed interrogators.

The current generation of radio tags have data rates of 9600 bits per second. Each tag needs to transmit a unique ID, a time-stamp, and any other supplementary dynamic

information about the soldier. Static information about the soldier, e.g., rank, serial number, name, etc., can be recorded onto memories residing in each interrogator or even each tag. Dynamic information might include latest GPS location and time (which can be used as another fixed reference point if recently updated), physiological condition (from body instrumentation), and status of munitions and weapons (whether they have been fired and how much). Therefore the total number of bits is probably a minimum of 64 (32 bits for the ID and 32 for a time-stamp) and a maximum that could easily reach to 256.

Therefore, the total time for a tag to respond to a query is on the order of 5-25ms. If we allow approximately the same amount of time for an interrogator query (the latency of a response should be negligible in comparison) then we can assume that it will take 10-50ms for a distance reading to each tag. The batch collection algorithm will enable each interrogator to know the tags in its range within an amount of time approximately a factor of ten times this value [7]. Therefore, the distance to 10 tags can be computed in approximately 1 second. The time required for the collected data to be propagated to every interrogator is probably comparable although it depends on the maximum number of hops required. The triangulation computations should be relatively short (i.e., much less than 1 second). The rebroadcast should take less time than the data collection phase as multiple hops are not necessary (each interrogator has all the data that must be transmitted). Thus, at current data rates, it should be possible to obtain new position information every 3-5 seconds. Clearly, this needs to be improved by at least a factor of two and more ideally by a factor of four. There is no inherent reason why this should not be possible. The reason for this choice of transmission speed in current tags is due to the common availability of serial line controllers at this data rate and not because of any fundamental limitations.

Another concern is collisions in the broadcast of queries from multiple interrogators. This is only an issue for initial broadcasts from interrogators collecting the IDs for tags in their range. A straightforward way to address this problem is to either adopt a random backoff algorithm (as is done in the Ethernet local area network) or assign a round-robin schedule to the interrogators. Which will work better is dependent on the density of interrogators. For a sparse network, the random backoff will probably be adequate.

There is no collision problem when tags respond. This is the case, because, a tag need respond only once for its signal strength to be measured by multiple interrogators. The order of queries to tags can be in ID number order. This is especially effective if all

interrogators know the ID of all the tags to be considered by the system. New tags that enter into the operating region can take a bit more time to process while their entry is propagated to all the interrogators.

The deployment of fixed and mobile interrogators is another issue that must be addressed. Fixed interrogators should be placed as much as possible around the perimeter of the region of operation as well as within the area. These could either be placed *a priori* (as would be the case for extended operations in the same region) or could be seeded as the unit moves through an area in an arrangement where interrogators move from back to front to form a moving wave. Mobile interrogators are probably most appropriate in these types of situations as their GPS position can be periodically updated.

Lastly, the range of the current tags must be increased by at least a factor of 5 in order to cover a 1km² area effectively and efficiently. The current limitations on range arise from FCC regulations. In order to make the tags exempt from FCC licensing rules, the power output had to be limited to a point corresponding to a range of approximately 100m. There are no known technical difficulties with extending the range and correspondingly increasing the resolution of the signal strength measurement. A more important issue for range is the variation of signal strength with different types and amounts of obstructions. Experiments need to be conducted to verify that multi-path effects compensate for much of the effect of obstructions [7]. Especially troublesome would be rapid fluctuations in signal strength as a tag moves from inside a room or building to outside it.

D. VULNERABILITIES

The proposed system has vulnerabilities that range from jamming to the compromising of soldier's positions. It should be possible to alleviate all these concerns.

The system uses RF frequencies that can be easily jammed. This may appear to make it just as vulnerable as GPS to this sort of countermeasure. However, due to the short range and omni-directionality of the system, it will require a much higher output jammer that will also be commensurately more expensive [2, 7]. Furthermore, the jammer will necessarily have to be close to the region of operation of the system. This will work to ensure that it will present an easily detectable target that can more quickly be reached and disabled. An additional aspect working in favor of the system is that the enemy will not know when the system will be operational. Thus, before entering the

region, a simple false activation can be used to force the enemy to activate the jammer so that it can be located and destroyed.

Spoofing is not really a concern as an enemy would have a difficult time causing a tag response to emanate from an arbitrary location. Encryption along with time-stamps and random numbers can be used to make this virtually impossible. A time-stamp is included in every query making it valid exactly once. A random number included in the query and encrypted along with the rest of the data is returned by every responding tag. Therefore, an adversary would have to decode the query in short order so as to frame a correct reply in time.

Probably the most serious concern is that the RF emissions may divulge a soldier's position. This is mitigated by the low-range nature of the signal and its omnidirectionality. The enemy would not also have to be at close range already but would also have to be able to triangulate the position of a responding tag.

Finally, captured units are also of concern. Fixed interrogators that are not associated with a soldier must be designed to disable their functions if not properly handled. This is a fairly straightforward package design problem. Tags and interrogators associated with soldiers can disable their location functions after the detection of a loss of vital signs and begin sending distress signals. Tags can even be specialized to an individual soldier's vital sign signature so that they cannot be easily moved to another individual [8].

III. CONCLUSION

Technology available today can provide a system for continuously updating the location of individual soldiers in real-time. This is especially advantageous in urban environments or where there is dense foliage. The utility of this location information is manifold and should prove invaluable to the soldiers themselves, their commanders, and civilian authorities. The soldiers can use the information to avoid possible fratricide scenarios; the commanders, for optimizing the re-directing of individuals or entire units; finally, civilian authorities may use the information in investigations of political incidents.

It has been shown above that a system combining radio tags and GPS receivers is feasible and could provide the information required. It can be prototyped in a very short period of time as all the elements (with the exception of the system's fairly straightforward software) are currently available. There are several technical issues in the deployment of such a system including the feasibility of extending the range and data rate of the radio tags and adequately dealing with transmission collisions. Verification of the system's efficacy can be obtained through video analysis of training exercises.

REFERENCES

1. Beyer, J. C. ed., *Wound Ballistics*, Office of the Surgeon General, Department of the Army.
2. Bober, J., personal communications, Institute for Defense Analyses, Alexandria, VA, 6 June 1995.
3. Fox, C. S., *U.S. Army Battlefield Combat Identification Program*, Briefing for Defense Science Study Group, Fort Belvoir, VA, 5 June 1995.
4. Joint Technical Coordinating Group for Munitions Effectiveness, *Evaluation of Wound Data and Munitions Effectiveness in Vietnam*, Aberdeen Proving Ground, MD.
5. Lundgren, A. (Maj.), *U.S. Marine Corps Applications of Microcircuit Technology in Logistics Applications*, Briefing and Demonstration for Defense Science Study Group, Camp Lejeune, NC, 11 August 1994.
6. Savi Inc., Tag and Interrogator brochures, Mountain View, CA, 1995.
7. Schwartz, S., personal communication, Savi Inc., Mountain View, CA, 30 August 1994.

**D. COMPARATIVE ADVANTAGES OF TITANIUM ALLOY
AND HIGH-STRENGTH STEEL SUBMARINE PRESSURE HULLS**

William C. Johnson
Department of Materials Science and Engineering
University of Virginia
Charlottesville, VA

ABSTRACT

PURPOSE AND SCOPE

The comparative advantages of using titanium alloys and high-strength steel as the structural material for constructing submarine pressure hulls are examined using simple design criteria. The availability of titanium alloys possessing the necessary materials properties is investigated and the cost of such alloys relative to that of high-strength steels indicated.

Although the selection of a material for a submarine pressure hull is strongly dependent upon the intended submarine use, no effort is made to examine naval operational doctrine and current or future submarine requirements.

FINDINGS

Titanium alloys offer significant advantages over high-strength steels as a structural material for submarine pressure hulls owing to their high strength-to-weight ratio, low density, good corrosion resistance and non-magnetic character. Increased operating depths of up to 80 percent over those obtained from steel hulls are theoretically possible for submarines of the same hull diameter and weight. Titanium alloys afford enhanced flexibility in hull design permitting better optimization of numerous submarine attributes including speed, detectability, payload, and independence of operations.

Ti-100 (Ti-6Al-2Nb-1Ta-0.8Mo) was identified as a candidate material for submarine pressure hulls in the 1970s. Base plates of Ti-100 and weldments possess material properties which meet or surpass most design requirements. Problems remain in that the fracture toughness of some welds are unacceptably low and titanium alloys are susceptible to stress corrosion cracking and room temperature creep. By the early 1980s, flat and horizontal 10-cm thick plates of Ti-100 could be welded providing an adequate hull thickness for some combat submarine pressure hulls but would be insufficient for special-purpose submarines. Adequate supplies of Ti can be obtained.

Use of titanium alloys as a structural material for pressure hull is probably feasible economically only for special-purpose submarines. A research program examining fabrication and joining of titanium alloys is recommended.

I. INTRODUCTION

Russia and the former Soviet Union (FSU) have deployed nuclear-powered attack submarines constructed with pressure hulls of titanium alloys since 1970.¹ Submarine classes deployed with titanium hulls include the "Alfa," "Papa," "Mike," and "Sierra" classes. The most recent Russian titanium-hulled boat, a Sierra II SSN attack submarine, began sea trials in August 1993. It is believed that no additional titanium-hulled attack submarines are planned and that the newest nuclear attack submarines, the Akula and the follow-on generation, are constructed with HY-130² steel hulls [3, 4].

Perceived advantages of employing titanium as a material for pressure hulls include increased maximum operating depth of the submarine, enhanced flexibility in hull design permitting better optimization of numerous submarine attributes, and increased difficulty of detection. These factors are interrelated and influence directly such operational factors as firepower (through submarine dimensions), flexibility and independence of operation, stealth and covertness. The FSU viewed titanium-hulled boats as offering superior performance for deep-diving missions: Soviet doctrine called for the "Alfas" to use their high speed and deep-diving capabilities to avoid pursuit by enemy forces after carrying out torpedo attacks [1].

The use of titanium alloys for pressure hulls reduces submarine detectability. Maximum reported operating depths on the order of 700-900 meters for the Russian Alfa-class submarine [1] are believed more than sufficient to avoid detection by surface and airborne platforms using magnetic anomaly detection (MAD).³ (Operating depths on the order of 400-500 meters for steel-hulled boats are considered sufficient to avoid detection by MAD [5]). Such deep operating depths have spurred sensitivity improvements in MAD on the part of the United States [6]. The nonmagnetic titanium hull further reduces

¹ The first boat of Alfa class was laid down in 1967, completed in 1970, and scrapped in 1974 owing to problems with either the titanium pressure hull or an overstressed reactor [1, 2].

² The number following the alloy designation, HY or Ti, refers to the yield stress in thousands of pounds per square inch (psi). The designation HY means high-yield-strength steel.

³ MAD identifies a submarine's position by detecting perturbations in the earth's magnetic field caused by iron-alloy components.

the operating depth necessary to avoid detection by MAD and provides additional protection against magnetic-type mines. Increased operating depths would also allow Russian submarines to operate in deeper waters as a stationary "bottom dweller" making their detection almost impossible [7].

Although Russia is no longer building attack submarines with titanium hulls, it is believed that they are currently investing heavily in special-purpose submarines with titanium hulls [8]. Russia is believed to operate at least six deep-diving submarines that can be used for eavesdropping or disruption of undersea communication links. Construction of such vessels is continuing. The most advanced hull designs are made from rolled titanium over 20 cm thick employing technologies that reportedly do not currently exist in either the West or Japan [8].

The purpose of this report is to examine some of the comparative advantages and disadvantages of constructing submarine pressure hulls from titanium as opposed to high-strength-steels. In Section II, issues of structural design are addressed and the relative advantages of titanium with respect to three high-strength steels are discussed. Other materials issues and a simple cost analysis are discussed in Sections III and IV, respectively.

II. STRUCTURAL DESIGN

The two most important aspects of structural design for pressure hulls are the external pressure at maximum depth (static collapse pressure) and the hull's resistance to shock (dynamic loading). Submarines are considered to be volume-limited: The submarine dimensions sometimes impose significant constraints on what types of weapon systems can be carried.

Submarines can be considered as hollow, thin-walled, ring-stiffened cylinders subjected to high external pressure. Under this external load, the hull experiences various types of forces including normal (compression and tension), shear, bending, and buckling. Each of these forces imposes different constraints on the hull's structural design and different demands on the material properties of the pressure hull. Figure 1 illustrates some of the various failure mechanisms associated with the pressure hull.

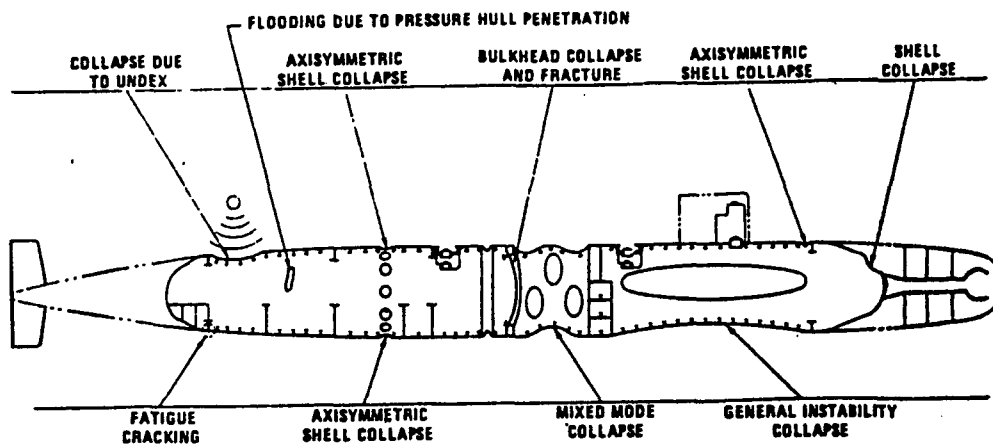


Figure 1. Different Types of Submarine Failure Modes Necessitate That Properties of Submarine Hull Material Meet Broad, Demanding Minimum Criteria [9]

The two material parameters which affect most directly the static collapse pressure are the yield stress (σ_y) and the elastic modulus (E) of the pressure hull. The yield stress is that stress at which the material begins to deform plastically. It imposes major design constraints on the maximum hull diameter and minimum hull thickness. The elastic modulus is a proportionality constant linking the (uniaxial) stress (σ) and

elastic strain (ϵ), $E = \sigma/\epsilon$, and is a measure of a material's rigidity. As such, the elastic modulus gives a measure of the deflection of the hull and hull stiffeners owing to the external pressure and relates to the bending and buckling of the components.

In the following two subsections, the relation between yield stress, elastic modulus, and hull design are presented simply. The structural advantages of using titanium alloys over high-strength steel as a pressure hull material are examined.

A. DESIGN ISSUES: YIELDING

The constraints imposed on submarine dimensions and operating depth by the yield stress of the hull material are obtained qualitatively by considering the submarine as a thin-walled cylinder subjected to an external pressure P_o , see Figure 2. Neglecting for the moment the inclusion of a design safety factor, the collapse depth is assumed to be determined by that external pressure which causes the pressure hull to begin to deform plastically. If the onset of plastic deformation is determined from the von Mises yielding criterion [10], it is shown in the annex that the collapse depth, $d_{collapse}$, can be approximated as

$$d_{collapse}(meters) = \frac{200t\sigma_y}{R} \quad (1)$$

where t and R are the hull thickness and radius in meters, respectively. The yield stress of the hull material, σ_y , is expressed in units of MPa.

For a given material (or σ_y), the collapse depth increases as the hull thickness increases and as the hull radius is decreased. *Assuming appropriate design of the pressure hull*, diving depth can be assumed to increase linearly with the yield stress.

Diving depth cannot be increased arbitrarily by continuing to increase the hull thickness. The weight and displacement of the boat increase approximately linearly with the thickness for a given hull radius. If a submarine is to remain buoyant, its total weight cannot exceed the weight of the water displaced by the volume of the submarine. Assuming the weight of the pressure hull cannot exceed the fraction h_r of the total weight of the submarine (h_r is the hull-weight ratio⁴), and that the total weight of the submarine

⁴ The hull-weight ratio can differ substantially between submarine classes. For the Dolphin class, $h_r = 0.5$. In comparison, the hull-weight ratio for the Los Angeles class is $h_r = 0.18$ [9].

cannot exceed the weight of the water displaced by the submarine, the following constraint on the hull dimensions is obtained,

$$\frac{\text{weight of pressure hull}}{\text{weight of displaced water}} = \frac{2\pi R t \rho}{\pi R^2 \rho_{H_2O}} = \frac{2ts}{R} < h_r \quad (2)$$

where ρ and s are the density and specific density of the hull material, respectively. Combining Eq. (2) with Eq. (1) gives an approximation for the collapse depth of a submarine:

$$d_{collapse}(\text{meters}) = \frac{100 h_r \sigma_y}{s} \quad (3)$$

where σ_y is expressed in MPa. Equation (3) clearly shows the importance of the strength-to-weight ratio of the hull material (σ_y/s) in the structural design of the submarine.

The maximum operating depth, d_{max} , is obtained by multiplying the collapse depth, Eq. (1), by a safety factor, f_s :

$$d_{max}(\text{meters}) = f_s d_{collapse} = \frac{200 f_s t \sigma_y}{R} \quad (4)$$

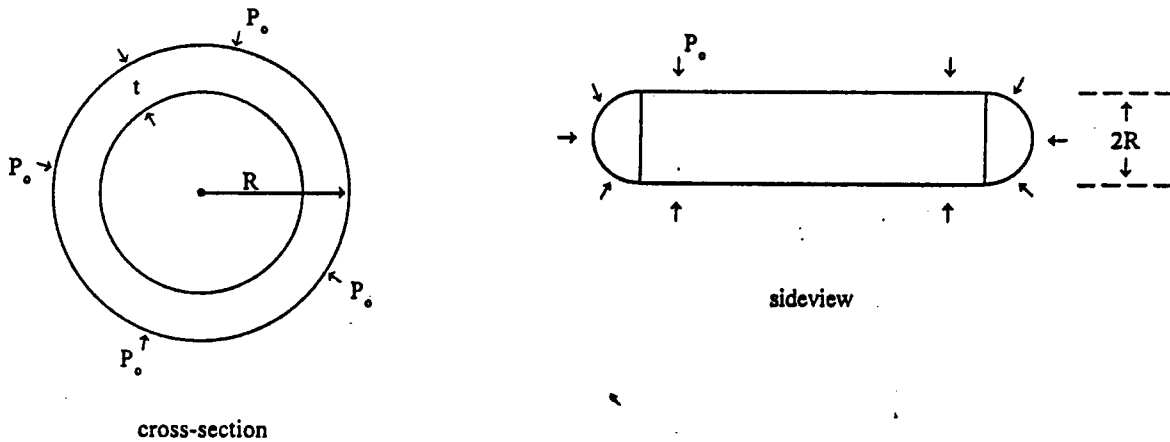


Figure 2. Cross-Sectional and Sideview of an Idealized, Thin-Walled, Submarine Pressure Hull
(The external pressure is P_o and the hull radius and thickness are R and t , respectively.)

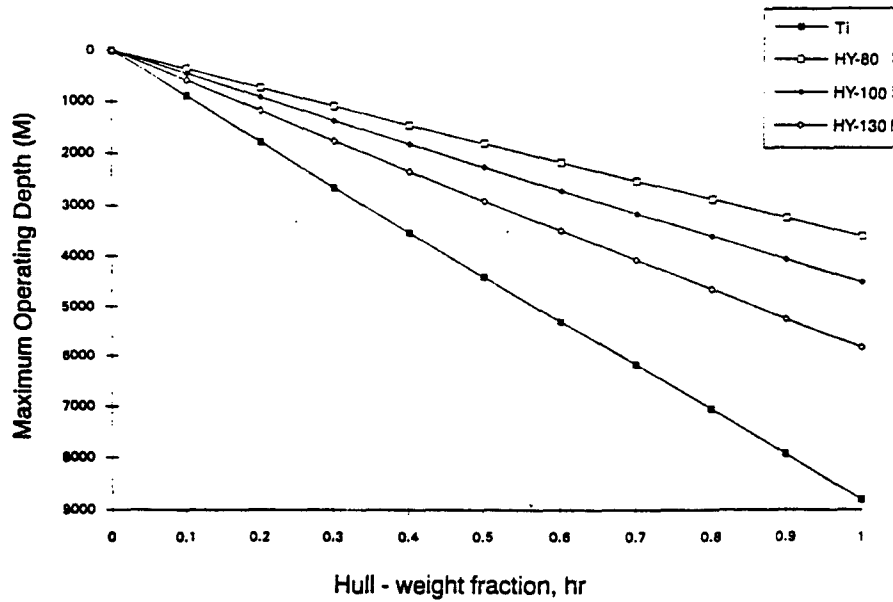


Figure 3. Maximum Operating Depth for Submarines With Pressure Hull of High-Strength Steel or Ti-100 Shown as Function of Hull Weight Fraction, h_r
(according to Eq. (3). A safety factor of 1/2 has been employed.)

B. DESIGN ISSUES: DEFLECTION AND BENDING

The dimensions of the pressure hull under load depend on the internal and external pressures; the change in dimension being a result of elastic displacement. Large external pressures relative to the internal pressure result in elastic deflection (compression) of the hull, beams, stiffeners, and other structural components. These elastic displacements cannot be arbitrarily large, and design criteria impose constraints on the deflections. The magnitude of the deflection for a given load is determined by the elastic modulus (E) and Poisson's ratio (ν).

In practice, pressure hull design is complicated by openings that must be present in the hull including the entrance, torpedo tubes, sonar window, and missile ports. Stresses are concentrated in such areas and must be carried by the surrounding material. If different structural materials are in contact with one another and their elastic moduli are sufficiently different, they will have different deflections for a given load and, where the materials come into contact, large stresses will develop. For this reason, critical inside structures often must be constructed of the hull material.

For illustrative purposes, the radial deflection of the hull (δR) can be considered proportional to the radial strain ϵ_{rr} of the hull. Using $\epsilon_{rr} = \sigma_{rr}/E$, one obtains

$$\delta R \propto P_o/E. \quad (5)$$

The smaller the elastic modulus, the greater the elastic deflection of the hull for a given external pressure.

C. COMPARISON OF TITANIUM ALLOYS AND HIGH-STRENGTH STEELS

Table 1 presents a few materials parameters relevant to the structural design of pressure hulls for steel alloys HY-80, HY-100, and HY-130⁵ and two titanium alloys [11,12]. Of particular importance is the low density of titanium relative to steel (4.5 g/cm³ vs. 7.8 g/cm³) and the low elastic modulus of titanium with respect to that of steel (110 GPa vs. 210 GPa).

Table 1. Materials Parameters of Steel and Titanium Alloys

Property	HY-80	HY-100	HY-130	Ti-5Al-2.5Sn	Ti-100
density (g/cm ³)	7.8	7.8	7.8	4.5	4.5
yield stress (MPa)	560	700	910	800	740
modulus (MPa)	2.1×10^5	2.1×10^5	2.1×10^5	1.1×10^5	1×10^5
elongation	20%			10%	8%
Poissons ratio	0.3	0.3	0.3	0.45	0.45?

For the same weight and safety factor in hull design, a titanium hull offers a significant increase in the maximum operating depth over a steel hull. Assuming identical hull diameters and weight, the ratio of the hull thickness for a submarine constructed of titanium (t_{Ti}) to that of steel (t_s) is proportional to the material densities, p :

$$\frac{t_{Ti}}{t_s} = \frac{\rho_s}{\rho_{Ti}}. \quad (6)$$

Using the yield criterion of Eq. (3) and assuming the same safety factor, the maximum operating depth of a titanium-hulled submarine (d_{max}^{Ti}) is related to that of a steel-hulled submarine (d_{max}^S) by

$$d_{max}^{Ti} = \left(\frac{\sigma_y^{Ti}}{\sigma_y^S} \right) \left(\frac{\rho_s}{\rho_{Ti}} \right) d_{max}^S. \quad (7)$$

⁵ The designation HY-80 refers to a high-strength steel with a yield stress of 80,000 psi.

Under these assumptions, a pressure hull constructed of Ti-100 would have a maximum operating depth of about 1.8 times that of a hull constructed of HY-100. The effect of different hull thicknesses is shown in Figure 4.

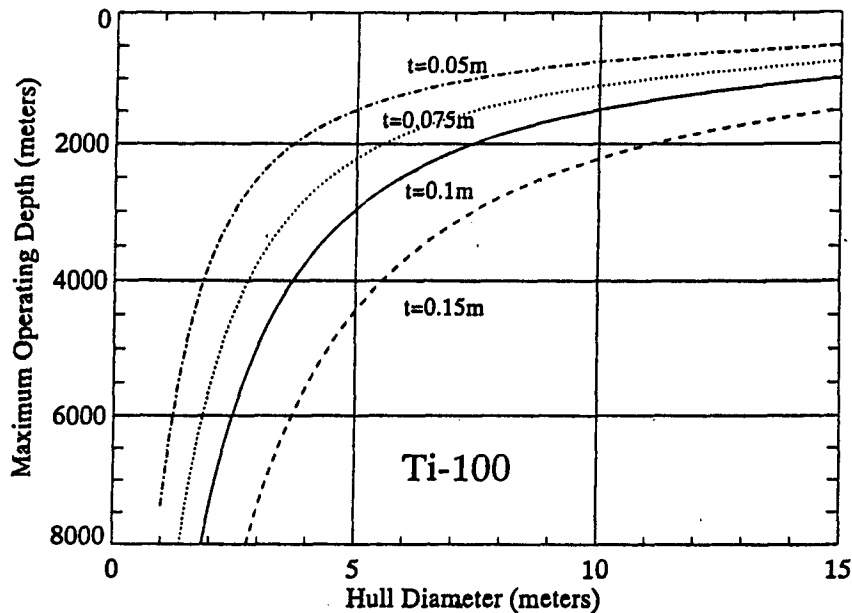


Figure 4. Maximum Operating Depth (d_{max}) for Submarine with Pressure Hull Constructed of Ti-100 Shown as Function of Pressure Hull Diameter for Several Different Hull Thicknesses (t)
(Using Safety Factor of 1/2)

The elastic deflection (displacement) of a titanium component under a given load with respect to that of a steel component possessing an identical geometry is given by the ratio of their elastic moduli, Eq. (5); a titanium component will undergo a displacement about twice as large as that of a steel component for a given load. The deflection at the onset of plastic yielding for a given geometry also depends upon the yield stress. At yielding, a titanium component would undergo a deflection of about 2.6 times that of a component constructed of HY-80. Comparable values for components constructed of HY-100 and HY-130 are 2 and 1.5.

Maximum tolerances for elastic displacements impose additional constraints on the hull design. For example, replacing HY-80 with HY-130 without structural changes to the hull would appear to increase the maximum operating depth by a factor of about 1.6 according to Eq. (1). However, since the elastic moduli of HY-80 and HY-130 are equal, the same elastic deflections would result when both materials are exposed to the same load. If the operating depth were to increase by a factor of 1.6, the hull deflection

under this new pressure would be a factor of 1.6 greater than the hull deflection at the operating depth of the HY-80 hull. If the design criteria do not permit a deflection greater than that encountered at the maximum operating depth of the HY-80 hulled boat, then the HY-130 hulled boat would still be constrained to the same maximum operating depth as the HY-80 hulled boat. For the same reason, simply replacing HY-80 with a titanium alloy without a concomitant change in the hull design would actually reduce the maximum diving depth.

From the viewpoint of structural design and on the basis of the above simple calculations, titanium alloys offer significant materials advantages over the high-strength steels, especially the commonly used HY-80 and even HY-100 currently used in the Sea Wolf submarine. Titanium's density is only 58 percent that of steel while its yield stress is 10-15 percent greater than HY-100. These strength-to-weight advantages are counteracted by the low elastic modulus of titanium as compared to steel. The lower elastic modulus requires thicker plates for the hull as well as thicker beams and stiffeners [11]. However, thicker sections for a particular weight actually allow for greater flexibility in structural design for operation at all diving depths [5] and also permits the production of more stream-lined shapes⁶ for noise reduction [13].

Russian submarine pressure hulls (Alfa class) are constructed of the Russian titanium alloy VTG. VTG is based on the composition Ti-6V-4Al which has good weldability [13].

It is noted that materials other than steel and titanium alloys have been considered for submarine pressure hulls. Aluminum alloys were abandoned owing to corrosion problems [9]. Glass reinforced plastics and carbon reinforced polymers could significantly increase diving depths over those obtainable using a titanium hull. However, submarines constructed from composites would need to be almost 30-cm thick and are extremely expensive. In addition, it is difficult to edge-bond the composite sections to metallic sections [14].

⁶ The lower density of titanium allows for thicker plates which permits easier structural design.

III. MATERIAL PROPERTIES

A. PROPERTIES

Titanium alloys have long been of interest to the Navy for use as a structural alloy, owing to their high strength-to-weight ratios and their superior resistance to corrosion in sea water. Beginning in 1975, an intensive investigation of titanium alloys for use as a structural material for combatant submarine pressure hulls was initiated [9]. Emphasis was placed on the alloy Ti-100 with a composition Ti-6Al-2Nb-1Ta-0.8Mo. This alloy had been used previously and successfully as a hull material for the deep submergence vehicle ALVIN.

In order to be used as a pressure hull material, Ti-100 must satisfy a great many material design requirements including, but not limited to, minimum values for the compressive and tensile yield strengths, ultimate tensile strength, fracture toughness, ductility, stress corrosion and sustained-load cracking, galvanic corrosion, erosion and cavitation resistance, creep behavior at room temperature, and low-cycle fatigue. Minimum design requirements are equally applicable to the plate material and welded areas joining the plates together.

A complete review of the state-of-the-art in titanium processing and joining is not possible here.⁷ It is important to recognize, however, that the material properties of Ti-100 (with appropriate processing) satisfy many of the minimum design requirements for submarine hulls. A few basic materials properties of Ti-100 are discussed below.

As a result of materials research and development efforts in the 1970s and early 1980s, some of the minimum design requirements have been achieved for Ti-100 [7,17]. Figure 5 indicates the yield strength⁸ of Ti-100 obtained in tension (a) and compression (b) for 2- and 4-inch thick plates using optimized heat-treating procedures [17]. Both

⁷ For a recent review of titanium alloys, see References 15 and 16.

⁸ The yield strength is the stress, in either tension or compression, at which plastic flow is initiated.

plates and corresponding welds surpass the minimum design goal of a yield stress of 100 ksi (700 MPa) in compression and tension.

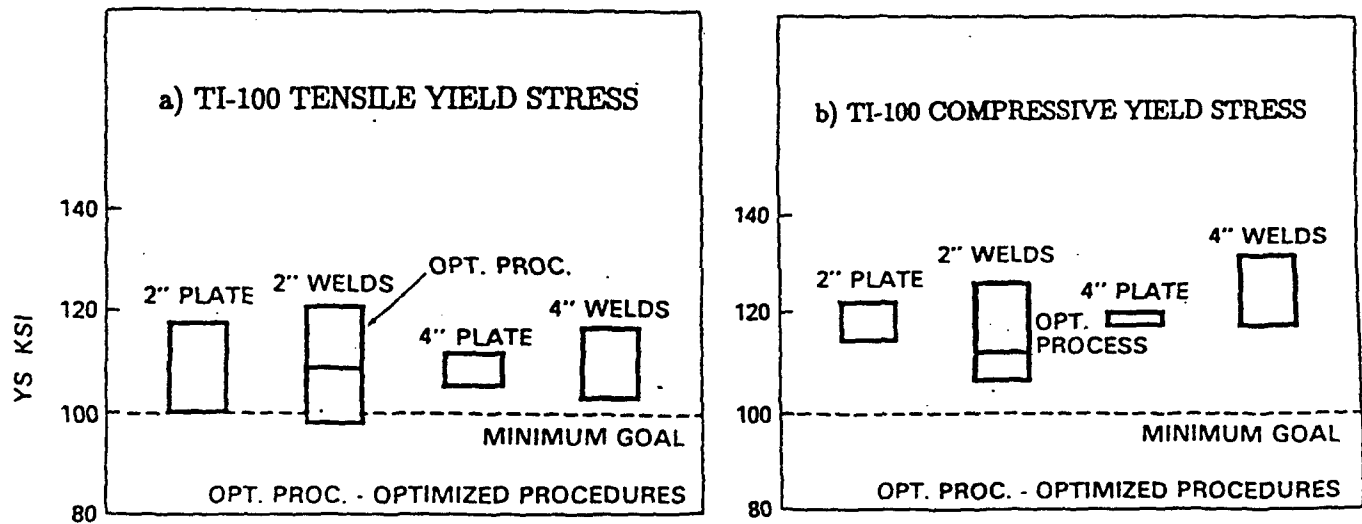


Figure 5. Yield Stress of Titanium Alloy Ti-100 in (a) Tension and (b) Compression Shown for 2- and 4-Inch Base Plates and Corresponding Welds

The ultimate tensile strength,⁹ (UTS) for both 2-inch and 4-inch base plates and corresponding welds are shown in Figure 6. The percent elongation to failure, a measure of the material's ductility, is shown for the same plate thicknesses and welds in Figure 7. In each case, the weld properties matched those of the base plate materials. Joining was accomplished using gas-metal-arc welding (GMA) on flat and horizontal pieces [17].

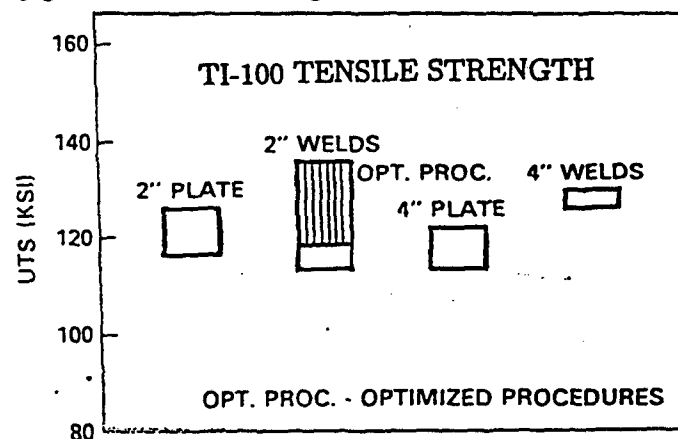


Figure 6. Ultimate Tensile Stress of Ti-100 Shown for 2- and 4-Inch Base Plates and Corresponding Welds

⁹ Ultimate tensile stress is the maximum tensile stress the material can sustain without fracture.

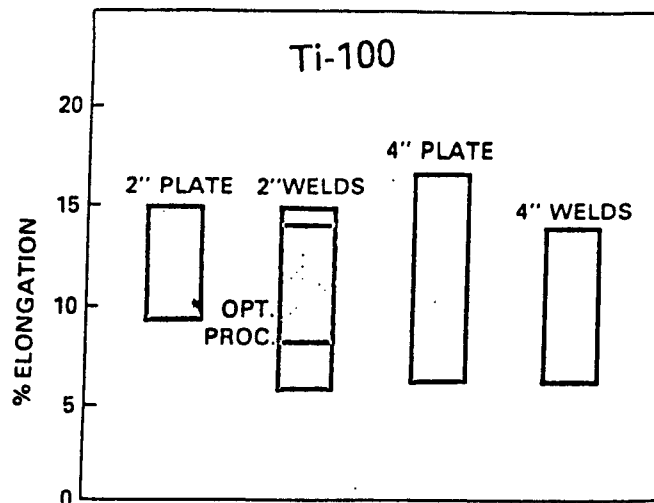


Figure 7. The Percent Elongation to Failure, a Measure of Ti-100's Ductility, Shown for 2- and 4-Inch Base Plates and Corresponding Welds

Titanium is unique in that it can experience a small, permanent deformation as a function of time under applied load (creep) at room temperature. Submarine design requires creep strains not to exceed 0.002 in/in over a period of 60 days when the applied load is equivalent to 90 percent of the yield stress in either compression or tension [17]. Reported creep strains for the bulk plate exceeded the maximum allowable creep strain by a factor of three while reported creep strains for the welds, 0.00024 in/in, were well within design criteria. In contrast, however, the fracture toughness of the welds does not meet the minimum design goal as established by 1-inch, drop tear fracture toughness tests illustrated in Figure 8.

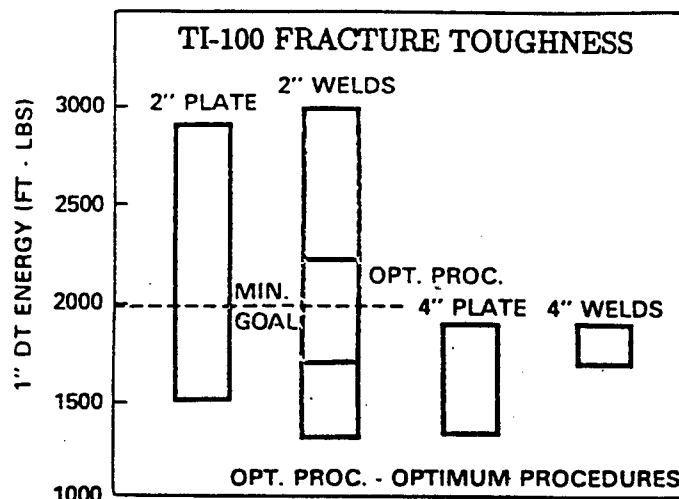


Figure 8. Fracture Toughness of Ti-100 Shown for 2- and 4-Inch Base Plates and Corresponding Welds

(The welded material does not display sufficient toughness for use as hull material.)

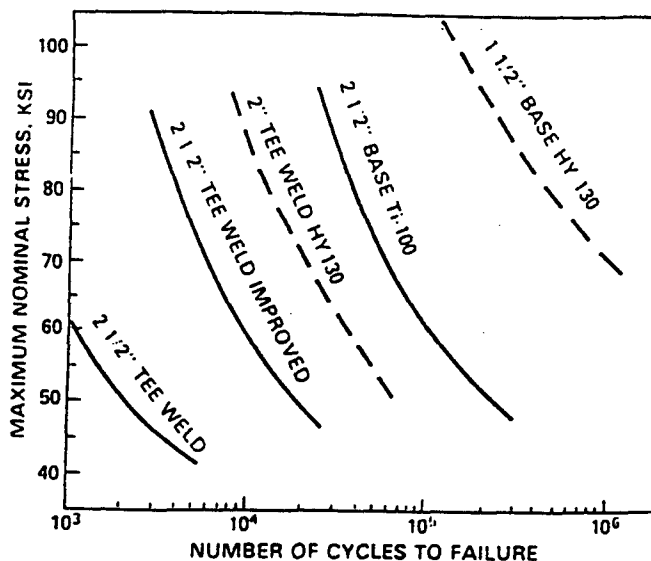


Figure 9. Low-Cycle Fatigue Behavior of Ti-100 Base Plate and Welds Compared to High-Strength Steel HY-130

B. FABRICATION

Titanium is a relatively abundant alloy found in rutile (TiO_2) and ilmenite (FeTiO_3). Rutile has a titanium dioxide content of about 94 percent [18]. Titanium is produced from these ores using the well-established Kroll process based on the reduction of TiCl_4 by Mg: $2\text{Mg} + \text{TiCl}_4 \rightarrow 2\text{MgCl}_2 + \text{Ti}$. The TiCl_4 is prepared by the chlorination of a mixture of carbon and the titanium-containing ore [19]. Most ore is imported from South Africa and Australia of which 90 percent is used as pigment for paint [20]. Melt processing of titanium and titanium alloys must contend with the extreme reactivity and volatility of titanium. Usually sophisticated vacuum processing is required.

The production of titanium-hulled submarines would require the development of extensive fabrication support services. The handling of titanium alloys would probably employ standard methods used for nonmagnetic materials including slings, hooks, and vacuum heads. Some machining operations would also be able to make use of standard shop practices. It is estimated that bending plates and components to hull contours could be accomplished using cold forming and roll bending techniques at costs similar to steel forming. However, titanium is known to stick to other tool materials [21], and other procedures, such as welding, would require an extensive investment in infrastructure and training of new technicians. A significant learning curve would be anticipated [7].

Several significant problems remain concerning the weldability of titanium alloys. The material performance of the weldment (contiguous base plate material, heat affected zone and actual weld material) is adequate for yield strength and ductility, but the fracture toughness of the weld requires further work. In addition, crack propagation in the weldment is extremely sensitive to weld contamination (primarily oxygen): Existing work has shown that usually no cracks are found in uncontaminated welds while extensive cracking is found in contaminated welds [21].

Problems with oxygen pickup (contamination) are usually taken to mean that the titanium alloys used in submarine hull construction would need to be welded in inert environments requiring very large and expensive buildings. Indeed, apocryphal reports exist of the FSU constructing entire buildings filled with inert gases for the expressed purpose of welding titanium hull plates. (The FSU probably employed tungsten electrode inert gas welding [18].) However, other possible solutions have been suggested [7]. Shielding devices that would envelop the weld material and arc in an inert gas could be developed to prevent contamination of the weld. The hull plates could be welded together to form cylindrical rings. Rings are then joined together to form the hull. Each ring would be constructed in a container with a controlled environment. In this scenario, only the weld area joining the rings would need to be contained in an inert atmosphere. Internal welding could be performed subsequently in an inert atmosphere once the hull had been sealed. Such procedures would reduce production costs.

It is also felt in some quarters that the demands placed on weld performance are over-matched [7, 21], i.e., the demands placed on the weld are unnecessarily high. The weld need not be as tough as the base metal so long as the weld is sufficiently ductile.

The FSU was apparently successful in welding titanium plates in the complicated geometry found in submarine pressure hulls. The material properties of welds reported above were based on the joining of flat, horizontal pieces. Successful welding of submarine hull plates would require developing the technology to join curved pieces in vertical positions. The use of titanium alloys as a hull material in special purpose submarines would require developing technologies for the joining of plates up to 20-cm thick.

Efforts were made in the early 1980s to introduce titanium welding technology for thicker plates to industry. This effort was not successful owing to union and labor

problems and the lack of commitment to actually support and purchase the finished product [7].

IV. COST ANALYSIS

The purpose of a cost/benefit analysis is to express various performance characteristics and attributes of the submarine in terms of a single factor: the monetary cost of the submarine over its anticipated lifetime. This process necessitates knowledge of the cost of obtaining such performance characteristics as desired operating depth and speed, construction costs, and lifetime operating costs. Construction costs would include development, certification, fabrication, and test costs as well as the cost of labor and materials. Estimating life-cycle costs of a new platform such as a titanium-hulled submarine is extremely difficult owing to the lack of operating experience.

Table 2 gives a breakdown of platform costs for an SSN-21 Sea Wolf attack submarine [22]. The total cost was estimated to be \$1.8 billion in 1991; 19 percent of which was attributable to the payload (\$344 million) and 81 percent to the platform (\$1.47 billion). Life-cycle costs, of course, would be much higher.

Table 2. Platform Costs of an SSN-21 Sea Wolf Submarine

System	Cost (\$Millions)	Percent Platform Cost
Hull Structure	132	9
Propulsion Plant	440	30
Electric Plant	44	3
Command & Surveillance	59	4
Auxiliary Systems	191	13
Outfit & Furnish	88	6
Integration/Engineering	294	20
Assembly & Support Services	161	11
Total	1,467	100

The hull structure comprises about 8 percent of the total cost of the Sea Wolf submarine. Of this, about 40-50 percent can be attributed to labor costs. Thus doubling the cost of the pressure hull material and labor would add about 8 percent to the total construction cost of the submarine.

Table 3 provides the representative life-cycle costs for a typical Navy ship other than a submarine [23]. Ship retirement is not included in the analysis, which could be as high as a third of the acquisition cost [24]. Maintenance costs would be expected to be decreased for a titanium ship owing to much improved corrosion resistance.

**Table 3. Representative Life-Cycle Costs
for a Typical Navy Ship (1991)**

Cost Component	Percent Total Cost
Acquisition	23
Design	2
Fuel	4
Personnel	37
Maintenance	21
Modernization	13

The original construction and subsequent maintenance of a titanium-hulled submarine would entail significant costs. The infrastructure for processing, forming, and welding large plates of titanium does not currently exist, although it has been argued (probably erroneously) that existing steel facilities could be converted simply to titanium use [20]. There is the additional problem of training sufficient numbers of technical people in the handling of titanium and having such new procedures adopted in practice.

Comparing the life-cycle costs of titanium- and steel-hulled submarines is not possible without understanding the operational requirements of the submarine [9]. This is because the operational requirements determine the performance characteristics of the submarine: In certain scenarios, the cost-benefit of titanium-hulled submarines could exceed that of a steel-hulled submarine; the opposite could be true for a different set of operations.

A cost-benefit model was developed in 1978 at IDA for nuclear attack submarines of the Sturgeon class [9]. The purpose was to examine the relative costs of using titanium and steel for combat submarine pressure hulls. Neither life-cycle costs nor the cost of submarine disposal¹⁰ were examined. The IDA study concluded:

¹⁰ Submarine disposal has been estimated to run 30-40 percent of the construction cost [24].

1. The use of HY-130 as a pressure hull material would offer significant cost benefits over HY-80 for operating depths close to 3,000 feet.
2. Titanium alloys offer few or no advantages over HY-130 unless operating depths approaching or exceeding 3,000 feet are required.
3. A *conventionally* designed titanium-hulled submarine designed for operating depths of 4,000 feet or more would be both large and very expensive.

The conclusions drawn by the 1978 IDA study are probably still qualitatively correct, although, at the time of the study, it was believed that the major problems associated with the use of HY-130 (stress-corrosion cracking of weldments) were either essentially solved or would be in very short order.

V. SUMMARY AND RECOMMENDATIONS

There are several generic reasons to consider titanium alloys for use as a pressure hull material in terms of both its enhanced static and dynamic loading capabilities. Whether it is economically justified to do so depends critically upon the intended use of the submarine.

1. Titanium possesses a favorable strength-to-weight ratio with respect to all high-strength steels. The low elastic modulus of the titanium notwithstanding, this characteristic permits greater flexibility in the structural design of the submarine. For example, a greater payload weight can be carried or the top operating speed increased for a given operating depth. Submarine dimensions can be enlarged for a given hull weight.
2. Titanium pressure hulls allow for increased operating depths. A greater diving depth offers significant advantages to nuclear attack submarines [9] including (a) increased tactical flexibility for attack and subsequent evasion, (b) increased operating space, and (c) significantly improved effectiveness in using the ocean acoustic environment in order to detect, classify, and localize sources. This is especially true for submarines capable of operating at depths of 1,300 meters or more [9]. In addition, the increased diving depths could permit the development of special-purpose submarines currently available to Russia [8, 25].
3. Titanium hulls can enhance a submarine's stealth capabilities with respect to high-strength steels. The nonmagnetic hull reduces the chances of nonacoustic detection by MAD. The greater flexibility in hull design afforded by titanium provides better optimization of acoustic emissions and a more streamlined shape. The greater operating depth is achieved without sacrificing weapons payload or speed.
4. Titanium hulls might offer enhanced resistance to dynamic shock. However, any improved resistance depends upon the static design and operating depth, as well as the material parameters of the hull material.
5. Life-cycle costs for titanium-hulled submarines would probably be lower than those of steel-hulled submarines for most deep-diving (greater than 900-1,000 meters) applications.

Titanium alloys have been developed and used as a structural material for submarine pressure hulls in Russia. Titanium alloys are also reportedly being exploited for use in deep-diving and special-purpose submarines in Russia, and the technology necessary for the fabrication and joining of these alloys is continually being developed in both Japan and Russia. Offers to purchase this technology from the Russians have reportedly been declined [21]. It is recommended that modest efforts be undertaken in the United States to increase the technology base associated with the joining of titanium.

ACKNOWLEDGEMENTS

I am grateful to the following people for their support in preparing this report: H.I. Aaronson, G. Bell, R.P. Gangloff, W. Hurley, N. Licato, E. Metzbauer, D. Meyn, D. Michels, G. Sorkin, and Adm. Train.

Annex A
DETERMINATION OF COLLAPSE DEPTH

Annex A

DETERMINATION OF COLLAPSE DEPTH

In this annex, a simple relationship between maximum operating depth, submarine beam dimension, pressure hull thickness, and the yield stress of the pressure hull material is obtained. The submarine is considered to be a thin-walled, hollow cylinder with the dimensions shown in Figure 2. Assuming the external pressure (P_0) is much greater than the internal pressure and that the hull thickness (t) is much smaller than the hull radius (R), the equations of elasticity can be solved for the stress and strain fields of the pressure hull in cylindrical coordinates:

$$\sigma_{rr} = -\frac{P_o R[1 - (a/r)^2]}{2t} \quad (8)$$

$$\sigma_{\theta\theta} = -\frac{P_o R[1 + (a/r)^2]}{2t} \quad (9)$$

$$\sigma_{zz} = \frac{-P_o \pi R^2}{2\pi R t} = \frac{-P_o R}{2t}. \quad (10)$$

The collapse depth is assumed to be determined by that external pressure which causes the pressure hull to begin to deform plastically. The onset of plastic deformation is established by the von Mises yielding criterion [10] as

$$\left(\frac{1}{2}\sigma'_{ij}\sigma'_{ij}\right)^{1/2} = \sigma_y \quad (11)$$

where the deviatoric stresses are defined by

$$\sigma'_{ij} = \sigma_{ij} - \frac{1}{3}\sigma_{kk}. \quad (12)$$

Substituting Eqs. (8)-(10) into Eq. (11) gives for the collapse pressure, $P_{collapse}$,

$$P_{collapse} = \frac{2\sigma_y t}{R}. \quad (13)$$

The collapse depth can be expressed approximately as

$$d_{collapse}(meters) = \frac{200t\sigma_y}{R} \quad (14)$$

where t and R are given in meters and σ_y is expressed in MPa.

REFERENCES

1. M. Vego, "Soviet Alfa-Class SSN," *NAVY International*, 139 (March 1984).
2. "U.S. Has Alfa-type Titanium Sample," *Jane's Defense Weekly* (14 January 1984).
3. B. Starr, "Russia is Building Next Generation SSN," *Jane's Defense Weekly* (9 October 1993).
4. "HY-130 Steel, Better Design Makes Soviet Subs More Dangerous," *Navy News and Undersea Technology* (26 February 1990).
5. C. Dawson, "Advances in Submarine Technology," *International Defense Review*, 1211 (September 1989).
6. "Titanium Subs Spur MAD Sensitivity Improvements," *Defense Electronics*, 35 (December 1982).
7. Dr. George Sorkin, private communication (June 1995).
8. R. Holzer, "Russians Invest in Special-Ops Subs," *Defense News*, 3 (15-21 May 1995).
9. *An Assessment of the Materials/Structures Technology Programs for Submarine Pressure Hulls*, IDA Paper P-1188, Institute for Defense Analysis, March 1976.
10. Y. C. Fung, *A First Course in Continuum Mechanics*, (Prentice-Hall, Inc., 1969) Chapter 5.
11. K. M. Heggstad, "Submarine Hulls of Titanium," *Maritime Defense*, 105 (April 1983).
12. H. L. Sussman, "Literature Search Evaluation of Titanium for Submarine Seawater Piping," David Taylor Research Center DTRC-PAS-89/35 (December 1989).
13. M. Suisman and G. R. Daly, "Titanium Big in Soviet Subs," *American Metal Market* (5 June 1985).
14. "Navy Considers Composite Submarine for 21st Century Fleet," *Navy News Undersea Technology* (May 20, 1991).
15. E. W. Collings, *The Physical Metallurgy of Titanium Alloys* (ASM, Metals Park, OH 1984).

16. *Materials Properties Handbook: Titanium Alloys*, eds. R. Boyer, E. W. Collings and G. Welsch (ASM International, 1993).
17. *Technical Review of Ti-100 Titanium Materials and Structures Technology for Future Submarine Construction*, Naval Sea Systems Command, 1977.
18. W. Guan "Titanium Alloy Submarines: Advantages and Problems," *CONMILIT*, pp. 59-61, June 1985.
19. R. P. Gangloff, *Physical Metallurgy of Structural Alloys*, University of Virginia, 1995.
20. "Titanium Industry Recovering from Downturn," *Navy News and Undersea Technology* (12 September 1988).
21. Drs. E. Metzbauer, D. Meyn and D. Michel, Naval Research Laboratory, private communication, 1995.
22. J. Johnston, D. O'Colman and C. Mathas, *Where the SSN Dollars Go: An Affordability Focus*, Cost Estimation and Analysis Division (SEAO17) Naval Sea Systems Command (April 11, 1991).
23. J. C. Ryan and D. P. Jons, "Achieving Technical and Management Excellence," Annual Technical Symposium (28), U. S. Naval Sea Systems Command, 1991.
24. B. Hinderliter, private communication, 1995.
25. "Titanium Alloy Pressure Hull," *Maritime Defense*, 303 (August 1984).

E. LAND-BASED ACOUSTIC SENSOR

**Christopher S. Kochanek
Department of Astronomy
Harvard University
Cambridge, MA**

ABSTRACT

We discuss issues related to construction, data processing and targeting using sparse acoustic sensor arrays in the context of locating and neutralizing snipers.

1 INTRODUCTION

Acoustic threat location is the oldest of passive sensors, since it is one of the primary functions of the human ear. The capabilities of the human ear to locate and recognize threats is astonishingly good. The ear can determine the arrival time of a sharp acoustic pulse to within about 10 μ sec (microseconds = 10^{-6} seconds) in laboratory conditions (see Woods 1941), and trivially classifies the source (rifle, canon, helicopter, plane, etc.). Microphones and computers can match the timing accuracy and improve on the directional accuracy of the human ear, but they cannot match the brain's ability to accurately classify sources.

Acoustic location on the land battlefield dates to the late 19th century and the problem of locating enemy artillery and directing counter-battery fire.¹ The systems used in WWI, consisting of microphones and plotting tables, were reliable and (reportedly) surprisingly accurate (80 m at 12 kilometers!). Systems consisting of directional horns (i.e. like listening at the mouthpiece of a trumpet) were used to locate and track aircraft, again with surprising accuracy. When WWII started, the artillery tracking systems from WWI were (literally) dusted off and used again with minor improvements. It is estimated that 75% of confirmed weapon targets were acoustically located. When the Korean War started, the modified units from WWII were dusted off and used again. When the Vietnam War started, the modified units from the Korean War were dusted off and used again (although they were beginning to fall apart). Several projects were started in the late 1960s to update the systems, but all the projects died in the early 1970s. The accuracy of the systems in Korea and Vietnam was substantially worse than the performance of the earlier systems, in part due to terrain and equipment age.²

AAS WGAS macros v2.0

¹This historical review is a synopsis from *Acoustic Weapons Location 197?*.

²The earlier reports may also have been exaggerations. Aerial reconnaissance in Korea suggested that the accuracies were closer to 180 meters than 80 meters.

The competing technology of radar appeared in the Korean War. Counter-battery radar scans the horizon for shells in flight. When it finds a shell, it computes the trajectory and tracks it back to the firer's location. In both Korea and Vietnam, neither the acoustic nor the radar systems were very successful due to the terrain limitations, but statistics suggest that in both conflicts the ancient acoustic sensors frequently outperformed the "modern" radar technology: the probability of locating a battery was 1-2% for the acoustic sensors versus 0.3-0.6% for the radars. Currently the US military uses only counter-battery radars (acoustic sensors were retired in the early 1980s (Hewish & Pongelley 1990)), although acoustic artillery location systems are fielded in Sweden, Russia, and possibly England (Janes 19??). The other WWI application, locating aircraft, has also reappeared, particularly in the application of locating helicopters. Such systems are included in anti-helicopter mines, and may be used as a warning device for the FAADS (Forward Area Air Defense System) anti-aircraft system. Acoustic sensors are also used in some "smart" anti-tank mines, and in the air-dispersed Brilliant Anti-Armor Submunition (BAT) (see Battlefield Acoustic Sensors 1992, Hewish & Pongelley 1990).

This discussion focuses on the issues involved in deploying, calibrating, and targeting using acoustic arrays. It will not discuss the issue of identifying targets. The original context of the paper was the problem of locating snipers, but the techniques are much more general. The second section outlines the physical parameters governing the deployment of an acoustic array, primarily the size of the system and the accuracy required for the microphone locations. The third section discusses methods for internally "self-calibrating" or "self-cohering" small arrays (sizes less than 100m). The fourth section discusses calibration using external sources, and the fifth discusses using arrays to provide target corrections. The last section summarizes the possibilities.

2 PHYSICAL CONSTRAINTS ON ACOUSTIC ARRAYS

In this section we discuss the physical requirements for acoustic arrays given a desired level of measurement accuracy for the bearing and range to a target. Acoustic arrays find targets by triangulating their positions using the relative arrival times of acoustic emissions from a source across the array. The fundamental performance of an array depends on the size of the array (distances between microphones), the accuracy of the location of the microphones, the accuracy with which the arrival times (or phases) of signals can be measured, and the accuracy of the model for the propagation of signals (sound speed, wind speed, refraction, reflection ...). The first step is simply to estimate how big an array must be for a given accuracy.

2.1 *The Two-Element Array - Bearing*

The simplest case we can analyze consists of two microphones separated by the baseline distance b . Let microphone 1 lie at $\mathbf{x}_1 = (0, -b/2, 0)$, and microphone 2 lie at $\mathbf{x}_2 = (0, b/2, 0)$. For simplicity, consider the problem in the x - y plane with a spherical wave produced at range R and bearing angle θ (see Figure 1). The time delay, or difference in the arrival

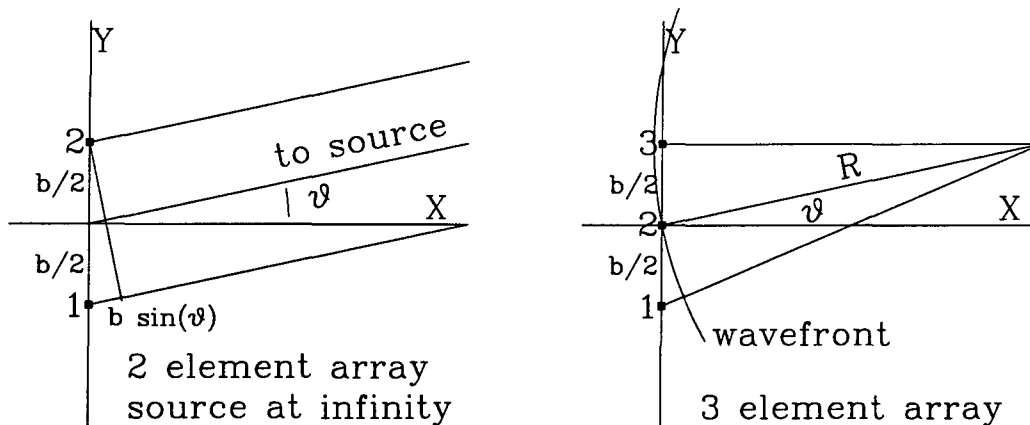


FIG. 1.—Schematic diagrams of 2 (left) and 3 (right) element arrays. In the 2 element array the source is at infinity ($b/R \gg 1$) and the time delay is $\Delta t_{12} = t_1 - t_2 = b \sin \theta / c_s$. Wavefronts of the signal are straight lines perpendicular to the direction to the source. In the 3 element array the source is at a finite (close!) range, and the wavefront is the labeled circle. The range information comes from determining the curvature of the wavefront, and the bearing information comes from the tangent to the wavefront.

times, between the two microphones is

$$t_1 - t_2 = \Delta t_{12} = \frac{b \sin \theta}{c_s} \quad (1)$$

for a source at infinite range ($b/R \ll 1$) where $c_s = 343.3 \text{ m s}^{-1}$ (at $T = 0^\circ \text{ C}$) is the sound speed. A two element array measures only the bearing of the source,

$$\sin \theta = \frac{c_s \Delta t_{12}}{b} \quad (2)$$

and the characteristic time difference is the sound crossing time over the baseline, $t_{\text{cross}} = b/c_s = 2.9(b/\text{m}) \text{ msec}$ (msec=milliseconds). If we simply use propagation of errors to estimate the uncertainties in the bearing,

$$\sigma_\theta^2 = \tan^2 \theta \left[\left(\frac{\sigma_c}{c_s} \right)^2 + \left(\frac{\sigma_b}{b} \right)^2 \right] + \sec^2 \theta \left(\frac{\sigma_t}{t_{\text{cross}}} \right)^2 \quad (3)$$

where σ_c , σ_b and σ_t are the errors in the sound speed, the baseline length, and the time measurement respectively. The expression does not include errors in the baseline orientation. The formal error diverges for source-array separations parallel to the baseline ($\cos \theta = 0$), so a mixture of baseline orientations is needed to be sensitive in all directions.

2.2 Three-Element Arrays: Range Estimation

Target range information requires triangulating from two element arrays. The simplest three element array adds a third microphone to the two element array discussed above (see

Figure 1). The bearing to the source is

$$\sin \theta = \frac{c_s \Delta t_{13}}{b} \quad (4)$$

(as before) and the range³ to the source is

$$R = \frac{b^2 \cos^2 \theta}{4c_s(\Delta t_{12} - \Delta t_{23})} \quad (5)$$

provided $|\Delta t_{12} - \Delta t_{13}| \gg \sigma_t$ where σ_t is the measurement error and $\cos \theta \neq 0$. The standard error in the range is

$$\left(\frac{\sigma_R}{R}\right)^2 = 4\left(\frac{\sigma_b}{b}\right)^2 + \left(\frac{\sigma_c}{c_s}\right)^2 + \frac{4}{\cos^2 \theta} \left(\frac{R \sigma_t}{b t_c}\right)^2. \quad (6)$$

The order of magnitude contributions of the baseline and sound speed errors are the same as in the bearing (see eqn. (3)), but the contribution from the timing errors is $2R/b$ larger.

2.3 Permissible Errors

A good target accuracy for the array is a standard error in the position of the source of 10 meters, the characteristic error of the GPS system. A $\sigma_R = 10$ m error at a range of $R = 1$ km corresponds to a bearing accuracy of $\sigma_\theta = 10$ mrad (milliradians), and the same accuracy at 10 km corresponds to $\sigma_\theta = 1$ mrad. There are two types of error terms in eqns. (3) and (6); terms that are independent of the baseline length, and terms that depend on the baseline length.

Errors in the sound speed, or any other error affecting the effective propagation rate of the acoustic signal, do not depend on the baseline length. The sound speed is only the most obvious of these effects. Errors in the wind speed, refraction, reflection, and "turbulent" broadening of the signal width can all act like errors in the effective sound speed. If we need $\sigma_\theta = \sigma_{\theta 1}$ mrad, then $\sigma_c \lesssim 0.35\sigma_{\theta 1}$ m s⁻¹ (1.2 km hr⁻¹). The sound speed varies with temperature by about 0.6 m s⁻¹ K⁻¹, so the mean temperature must be known to $\sigma_T \lesssim 0.6\sigma_{\theta 1}$ K. Winds advect the sound wave with them, so the effective sound speed between two points is the true sound speed plus the component of the wind velocity along the separation vector. An error in the wind speed will act like an error in the sound speed, so we must know the mean wind speed to $\sigma_v \lesssim 1.2\sigma_{\theta 1}$ km hr⁻¹.

If we consider impulsive events (shock waves and muzzle blasts), then the accuracy with which the arrival time can be determined depends on the characteristic frequency of the event. Sonic pulses from explosions have characteristic frequencies of $\nu_c \simeq 200Y_{Kg}^{-1/3}$ Hz, where Y_{Kg} is the size of the explosive charge in kilograms of TNT. The rough accuracy with which we can measure the arrival times is

$$\sigma_t \sim \Delta T \left(\frac{S}{N}\right)^{-1} \left(\frac{\delta t}{\Delta T}\right)^{1/2} \quad (7)$$

³The inverse range R^{-1} is a better defined statistical quantity when the errors are large.

where $\Delta T \simeq \nu_c^{-1} = 5Y_{Kg}^{1/3}$ msec is the width of the signal ($\Delta T \simeq 1$ msec for an automatic rifle), S/N is the peak signal-to-noise ratio, and δt is the sampling rate. Typical values are $S/N \sim 50$, and $\delta t \sim 0.1$ msec. Timing accuracies of 10 to 100 μ sec (microseconds) are achievable – comparable to the performance of the human ear. Signals from short range weapons (e.g. automatic rifles) will have higher timing accuracies than signals from long range weapons (canon) not only because they intrinsically produce higher frequencies, but also because of the differences in range. High frequencies (particularly above 10 kHz) are absorbed by the atmosphere and filtered out of the signal from a canon at 10 km, and turbulence in the atmosphere broadens the pulse.

The timing accuracy is set by the physical properties of the acoustic signal, so to achieve a given range accuracy we must adjust the size of the array. For a timing accuracy of $\sigma_t = 100\sigma_{t100}$ μ sec, a range $R = R_1$ km, and desired range error of $\sigma_R = 10\sigma_{R10}$ m, the minimum baseline is

$$b \gtrsim b_{min} = R \left[\frac{2c_s\sigma_t}{\sigma_r} \right]^{1/2} = 83R_1\sigma_{t100}^{1/2}\sigma_R^{-1/2} \text{ m} \quad (8)$$

and the uncertainties in the microphone positions must be smaller than

$$\delta b < b \frac{\sigma_R}{2R} = \frac{b}{b_{min}} \left(\frac{c_s\sigma_t\sigma_R}{2} \right)^{1/2} = 0.4 \frac{b}{b_{min}} \sigma_{t100}^{1/2} \sigma_{R10}^{1/2} \text{ m.} \quad (9)$$

A counter-battery array (10 m accuracy at 10 km) must be spread out over a kilometer or more, and the microphone positions should be known to approximately 4 meters. In fact, the Swedish SORAS 6 system has nine microphones spread in an area 8 km wide and 1 to 2 km deep (Jane's 19??) and the individual microphones must be at least 300 m apart. For ranges of up to 25-30 km, the system is advertised to be accurate to 0.5-2% of the distance. For artillery location it is difficult to obtain high accuracy timing because only the low frequencies can propagate such large distances. The SORAS 6 system only uses frequencies from 2 to 150 Hz, so its timing accuracy is probably of order $\sigma_t = 1$ msec. An anti-sniper array (10 m at 1 km) must be about 100 m in size and the microphone positions must be known to slightly better than 0.5 m. If we increase the precision to 1 meter at 1 kilometer, then the array must be at least 250 meters in size with position accuracies of 0.1 m.

Because counter-battery arrays must be spread over a wide area, the microphone positions have to be determined by surveying, external calibration (see §4), or GPS. Since the positions of the microphones need only be known to a few meters, any newly designed system (although not the deployed ones) should simply use time averaged GPS positions or differential GPS systems to determine their positions. An anti-sniper array is more difficult to deploy. It must still be spread over a wide area, but the microphone positioning requirements are near the limits for GPS based positions, particularly if there is any risk of jamming in the battlefield. Surveying the positions is labor intensive, slow, hard (because the fractional accuracy is 10^{-3}), and dangerous near the front lines.

In fact, it appears that none of the ARPA anti-sniper programs is seriously trying to determine accurate ranges. The acoustic systems are based on small fixed arrays of

microphones, and this limits the baselines to a few meters. Small baselines are hard to correctly deploy because the position uncertainties drop to centimeters (or even millimeters), and only fixed, rigid mounts can robustly provide this precision. The question we address for the remainder of this discussion is the deployment, calibration, and use of arrays intermediate in size between the huge counter-battery systems, and the small, rigidly mounted, bearing determination systems.

3 SELF-CALIBRATION OF ACOUSTIC ARRAYS

We are faced with the fundamental problem of safely erecting an acoustic array within one kilometer of the front line, using microphones separated by tens to hundreds of meters, each located to within ten centimeters, quickly, accurately, and safely. In this section we develop a method for determining the locations of all of the microphones in an array with no surveying of their locations and without using rigid connecting rods.

A real array must contain more elements than the simple two and three element arrays we analyzed in §2. A linear array is most sensitive to sources in the direction perpendicular to the array, so a full array must have enough microphones to be sensitive with equal accuracy in all directions. For example, the standard rigid designs consist of small arrays of four microphones on a rigid mount (separations of a meter). A dispersed array should have more elements to increase its sensitivity and to provide redundancy.

Suppose we have N microphones, and we want to determine the $3(N-1)$ components of the relative positions. The number of baselines in the array is $N(N-1)/2$. The number of unknown positions grows only linearly with N , but the number of measurable quantities grows as the square. Measuring the relative delays cannot determine the absolute position of the array, and the relative positions are known only up to rotations or reflections of the array. We measure $N(N-1)/2$ time delays, and (because of the degeneracies) we can measure $3(N-1)-2$ relative coordinates of the microphones (the -2 is for the rotational degeneracy), so we need at least $N = 5$ microphones to determine the positions. We must also include the effects of the wind. If v is the wind velocity along the 1-2 baseline, then the sound travel times between the two microphones depends on the direction, with $\Delta t_{12}(c_s + v) = b$ and $\Delta t_{21}(c_s - v) = b$. Calibrating the array without accounting for the wind leads to fractional errors in the baselines of order the wind Mach number $v/c_s = 0.008(v/10\text{km hr}^{-1})$ rather than of order the fractional timing accuracy $\sigma_t/t_{\text{cross}} = 0.0035(\sigma_t/10\mu\text{sec})(m/b)$. Simply by comparing the propagation speed in the two directions, the array can determine the velocity,

$$v = c_s \frac{\Delta t_{12} - \Delta t_{21}}{\Delta t_{12} + \Delta t_{21}} = c_s^2 \frac{\Delta t_{12} - \Delta t_{21}}{2b}. \quad (10)$$

The measurement accuracy for the velocity is $\sigma_v = 14(\sigma_t/100\mu\text{sec})(m/b) \text{ km hr}^{-1}$.

If we add to each microphone assembly the ability to transmit a sharp sound pulse or a quasi-periodic wave audible across the array, then we can measure the relative delays between all pairs Δt_{ij} . Given the relative times, the relative positions of the microphones can be found to accuracy $(\sigma_b/b)^2 \sim (\sigma_t/t_{\text{cross}})^2 + (\sigma_c/c_s)^2$. For a given time accuracy, the baseline

accuracy will be $\sigma_b = c_s \sigma_t = 3.4(\sigma_t/100\mu\text{sec})$ cm, and this is more accuracy than required for our model requirements in an anti-sniper array (10 cm to 1 m). The optimal calibration signal is a high frequency, pseudo-randomly coded sound wave because it is both hard for the enemy to detect and very precise. A pseudo-random code is easy to detect by receivers that know the coding pattern, but difficult to detect by an enemy searching for a signal.⁴ High frequencies allow precise measurements (for frequency ν the timing accuracy is $\sigma_t \sim \nu^{-1}$). Moreover, the atmosphere absorbs high frequency sound waves. The classical absorption in dry air for a frequency of f_{kHz} measured in kHz is $1.7 \times 10^{-4} f_{\text{kHz}}^{-2}$ dB m⁻¹, so a 10 kHz signal has a classical loss of 17 dB over one kilometer. The real absorption is higher (largely because of water vapor) (see Liepmann & Roshko 1957). At 10 kHz, the peak absorption of 270 dB per kilometer occurs at 20% relative humidity, and the typical absorption is 100 dB per kilometer. This means that an easily detectable high frequency calibration signal inside the array quickly becomes undetectable away from the array. The absorption also means that the signals from distant sources do not contain such high frequencies, so the calibration system can be running at all times without interfering in the detection of signals. A low-pass filter extracts distant sources, and a high-pass filter extracts the calibration signal.⁵

We will assume that all N microphones have a transmitter. This is the simplest case to analyze, and the simplest system to deploy. Let the measured delays between microphone i and j be Δt_{ij} , with constant⁶ measurement errors σ_t . If the model microphone positions are $\mathbf{x}_i$, the sound speed is c_s , and the average wind speed is $\mathbf{v}$, then we model the delays by

$$\Delta t_{ij}^M = \frac{|\mathbf{x}_i - \mathbf{x}_j|^2}{c_s |\mathbf{x}_i - \mathbf{x}_j| + \mathbf{v} \cdot (\mathbf{x}_i - \mathbf{x}_j)}. \quad (11)$$

If there is no wind ($\mathbf{v} = 0$), then the delay is simply the sound travel time, $\Delta t_{ij}^M = |\mathbf{x}_i - \mathbf{x}_j|/c_s$, while if there is a wind, it includes the asymmetric delays seen in the upwind and downwind directions. We fit the delays using a χ^2 statistic

$$\chi_1^2 = \sum_i^N \sum_{i \neq j}^N \left(\frac{\Delta t_{ij} - \Delta t_{ij}^M}{\sigma_t} \right)^2 \quad (12)$$

where the sum extends over $i, j = 1$ to N . We cannot solve for all the positions using the relative delays. The χ^2 statistic is invariant under translations ($\mathbf{x}_i \rightarrow \mathbf{x}_i + \Delta \mathbf{x}$), rotations

⁴Effectively, if you know the code the signal-to-noise ratio of the signal is much higher than if you are just searching for a signal, because the code is the optimal filter for the signal.

⁵It may even be possible to use the high-frequency calibration signal to link the array. If frequencies below 10 kHz represent the signal, and we have N microphones, then the total bandwidth required to transmit all the signals is of order $100N$ kHz (10 bits per cycle uncompressed). Fewer bits and compression still could not compress the bandwidth enough to transmit the data at a frequency that has a reasonable propagation range. However, if we seek only impulsive signals, then the low duty cycle of impulsive events might allow the data to be squeezed into a bandwidth of 10 kHz and then permit it to be transmitted by the calibration signaling system.

⁶We do this for simplicity. A real system would use the estimated measurement error $\sigma_{t_{ij}}$ for each delay Δt_{ij} .

($\mathbf{x}_i \rightarrow R\mathbf{x}_i$, with R rotation matrix), and reflections ($x_i \rightarrow -x_i$ and/or $y_i \rightarrow -y_i$) of the array. We assume that the position of microphone $i = 1$ is the origin ($\mathbf{x}_1 = 0$) and that microphone 2 lies on the x axis ($\mathbf{x}_2 = (x_2, 0)$), to remove the translation and rotation degeneracies. The sound speed is degenerate with a rescaling of the array size, so it cannot be determined unless at least one microphone pair has a fixed, known separation.

The self-calibration procedure works in three steps. The microphones are emplaced, more or less randomly, over an area. The calibration system is activated, and the Δt_{ij} values are measured. Then we make a guess as to the microphone positions, with some accuracy σ_x for their positions, and minimize the χ^2_1 statistic. There are $N_{data} = N(N - 1)$ data for one calibration measurement. In three dimensions, including the wind velocity, there are $N_{par} = 3(N - 1) + 1$ parameters fit to the data, leaving $N_{dof} = (N - 3)(N - 1) + 1$ degrees of freedom (d.o.f.). A good fit to the time data has $\chi^2_1 = N_{dof}$, and a one standard deviation range of $(2N_{dof})^{1/2}$ for $N_{dof} \gg 1$. Although $N_{dof} > 0$ for $N < 5$ microphones, the χ^2_1 is still indeterminate because most of the extra d.o.f. are related to the wind speed not the microphone positions. The issues we must address are: (1) how well must we know the positions of the microphones to reliably find a good fit to the measured delays, (2) do good fits to the delay provide accurate positions for the microphones, and (3) can we reliably extract the mean wind speed. The second issue is the question of whether there is more than one microphone configuration that reproduces the observed pattern of time delays even after removing the translation, rotation, and reflection degeneracies.

To test the self-calibration idea, we generated random, "self-avoiding," two-dimensional arrays, with the coordinates of the microphones uniformly distributed in a square box $\Delta X = 50$ m on a side. We required all microphones to be separated by at least one meter. Then for a given time accuracy σ_t , and an initial position guess accuracy of σ_x , we examined how often we found a good model for the measured time delays, and how well good models for the time delays fit the true microphone positions. Minimization was done using the Levenberg-Marquardt method (see Press et al. 1992). We examined a wide range of parameters for the timing accuracy, uncertainties in the initial positions, and wind speeds. The problem is remarkable robust, converging to the correct solution even when the uncertainties in the initial microphone positions are larger than the typical microphone separations. The algorithm also measures the wind speed very accurately. If the initial position uncertainties are too large, the algorithm sometimes converges to a reflection of the array or fails to converge. Both of these problems are easily fixed. Very coarse extra constraints can prevent convergence to reflections, and a new random starting guess will typically converge. We show the results from some random trials in Figures 2, 3, and 4.

In short, by adding emitters to the microphone assemblies, we can robustly achieve accurate relative positions for the microphones and accurately determine the average wind speed across the array. If the timing accuracy of the array is limited by unknown timing offsets in the transmit and receive parts of the array (e_{it} and e_{ir}) then an array with 10 microphones can also calibrate these internal errors of the microphone system.⁷ There is simply no need

⁷In radio astronomy this particular problem is called self-calibration (Pearson & Readhead 1984), which was

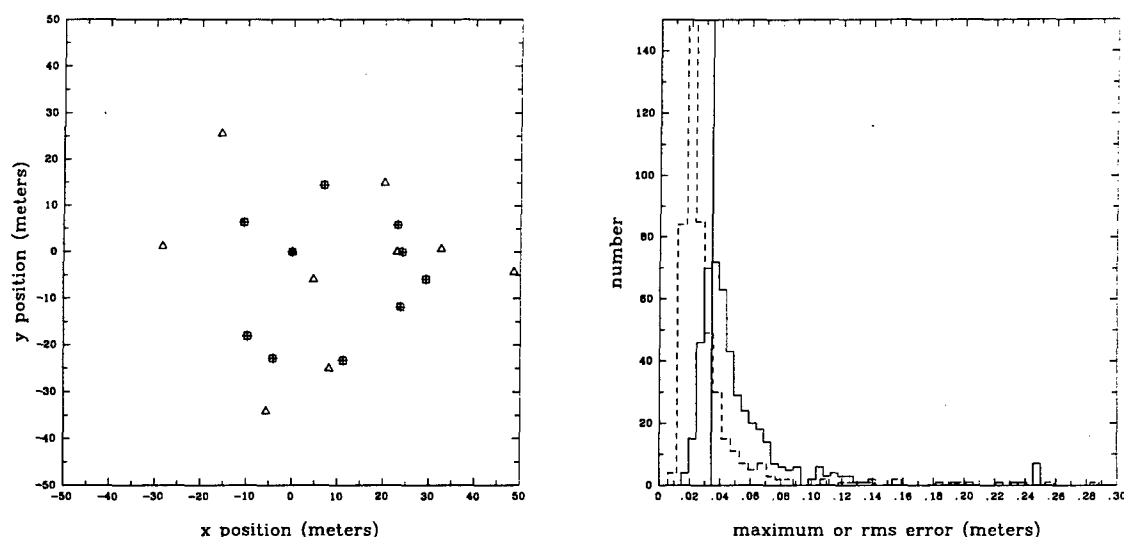


FIG. 2.—An example of convergence. A $N = 10$, self-avoiding, random array in a box $\Delta X = 50$ meters on a side with $\sigma_t = 100 \mu\text{sec}$. The rms distance of the initial positions from the true positions is 10 meters in each coordinate. The squares show the true positions, the triangles show the initial guess, and the crosses show the best fit positions. The rms distance between the true and best fit positions is 3 cm. The reference microphone is fixed at the origin, and the array is forced to be oriented along the x axis.

FIG. 3.—Histograms of the maximum (solid) and rms (dashed) errors for the parameters in Figure 1. The histograms show the 480 of 500 realizations that converged. The vertical solid line marks the characteristic error $c_s \sigma_t = 0.034$ meters.

to restrict the design of ground based acoustic arrays to rigidly mounted assemblies with short baselines, because it is simple to deploy and calibrate these intermediate size arrays.

4 EXTERNAL CALIBRATION

The self-calibration technique only determines the relative positions of the microphones up to a rotation, and cannot determine the absolute position of the array. The absolute position we assume will be set by GPS, and the 10 meter positioning accuracy of the GPS system sets a minimum error throughout the system. Averaging the GPS signal over long time periods or using differential GPS can reduce this uncertainty but the orientation of the array is difficult to determine with sufficient accuracy using GPS. Recall from §2 that we want bearing accuracies of one mrad (0.06°). If the GPS error uncertainty is σ_{GPS} and the longest baseline in the array is b_{max} then measuring the GPS positions of the two maximally separated microphones only determines the orientation of the array to $\sqrt{2}\sigma_{GPS}/b_{max}$. For an intermediate size array with $b_{max} \sim 10^2$ meters, the GPS accuracy must exceed $7(100\text{m}/b_{max})$ cm. While such accuracy can be achieved (with some difficulty) using differential GPS (D-

the genesis of this scheme.

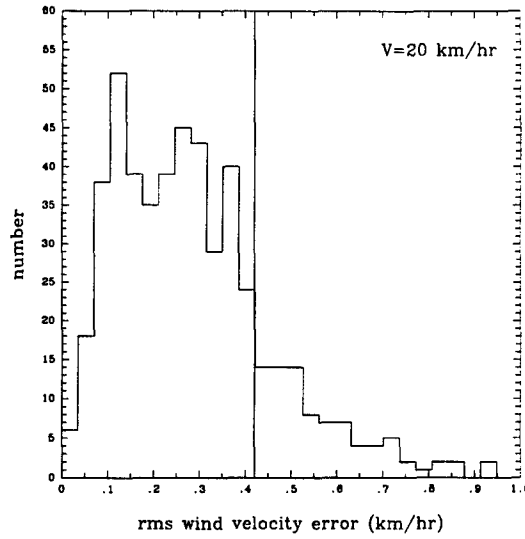


FIG. 4.—Histograms of the error in the wind velocity ($|\mathbf{v}_{mod} - \mathbf{v}_{true}|$) for a $|\mathbf{v}_{true}| = 20 \text{ km hr}^{-1}$ wind. The model uses $N = 10$, $\Delta X = 50$ meters, self-avoiding random arrays with $\sigma_t = 100 \text{ } \mu\text{sec}$ timing accuracy. The initial microphone positions are uncertain to 10 meters, and the initial wind speed is set to zero. The vertical solid lines marks the characteristic wind speed error of $c_s^2 \sigma_t / 2\Delta X = 0.42 \text{ km hr}^{-1}$.

GPS), it is not simple and it will drive up the cost. The reason that GPS (or D-GPS) cannot either calibrate the array directly or determine its orientation with sufficient accuracy for the error limits required of the array is that the separations of the microphones are too small compared to the accuracy of the GPS positions. However, if we can increase the calibrating baseline to $10^3 \sigma_{GPS}$, then the orientation can be calibrated accurately, or in the case of a counter-battery array, the position of each microphone can be accurately determined. External sources are used to calibrate some of the commercially available counter-battery arrays (Jane's 19??).

Suppose that we have a “calibrator,” a device that consists of a GPS receiver, an acoustic generator, and a radio. The acoustic generator can either be a repeating source (like an air-horn, or device to automatically fire blanks), or a self-destructing explosive source for longer ranges. The calibrator is placed some large distance $\mathbf{X}_c$ from the array ($b \ll |\mathbf{X}_c|$). It could be placed by hand, air-dropped, or even fired by a mortar. After placement, it determines its GPS position, and transmits the position and an agreed on execution time to the array.⁸ At the execution time the calibrator sets off its acoustic pulse, and the array determines the N arrival times t_i .

⁸The array need not know the exact execution time to calibrate positions, although there are two advantages to knowing the emission time. The first advantage is that the mean sound speed between the calibrator and the array can be determined to accuracy $\sigma_{GPS}/|\mathbf{X}_c|$. The second advantage is that the array knows both where and when to look for the calibration signal.

Assume we have M calibrators with positions $\mathbf{X}_k$ and GPS measured positions of $\mathbf{X}_k^{GPS}$, N microphones with positions $\mathbf{x}_i$ and GPS measured positions of $\mathbf{x}_i^{GPS}$, and that GPS positions are uncertain to σ_{GPS} . Assume we still have a self-calibrating array for which we measure the Δt_{ij} , and that we measure the time delay between the emission of the signal at calibrator k and microphone i , Δt_{ik}^C , where our model for this delay is

$$\Delta t_{ik}^{CM} = \frac{|\mathbf{x}_i - \mathbf{X}_k|^2}{c_s |\mathbf{x}_i - \mathbf{X}_k| - \mathbf{v} \cdot (\mathbf{x}_i - \mathbf{X}_k)}. \quad (13)$$

The simplest way to solve for the microphone positions to the (high) internal accuracies permitted by the timing errors and the lower absolute accuracy permitted by the external calibrators is to simultaneously determine all the coordinates (both $\mathbf{x}_i$ and $\mathbf{X}_k$) and the wind velocity $\mathbf{v}$ simultaneously. This will automatically determine relative coordinates as precisely as the internal timing requires while determining the absolute coordinates as well as possible given the distance to the calibrators. This defines the χ_2^2 statistic

$$\chi_2^2 = \sum_i^N \sum_{j \neq i}^N \left[\frac{\Delta t_{ij} - \Delta t_{ij}^M}{\sigma_t} \right]^2 + \sum_k^M \sum_i^N \left[\frac{\Delta t_{ik}^C - \Delta t_{ik}^{CM}}{\sigma_t} \right]^2 + \sum_i^N \left[\frac{\mathbf{x}_i - \mathbf{x}_i^{GPS}}{\sigma_{GPS}} \right]^2 + \sum_k^M \left[\frac{\mathbf{X}_k - \mathbf{X}_k^{GPS}}{\sigma_{GPS}} \right]^2. \quad (14)$$

The first term in this expression is the self-calibration χ_1^2 statistic from the previous section. The second term is the normalized error for fitting the arrival times of the M calibration signals, and the last two terms are the normalized errors for the absolute positions of the microphones and calibrators given their GPS positions. Note that the errors in the calibrator timing are the same as the errors in the internal timing – the position of the array relative to the timing position of the calibration signal is very accurate, but the absolute position and orientation of the total array+calibrator system is only as accurate as the errors in the GPS coordinates. There are, of course, variant system designs, the most obvious of which, and the one used in existing counter-battery arrays, is to dispense with the internal calibration signals (the first term) and use only the external calibration signals.

Since every microphone and calibrator has a GPS position, this model is never degenerate. However, without at least one external calibration signal, the array accuracy drops to that permitted by GPS microphone positions, and without distant external calibrators the absolute orientation of the array will be poor. It is worthwhile using more external calibrations then are necessary to orient the array for two reasons. First, with several external calibrators in different directions, the mean wind speed can be determined over large distances. While determining the wind speed over the array using internal calibration is better than a single meteorological station measuring the wind speed at a particular point, the still larger distances to the calibrator give even better estimates of the mean wind speed expected for real signals.

The second advantage of extra external calibrators is that they constrain the absolute positions of the array relative to the calibrators more accurately and allow a direct measurement of the sound speed. Recall that changes in the sound speed (with temperature primarily) are degenerate with rescaling the array size, so that the sound speed cannot be measured unless there is a fixed, known baseline in the system. While a rigidly mounted

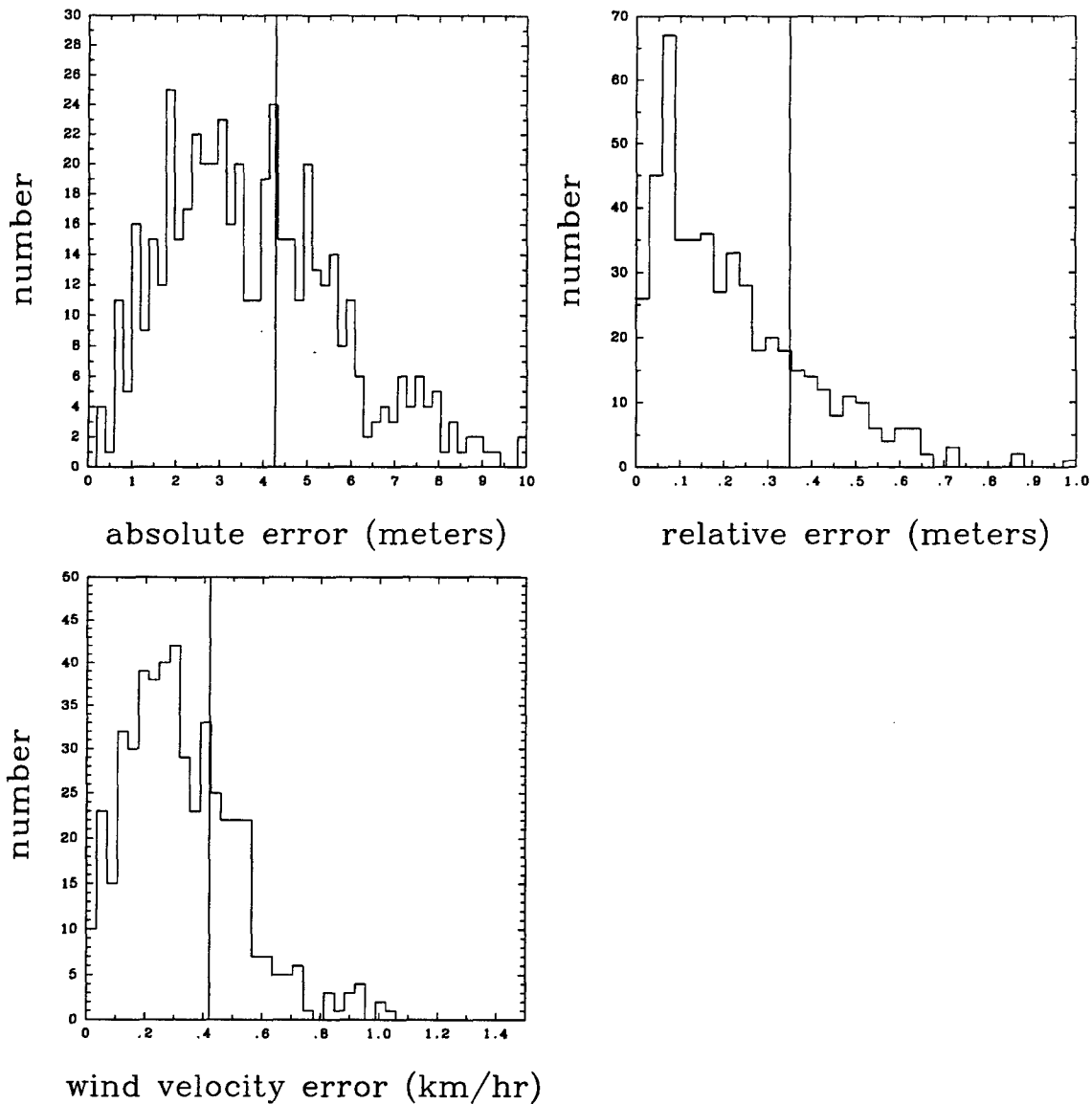


FIG. 5.—The effects of an external calibrator. The array parameters are the same as in Figure 4, but we have added one external calibrator at a range of 1 km, and assumed $\sigma_{GPS} = 10$ m. The top left panel shows a histogram of the mean error in the absolute position of the array, and the vertical line shows the expected uncertainty. The top right panel shows a histogram of the mean relative uncertainty in the position of the array elements. This is dominated by the ability of the calibrator to determine the absolute array orientation, and the expected error is marked by the vertical line. The bottom left panel shows a histogram of the errors in the wind speed measurement and its expected uncertainty.

microphone pair adds a fixed baseline, it is difficult to have a rigid mount large enough to supply the desired accuracy of $\delta c/c_s \sim 10^{-3}$. The sound speed varies by about $0.6 \text{ m s}^{-1} \text{ K}^{-1}$ so the mean temperature along the path must be known to approximately 0.5 K for $\delta c/c_s \sim 10^{-3}$. While a meteorological station can determine the temperature and pressure (and hence the sound speed) at a point, we really need some mean sound speed along the ray. The calibration signals supply a fixed baseline, which is uncertainly only to the GPS positional uncertainty. Hence the sound speed can be directly measured to $\sigma_{GPS}/|\mathbf{X}_c| \sim 10^{-2}$ (10^{-3} for D-GPS) *averaged over the long path to the calibrator*. Once the array is accurately calibrated, it can determine changes in the sound speed by redoing the internal calibration and looking for apparent changes in the array size.

We experimented with adding one external calibrator to the standard model array used in the early sections. Figure 5 shows some of the results, and illustrates the hierarchy of residual errors. The absolute position of the array is still uncertain to $\sqrt{2}\sigma_{GPS}/\sqrt{N+1} \sim 4 \text{ m}$ if the individual GPS errors are uncorrelated. Relative to the position of one microphone, the positions of the other microphones are uncertain to the accuracy with which the orientation of the array was determined, $\Delta x \Sigma_{GPS}/|\mathbf{X}_c| \sim 0.35 \text{ m}$, and relative to a particular microphone baseline the uncertainties are that of the self-calibration system, or about 0.04 m. Again, convergence is very robust, and with the external calibrator we can no longer find the reflection symmetric arrays because the calibration signal provides a direction.

The beauty of adding external calibration sources is that the resulting targeting accuracy of the array is exactly the GPS accuracy used to calibrate it *if the physical size of the array is large enough to reach that accuracy and ignoring other sources of error*.

5 TARGETING CORRECTIONS

All the previous discussion has focused only on the problem of deploying an array that is sufficiently large and well characterized to have the capability of acoustically locating targets with 10 meter accuracy at 1 km. We now want to use the array to target an acoustic source in the presence of all the other errors that interfere with the operation of the array such as inhomogeneous/turbulent winds, temperature gradients/refraction, and reflections. The actual targeting accuracy of a well-calibrated array is always worse than the calibrated accuracy.

The only means of removing these errors is "sound-on-sound" ranging or a related method. Sound-on-sound ranging has certainly been used in counter-battery acoustic arrays (particularly by the Russian, then Soviet, army – see *Acoustic Weapons Location*). In this approach, the array provides a predicted position for the enemy artillery, and is used to target counter-battery fire. The array listens for the explosions of the counter-battery fire, and compares the position of the explosions to the position of the original source. The difference in the apparent positions provides a correction for a second counter-battery salvo that corrects not only for the errors in the original acoustic position but also for the errors in the aiming of the counter-battery fire itself. For counter-battery use, the problem with this approach is the slowness of the iterations. It takes 30 seconds for sound to travel 10

km, so the time for each cycle of firing, sensing, and retargeting is a minute or more. A self-propelled gun can move faster than you can correct the positions.

The idea of sound-on-sound ranging should work very well against snipers or other close in targets because the cycle times drop from minutes to seconds. We can use either the "classic" sound-on-sound method, or a variant using pre-emplaced calibrators. The fast cycle times require an automated or semi-automated system. Conceptually we imagine a mode in which the the array can directly aim and fire the the 20 mm canon of a Bradley AFV or even the main gun of an M1-A1 subject to a "go/no go" decision from the operator. The particular weapon system would depend on the level of collateral damage that could be tolerated.

The sound-on-sound method assumes that the array can detect the impact or explosion of the fired rounds, compare them to the position of the original source and then make corrections. The goal is not to determine the absolute position of the source, but simply to make the acoustic positions of the original source and the response coincide. Between the uncertainties in the original acoustic position and the errors of the responding weapon, there will probably be collateral damage in a region some tens of meters around the target. If wind conditions are statistically steady during the response, the convergence rate of the response onto the target should be quadratic in the number of detected impacts. This is a standard minimization problem, with the parameters being the two angles specifying the weapon's firing direction.

If collateral damage is not acceptable, we need to try to remove the propagation errors from the position of the source. We can do this using the calibrator sources discussed in §4. Suppose we (secretly) plant extra calibrators in the region we believe the snipers may lie. These calibrators determine their positions and then lie quiescent. When the array picks up a target, it signals the calibrators near the target to send their GPS positions and trigger their acoustic signals. This gives the array one or more sources at known positions "near" the true target. By comparing the acoustic positions and known positions for the calibrators, the array can compute a corrected position for the target source including the effects that are not present in the simple models for signal propagation used to calibrate the array. In essence, we know the acoustic position of the target $\vec{x}_A$ and both the acoustic $\vec{x}_{ai}$ and physical $\vec{x}_{pi}$ positions of the calibrators, so the physical position of the target is

$$\vec{x}_p = \vec{x}_A + \langle \vec{x}_{pi} - \vec{x}_{ai} \rangle \quad (15)$$

where the $\langle \dots \rangle$ denotes an appropriate average over the differences between the acoustic and physical positions of the calibrators.

6 CONCLUSIONS

We discussed how to deploy, calibrate, and target using "intermediate" scale acoustic arrays. These are arrays of microphones in a region 100 meters on a side, and they are intermediate in size between counter-battery arrays and rigidly mounted arrays. A calibrated intermediate array can locate acoustic sources to 10 meters or better at a range of 1 kilometer.

The initial deployment of the array consists of scattering the sensors over the area. Ten microphones in any pseudo-random, self-avoiding pattern will suffice.

For the array to be useful, we must calibrate the array by accurately determining the positions of the microphones. For intermediate arrays, GPS positions are not sufficiently accurate, although very precise differential GPS (D-GPS) positions may be sufficiently accurate. Roughly, we need to know the microphone positions to 10 centimeters for 100 meter separations to achieve targeting accuracies of 10 meters at ranges of 1 kilometer.

Self-Calibration: If each microphone also contains an ultrasonic transducer that can transmit a pseudo-random code, then the array can self-calibrate or self-cohere. If there are more than 5 microphones and the array can measure the time it takes a sound wave to travel between every microphone pair, then the array can determine the relative positions of all microphones to the timing accuracy of the array – typically a few centimeters. In the simplest design, every microphone consists of a microphone/transducer pair, and this allows accurate determinations of the mean wind speed over the array (accuracy $\sim 0.5 \text{ km hr}^{-1}$). Even if D-GPS system is available and can locate the microphones to the requisite accuracy, it is worth maintaining the acoustic calibration system because of its ability to continuously monitor the average wind speed and variations in the sound speed, both of which are key elements of accurate targeting.

External Calibration: External calibrators consist of a GPS receiver, a radio to communicate with the array, and a transducer. At least one external calibrator is needed to accurately determine the orientation of the array – the longest baselines of the intermediate arrays are not long enough to use GPS to determine the orientation. The size of the array should be chosen to match the accuracy of the GPS position of the external calibrator. External calibrators fix the orientation of the array, and determine the mean sound and wind speeds over paths similar to those for a real source, rather than only near the array. Maintaining target accuracy may require periodically recalibrating the array.

Self-calibration and external calibration easily, rapidly, and precisely calibrate intermediate arrays. We did not discuss the detection and location of target sources, but the process of targeting a response at a source is closely related to calibration. Even though the calibrated array has an intrinsic physical accuracy of 10 meters, errors in the propagation model (spatially varying winds, spatially varying temperatures, and reflections) degrade the precision with which sources are located. If the physical effects leading to the errors vary slowly with time (temperature/reflections/steady wind gradients) or are statistically stable (some gusty winds), then the calibration methods can be used to give precise targeting information.

Sound-on-Sound Targeting: In this case we assume a semi-automated response to the source of the sound using a weapon whose impact in the target area can be detected by the array. By minimizing the difference in the acoustic positions of the source and the weapon impacts, the response can be brought onto the target. Such methods have been used for counter-battery fire, where the cycle times of a minute or more reduce the utility of the method. At the 1 km ranges probed by an intermediate array, the cycle times are much faster. However, if the accuracy of the acoustic positions is degraded by propagation effects,

then there will be considerable collateral damage caused by the badly targeted rounds at the start of the process.

Recalibration: If extra calibration sources are implanted in the area where targets are expected, then they can be used to quickly recalibrate the array to include the current environmental effects. When a target is detected, the array directs the calibrators near the target to relay their positions and send an acoustic signal. By comparing the true and acoustic positions of the calibrators, and updated position estimate for the source can be made. This can be used to initiate the sound-on-sound targeting procedure with reduced risks of collateral damage.

REFERENCES

- Acoustic Weapons Location, 197?, (McMillan Science Associates: Los Angeles)
- Battlefield Acoustic Sensors, Advanced Technology Assessment Report, 1992 Army Material Command
- Hewish, M., & Pongelley, R., 1990, Defense Electronics & Computing, IDR Supplement 4/1990, 47
- Pearson, T.J., & Readhead, A.C.S., 1984, Annual Reviews of Astronomy & Astrophysics, 22, 97
- Press, W.H., Teukolsky, S.A., Vetterling, W.T., & Flannery, B.P. 1992, Numerical Recipes, (Cambridge Univ. Press: Cambridge)
- Janes Sensors and ??
- Liepmann, H.W., & Roshko, A., 1957, Elements of Gas Dynamics, (Wiley: New York)
- Woods, A., 1941, Acoustics, (Interscience: New York)

**F. TECHNICAL ISSUES PERTAINING TO
ENVIRONMENTALLY SOUND SHIPS**

**Ann R. Karagozian
Mechanical, Aerospace and
Nuclear Engineering Department
University of California, Los Angeles
Los Angeles, CA**

TECHNICAL ISSUES PERTAINING TO ENVIRONMENTALLY SOUND SHIPS

1 Abstract

This "thinkpiece", written in partial fulfillment of the requirements for participation in the Defense Science Study Group, is designed to be a brief overview of military shipboard waste problems and potential solutions. This document summarizes the types of waste streams generated on military vessels, national and international restrictions on the pollution of waterways, legislation concerning the performance of waste treatment systems, current shipboard waste treatment systems, and future technologies that could play a key role in the development of environmentally sound ships. The study is based on a literature review as well as fairly extensive interviews with individuals involved with shipboard waste treatment issues within the government, at various research labs, and in industrial organizations. While clearly not intended to be a complete treatise on shipboard waste issues, many of the advantages and disadvantages of various treatment technologies are discussed, and recommendations for future directions are made. It is hoped that this study will provoke useful dialog toward the attainment of realistic solutions to the complex environmental problems associated with military vessels.

2 Introduction

With increasing restrictions placed on ship discharge by the international community, U.S. maritime organizations are actively seeking to develop technological solutions for the treatment of shipboard wastes. The types of wastes typically generated on board military ships can be classified as solid (plastics, food wastes, paper, cardboard, medical wastes), liquid ("non-oily" wastes such as blackwater and greywater and "oily" wastes such as bilge cleaning and de-rusting fluids, antifouling paints, and waste oils), and gaseous (air emissions from marine diesel or gas turbine engines). These waste streams are similar to those found and treated to a limited extent on large commercial ships such as cruise liners. The difficulties in implementing waste treatment systems on board existing military ships (Navy as well as Coast Guard vessels) have to do with the volume of wastes generated during long military missions, the limited space and power available on board ship for waste treatment systems,

and the fact that the waste treatment system(s) must not hamper the vessel's fundamental military mission, whether during peacetime or during war.

The purpose of the present study is to survey the environmental problems specific to military marine vessels, to assess the technical approaches currently being implemented or under evaluation for shipboard waste applications, and to suggest alternatives that the military might want to consider for on-board waste destruction.

It should be noted that this study does not investigate other environmental problems that face the Department of Defense at large, including the destruction of chemical warfare agents, propellant devices, or unexploded ordinance, nor does it consider the development of alternatives to ozone-depleting fire retardants (chlorofluorocarbons). In addition, this study will focus primarily on hazardous and "municipal" waste streams generated on ships and will not discuss radioactive wastes to any significant degree, although nuclear waste disposal is clearly a very serious problem and does have relevance to nuclear-powered ships and submarines.

3 Environmental Restrictions and Legislation for International Waterways

The 1973 International Convention for the Prevention of Pollution from Ships (MARPOL) as well as its associated 1978 Protocol have provisions which significantly restrict current and future discharge from maritime vessels. Provisions of MARPOL requiring enforcement by maritime nations include[1]:

- Annex I: Regulations for the prevention of pollution by oil
- Annex II: Regulations for the control of pollution by noxious liquid substances in bulk
- Annex III: Regulations for the preventions of pollution by harmful substances carried by sea in packaged forms, or in freight containers, portable tanks, or road and rail wagons
- Annex IV: Regulations for the prevention of pollution by sewage from ships
- Annex V: Regulations for the prevention of pollution by garbage from ships
- Annex VI: Regulations for the prevention of pollution by air emissions from ships

While this Convention specifically exempts warships, and specifically states that government-owned vessels should comply so far as reasonable and practical, many nations expect their own and other navies to meet these restrictions[2].

MARPOL Annex I, for example, restricts the discharge of oily waste from ships, limiting the liquid dumped within 12 nautical miles from shore to no more than 20 ppm oily waste. The U.S. Clean Water Act (1977), in fact, can place further restrictions on seafaring vessels

in that discharge regulations are governed by local state water quality standards. As a consequence, it is expected that the International Maritime Organization (IMO), which convened MARPOL, will further restrict its oily waste dumping regulations[3].

MARPOL Annex V regulations concerning solid waste include a total ban on plastics discharge anywhere at sea, a prohibition of all solid waste discharge in specially designated maritime areas (e.g., the North Sea, Baltic Sea, and Mediterranean Sea), and a restriction on the disposal of food and other processed solid waste to at least 12 nautical miles from the nearest point of land. The U.S. Congress voluntarily enacted MARPOL Annex V requirements in its 1987 Marine Plastic Pollution Research and Control Act (MPPRCA). Deadlines for compliance have been set for surface ships at 12/31/98 for plastics and 12/31/00 for special areas; for submarines these deadlines are 12/31/08 for both plastics and special areas. It is not expected that the restrictions on sea pollution from American vessels will be altered in the near future, even with alterations in the U.S. political climate, due to international pressures.

Air emissions from ships are also to be restricted in the future. By 1999, for example, new marine diesel-powered vessels will have specific restrictions on the generation of nitrogen oxides or NO_x (9.2 g/kW-hr), unburned hydrocarbons (1.3 g/kW-hr), carbon monoxide (11.4 g/kW-hr), and particulate matter (0.54 g/kW-hr). Low emissions diesel and marine gas turbine engines are currently being developed for future generation ships, since backfit of existing emissions reduction systems onto warships is prohibitively expensive.

The importance attributed to maritime environmental protection by the international community is evidenced by the formation of a special working group (SWG/12) within the NATO Naval Armaments Group, which provides expertise for the evaluation of alternative shipboard waste solutions.

4 Shipboard Waste Streams

While solid wastes generated on military vessels are generally of the same type as typical municipal solid wastes, shipboard hazardous liquid wastes are rather unique to seafaring vessels, both in type and volume. Solid waste generation on board a ship is of the order of 3 lbs./person/day, which includes food wastes, paper/cardboard, metals, glass, and plastics.

In many respects, large military vessels such as carriers, with thousands of military personnel, may be considered to be small cities while simultaneously "floating machine shops". The on-board technological processes associated with the highest rates of disposal of hazardous liquid wastes between 1988 and 1992 include, in order of disposal rate[4]: industrial waste bilge cleaning, fluids changeout, painting operations, abrasive blasting, chemical stripping, industrial operations, boiler cleaning, electroplating, solvent degreasing, ordinance operations, spill cleanup, and pipe flushing. Many of these processes are a part of normal ship operations, such as ship boiler cleaning, while other processes are associated with ship's military mission, including ordinance operations. There are ongoing efforts to reduce the toxicity of oily liquids[2]. For example, the development of high performance yet environmentally friendly anti-foulant paints and hull coatings is currently of interest. Research is

focusing on silicone based coatings for the near term and natural bio-inhibitors (such as eel grass and coral) for the long term.

Most often, hazardous liquids are stored on-board in 55-gallon drums and off-loaded at port, where they are typically transported to landfills or land-based incinerators such as rotary kilns. While the volume of hazardous liquids generated on military vessels is much smaller than that of plastics and solid food wastes, the potential for health hazards due to spillage or improper handling is significant.

5 Requirements for On-Board Waste Treatment Systems

While the basic goals of current and potential on-board waste treatment systems are the fulfillment of MARPOL and Congressionally mandated discharge standards, it is logical to assume that U.S. EPA and/or international regulations concerning benchmarks and continuous monitoring of equipment are also likely to apply. Table 1 contains performance standards for emissions monitoring levels proposed or required of U.S. municipal, medical, and hazardous waste incinerators, in addition to requirements of hazardous waste incinerators in the European Union[5].

Table 1. Performance standards for municipal and hazardous waste incinerators.

Pollutant	U.S. Municipal Waste Combustors (Proposed)	U.S. Medical Waste Incinerators	Hazardous Waste Combustors	European Union Haz Waste Combustion
(ng/dscm) TEQ and Total Congeners	0.20 TEQ or 13 Total	1.9 TEQ or 80 total	0.2 TEQ	0.19 TEQ
Particulate Matter (mg/dscm)	15	30	23 (9 hr avg)	13 (24 hr avg) 13-39 (30 min avg)
Hg (µg/dscm)	80 or 85% Reduction	470 or 85% reduction	30 (12 hr avg)	6.5
Semivolatile metals (µg/dscm)	Cd: 10 Pb:100	Cd: 50 Pb:100	Cd+Pb:40	Cd: 3.25 Ti:3.25 Pb:65
Low Volatile metals (µg/dscm)	none	none	80 (As, Be, Cr, Sb)	565 (Sb, As, Cr, Co, Cu, Mn, Ni, V, Sn)
CO, ppmv	50-150 (4 to 24 hr avg)	50 (12 hr avg)	100 (1 hr avg)	82 (24 hr avg) 104 (30 min avg) 156 (10 min avg)
Total hydrocarbons ppmv	none	none	6 (1 hr avg)	6 (24 hr avg) 8-18 (30 min avg)
HCl and Cl ₂ , ppmv as HCl equivalent	25 or 95% reduction	42 or 97 % reduction	25 (9-hr avg)	6 (24 hr avg) 8-48 (30 min avg)
Residual Organics	none	none	99.99 % DRE and Site Specific Risk Assessment	
Continuous Emissions Monitoring Requirements	CO	CO	PM, CO, THC, HCl, H ₂	PM, CO, THC, HCl, H ₂

In this table, TEQ refers to total equivalent parts per trillion, PM refers to particulate matter, THC refers to total hydrocarbons emitted, and DRE refers to the destruction and removal efficiency of the hazardous waste treatment system. In addition to these specific regulations, the U.S. Resource Conservation and Recovery Act (RCRA) of 1980 identifies specific classes of hazardous wastes; it also mandates that trial burns in EPA-permitted incinerators be performed on the wastes' principal organic hazardous constituents (POHCs)[6]. As noted in the table, incineration (or other thermal destruction) of the POHC must yield DREs of at least 99.99 %, where the DRE is defined as

$$DRE \equiv \frac{(\dot{m}_{in} - \dot{m}_{out})100\%}{\dot{m}_{in}}$$

and where $\dot{m}_{in}$ and $\dot{m}_{out}$ are the mass flow rates of the POHC entering and leaving the system, respectively. DREs are usually reported in numbers of "nines", e.g., 99.99% destruction and removal efficiency is reported as "four nines DRE". Specific RCRA requirements also exist for *HCl* and particulate emissions, and for continuous monitoring of combustion temperature, waste feed rate, combustion gas flow rate, and CO in the exhaust gas stream. Additional restrictions on the incineration of polychlorinated biphenyls (PCBs) are covered in the Toxic Substances Control Act (TSCA).

Whether or not shipboard waste treatment systems are required to meet the above standards, it becomes clear that there are technical and societal issues here that are peculiar to military ships. If waste is treated and neutralized on board the ship, the impact on the ship's military mission must be examined. For example, if high temperature treatment systems are implemented, one must consider how the vessel's infrared signature is altered. If a plasma-based system is the major on-board treatment mode, electromagnetic interference could be a major concern as well. Noise and safety issues become critical on military vessels due to the large number of personnel and tight quarters, and operator expertise must be considered in treatment system selection. The fact that there is a limited amount of power available from the ship engines to drive an on-board treatment system is another major factor; requirements of auxiliary fuel-fired generators could increase air pollution while further taking up space. In fact, it may be desirable in some cases to be able to recover energy from the waste treatment system. Clearly, on-board waste treatment systems will have to be highly compact and lightweight, while having a high degree of reliability and significant volume reduction of the waste stream.

The realistic alternatives to treating waste on-board are few. When waste is discharged overboard, irrespective of international regulations, the waste could disperse and degrade or be rendered non-toxic by the saline environment. Yet before this occurs, the waste entrained into the wake of the vessel could become detectable. If the waste is made to sink quickly so as to avoid detection, it could disturb the ocean's plant/animal ecosystem. If the waste is simply separated and stored on-board for discharge at port, precious space becomes occupied by garbage, and the potential for spillage or contamination increases. Of course, the ability to off-load waste at port is very much dependent on the local country and the circumstances of the ship's presence.

6 Current Shipboard Treatment Systems

As with current approaches being taken to solve land-based hazardous and municipal waste problems, shipboard waste strategies have taken two paths: waste minimization and waste treatment. Some of these approaches have been quite successful. The hazardous waste minimization (HAZMIN) program instituted by the Chief of Naval Operations, for example, established a goal of 50% reduction (by weight) of hazardous waste disposal during the period 1988-1992; during this period the Navy exceeded the reduction goal by 5%[4].

Contributing to this degree of hazardous waste "source" reduction was the Integrated Liquid Discharge System (ILDS), implemented by the Navy to reduce liquid waste volume by segregating and de-watering waste. These techniques include filtration for oily and non-oily wastes[7]. Ceramic membranes are used to filter oily wastes, while polymeric membranes are most suitable for non-oily wastes (e.g., sewage or "blackwater" and shower/laundry discharge or "greywater"). Filtered water that is then discharged has concentrations of waste that are below 15 ppm, within the current acceptable discharge limit. Reducing the waste discharge further is an area of continuing research and development at the Carderock Division of the Naval Surface Warfare Center.

Waste minimization efforts for blackwater have also seen some success. Vacuum collection systems, which are commercially available, are in the process of being implemented by the Navy; these reduced volume fixtures could produce a 50-70% reduction in blackwater generation. Greywater minimization efforts include appliances which require reduced flow due to recycling and low flow shower heads that can reduce greywater volume by 90%.

Actual on-board thermal destruction systems are not very widely implemented at present. Most aircraft carriers have two incinerators installed on board; these are typically of 1935-45 vintage and were designed to incinerate paper. A few other ships, e.g., amphibious assault ships and submarine tenders, also each have one on-board incinerator. The incinerators were originally installed in order to be able to destroy classified documents. Besides being highly inefficient and producing high levels of undesirable air emissions, these older incinerators are not capable of burning much more than paper, and are not designed for large volumes of waste nor for operation on the order of hours. A recent materiel failure took place during operation of one of these incinerators; this occurred when cardboard was attempted to be burned on the U.S.S. Wasp in Haiti. An outright accident also recently occurred when rags lying close to an incinerator caught fire on the U.S.S. Roosevelt; this was apparently due to operator error.

An on-board incinerator that is used to burn sewage, particularly in restricted waters such as in the Mediterranean, is the "vortex incinerator". This device, of 1950s vintage, has been implemented on roughly 30 DD 963 and DDG 993 class destroyers, and is fired using the same JP5 fuel that is used to run the ship's gas turbine propulsion system. Auxiliary fuel and air are injected into the vortex incinerator off-center in order to create a hot, swirling flow region. An atomizing spray comprised of liquid waste (which can contain up to 1% organics and 1% suspended solids) is injected into the hot, combustive vortical flow, where it is burned or thermally decomposed at a rate of roughly 0.5 gal/min. Ash which is produced is either

stored on board then off-loaded at port, or discharged off the ship in permitted regions. This device is often used when the vessel is within 3 nautical miles of shore.

It should be noted at this point that large commercial marine vessels, such as cruise ships, encounter similar problems with respect to restrictions on dumping waste into foreign as well as domestic waters. Current systems implemented on cruise ships tend to be of the "off-the-shelf" variety, including crushers/compactors for plastics, pulpers for paper and food wastes, and a few very large incinerators (of the order of ship "stories" tall). Much of the organic waste, including blackwater, is treated by filtration and dumped overboard in permitted regions. Increasing restrictions on off-loading waste from cruise ships at foreign ports has led this industry to seek alternative solutions to waste treatment and disposal; in fact, the cruise line industry is looking to the DoD for guidance in this area.

7 Future Directions in Shipboard Waste Treatment Systems for the Military

Clearly, the solution(s) to military and commercial shipboard waste treatment and disposal problems can take two fundamental paths: 1) use of commercially available "off-the-shelf" treatment systems or 2) development of novel, advanced treatment systems specially designed for shipboard utilization. Interestingly, the U.S. Coast Guard has taken the first approach, while the U.S. Navy has decided to pursue the second approach. Based on the volume and variety of wastes generated by vessels in each service, as well as the existence (or lack thereof) of applied research activities, these alternative strategies are perhaps not unexpected. Descriptions of specific systems considered in these categories are provided below.

7.1 "Off-the-Shelf" Technologies: U.S. Coast Guard

Commercially available waste systems have been determined a viable option for Coast Guard vessels, including cutters, buoy tenders, and ice breakers. Most of these vessels are out to sea for 30-60 day periods, with 130-170 persons on board, so the requirements for waste disposal are not as extensive as, say, for an aircraft carrier, although ice breakers can be away from home port for up to six months at a time.

Solid waste incinerators currently under evaluation for the Coast Guard are, for the most part, manually fed refractory boxes, such as the Kvaerner Golar waste incinerator which is being tested as a prototype on the CGC Storis in Kodiak, Alaska[8]. This particular incinerator can be fired using either diesel or JP5 fuel, and burns waste which is either batch or continuously fed. Temperatures in the combustion chamber lie between $1000^{\circ}K$ and $1500^{\circ}K$. The incinerator cannot destroy glass, however, which is typically separated and off-loaded at port. The Kvaerner Golar incinerator is said to give off relatively little smoke due to the temperature range of combustion, and resulting ash is minimal. Although the manufacturer claims that the incinerator has a burn capacity of 176 lb/hr, when one accounts

for the actual time spent in waste separation, incineration, cool down, and ash separation, the actual destruction rate may be as low as 30 lb/hr. Another fuel-fired, refractory box type of incinerator, the Venta-matic, has been used on ships in the past and upgraded versions are also under consideration by the Coast Guard. In most instances, much of the waste that is burned in a shipboard incinerator needs to be pre-treated, for example, passed through a dehydrator or an oily water separator. It is not clear how well this type of incinerator actually burns hazardous liquids, however, and sludge tanks are still expected to be utilized on Coast Guard vessels.

The Coast Guard consciously avoids incinerating plastic waste, and has purchased a number of trash compactors which are used for plastic as well as metal wastes such as aluminum cans. These compactors typically reduce waste volume by a factor of 20:1.

Off-the-shelf alternatives that the Coast Guard (or others) might also want to explore could also include completely integrated incineration systems, such as the Deerberg, advanced systems such as that by Synthetica, or specialized medical waste incinerators. In contrast to the simple refractory box concept for most off-the-shelf incinerators, the Deerberg system is fully integrated, with autofeed and automatic de-ashing, and includes its own air pollution control system of filters, NO_x catalytic converters, and scrubbers[5]. The combustion processes take place in three chambers, with the possibility of fixed or moving grates in the primary chamber, at peak temperatures ranging from 1200°K to 1500°K. Ash is automatically collected and can be stored or dumped easily. The Synthetica system is one which uses steam reforming and gasification to pyrolyze waste which is injected; this has also been used successfully in mixed radioactive/hazardous wastes.

Medical waste, while perhaps occupying a small portion of a ship's total hazardous waste volume, holds extraordinarily importance in terms of the necessity for proper disposal, for health reasons as well as for public perception. On land, medical wastes which are incinerated by the generating facility (e.g., the hospital) are not required to be strictly tracked from generation to disposal under the Federal Medical Waste Tracking Act of 1988[9]. Hence there has long been the incentive to incinerate medical wastes, and most hospitals and clinics have at least one incinerator. Medical waste incinerators are typically batch process units, and are increasingly becoming smaller in size and thus more suitable for marine applications. Today's medical waste incinerators generally have two chambers: a fixed hearth primary chamber which partially burns the waste under starved air (fuel rich) conditions, then waste and products are more completely destroyed in a secondary chamber fired with supplementary fuel and air. The secondary chamber can become quite large due to high temperature and residence time requirements. Predominant stack emissions typically include acids, e.g., HCl from the burning of chloride-containing plastics, and particulate matter from carbonaceous material. These are mostly removed by fabric filters or venturi scrubbers.

7.2 Emerging Shipboard Technologies: U.S. Navy

The Navy appears to be considering two alternative philosophies in its evaluation and development of shipboard waste treatment systems. One philosophy is to develop separate

waste treatment systems for specific classes of wastes, e.g., processors for plastics, pulpers for food and paper wastes, incinerators and/or other thermal treatment systems for concentrated oily and non-oily liquids, etc. An alternative philosophy emphasizes the development of one large scale waste treatment systems, such as plasma arc pyrolysis, that potentially could handle nearly all types of waste streams generated on board Naval vessels. While the current emphasis within the research branches of the Navy appears to be placed on plasma arc pyrolysis, a great deal of time and funding has been spent on the development of certain specialized systems, such as the plastics processor, and since some are being implemented already on existing ships, they are likely to remain viable alternatives in the future. These technological systems are described below.

7.2.1 Plastics Processing Systems

Plastic waste occupies only 5% of shipboard waste generated by weight, yet it can occupy 30% of the waste by volume, hence on-board volume reduction techniques have received considerable attention in recent years. While the specific MARPOL restrictions on dumping plastic waste do not take effect until 12/31/98, current Navy policy already prohibits such dumping. Navy policy also currently prohibits the incineration of plastics, due to anticipated problems with HCl and other emissions arising from the combustion of chloride-containing plastics, e.g., polyvinylchloride. As an alternative treatment system, plastics processors shred and then thermally compress plastic waste into flat disks which can then be stored on board before being off-loaded at port. Shredding the waste achieves a volume reduction of 3:1, while thermal compression yields further volume reduction by a factor of 30:1. This treatment technology has been under development by the Navy for a number of years, and has been tested on many vessels, including the U.S.S. Roosevelt, Washington, and Wasp. Recently (July 1, 1995) the Navy awarded a major contract to a plastics processor manufacturer to install a number of systems on board large surface ships during the spring of 1996.

While plastics processors vastly reduce plastic waste volume, and while the product of the process may be re-used (e.g., for pier pilings), the treatment system is not without problems. Contaminated plastics placed in the processor cause the resulting plastic slugs to be considered toxic; they are currently placed in sealed odor barrier bags before being stored. In addition, this technology is viewed by some as being an interim solution, in that biodegradable materials such as cellulose are being examined as alternatives to plastics.

7.2.2 Shredders and Pulpers

Shredders have been tested for a variety of solid waste streams, including metals and glass, and pulpers have been studied for many years for the treatment of food and paper wastes. While shredders have not been found to be optimal for metal and glass waste[7], which is still stored on board and off-loaded at port for recycling, pulpers have been found to be quite efficient for the treatment of food and paper wastes. Pulpers are able to treat up to 600 lbs/hr of waste, and are in operation on two aircraft carriers at present. The pulpable slurry which is formed can be mixed with sea water and dumped overboard. The slurry

apparently disperses quickly, so that detection is not a problem; it is also found to be benign for plant and fish life. As food and paper waste constitutes roughly 60% of Navy solid waste by weight, this technology can have significant benefits to overall waste stream reduction. In terms of the system's application to a "zero discharge" ship, the pulpable slurry could potentially be injected into a thermal treatment system such as an incinerator.

7.2.3 Liquid Waste Filtration Systems

As noted in Section 6, the development of advanced filtration systems to reduce oily and non-oily waste discharge is continuing, principally at the Carderock Division of the Naval Surface Warfare Center. It is sought to design membrane material which would be compatible with both oily and non-oily wastes, and which could yield virtually zero waste discharge. Recent testing of small scale ceramic membrane units for bilge oily water aboard the U.S.S. Ly Spear demonstrated efficient liquid waste filtration, with less than 5 ppm oil in the effluent[10] and an overall concentration factor of 100:1. The liquid waste concentrate, however, needs to undergo further on-board destruction to avoid the necessity for storage.

7.2.4 Supercritical Water Oxidation

Supercritical water oxidation (SCWO) is under consideration by the Navy as a treatment technology for the destruction of these concentrated liquids that result from on-board filtration processes. SCWO is a process by which oxidation of waste species takes place in H_2O at very high temperatures and pressures, well beyond the critical point (e.g., temperatures between $350^{\circ}C$ and $700^{\circ}C$ and pressures between 250 and 400 atm). Organic carbons are oxidized to form CO_2 and carbonates within a SCWO system, for example, and organic nitrogens become N_2 gas, ammonia, and N_2O . With residence times of the order of seconds to minutes, SCWO can produce high DREs (of the order 99.9999% or "six nines") in a relatively clean, easy to control system. SCWO is highly energy intensive, however, and commercially available systems tend to be large, heavy, and relatively expensive. Most likely an external generator would be required to create such a high pressure, high temperature environment for waste treatment. There is currently a large project sponsored by ARPA for the evaluation of SCWO as a hazardous treatment system for ships.

Technologies related to SCWO which also hold potential for hazardous as well as mixed radioactive-hazardous wastes include steam reforming, wet air oxidation, and high pressure hydrothermal processing[11]. High pressure hydrothermal processing extends over a much wider pressure range (400 to 1100 atm) than does SCWO, and has been shown to dissolve and reform sludges without neutralization or acidification.

Another non-combustion oxidation technique which has been shown to be efficient for hazardous as well as mixed (radioactive-hazardous) waste is electrochemical oxidation[12, 13, 14]. This process involves the placement of an electrochemical cell in an aqueous solution containing waste, and allowing toxic species to be oxidized at either electrode: toxic metals, for example are deposited on the cathode, while organic species are oxidized onto the anode to form carbon dioxide and water. Treatment of chlorinated organics such as PCBs actually

allows separation of chlorine to form chloride salt, which is easily stored or disposed. This technology has many advantages in that it operates at room temperature, it is silent, it is easy to control, it is easily scalable, and it presents no risks of uncontrolled emissions or explosions. It is likely to require a scrubber, however, for NO_x emissions. Bench scale as well as pilot scale testing of this process is ongoing at Los Alamos National Laboratory, for example.

7.2.5 Plasma Arc Pyrolysis

Plasma arc pyrolysis is an ultra-high temperature process which can be used very effectively to destroy hazardous and non-hazardous wastes in the absence of oxygen. In this process, a DC power source is required to form an electrically discharged plasma arc; argon or helium are typically used as the working gas, although air can be used in certain configurations. Equilibrium temperatures of the plasma lie in the range $5000^{\circ}C$ to $15,000^{\circ}C$. Solid waste as well as liquids and slurries can be fed into the arc jet or into the molten slag formed by the plasma arc, which can reach temperatures exceeding $2000^{\circ}C$ [15].

The high temperatures attained in the plasma and molten slag are very effective in destroying wastes; the high temperature breaks the waste molecules down to their constituent atoms, which remain in the elemental form or recombine to form simple molecules such as CO and HCl. The combustible gas product usually is burned in an afterburner. Plasma-based systems have been found to be very successful in destroying dioxins and furans, which can be present in oily wastes, by breaking their chemical ring structures. Many of the chemical aspects of the pyrolysis process are not yet well understood, however, and active research in this field is ongoing at a number of institutions, including the Naval Research Laboratory, the University of Maryland, and Mississippi State University.

As noted above, the plasma arc pyrolysis process has many potential advantages for waste treatment in general. The slag which is formed after waste pyrolysis is inert and of relatively low volume; this can be readily extracted from the system. There is also little or no ash discharge, and the gaseous discharge from the process can be passed through an afterburner and/or air pollution control equipment such as a venturi scrubber.

Not least among the advantages of the plasma-based system is the fact that it is not officially classified as an "incineration" process, although it is clearly one which destroys wastes at very high temperatures and which produces a gaseous discharge. The absence of a physical flame structure is what apparently allows plasma systems to avoid the incinerator label, which can be a tremendous advantage from the point of view of public acceptance. If the plasma device requires an afterburner, however, then it is classified as an incinerator, and EPA regulations pertaining to the incineration of hazardous and municipal waste would apply. The plasma arc system also has been used successfully for the separation and destruction of mixed radioactive and toxic wastes at several sites, for example, at the DOE Hanford facility.

In terms of the application to shipboard waste destruction, however, there are disadvantages to plasma arc pyrolysis that may outweigh the advantages. Current plasma arc

systems are very large, viewed by some to be impractical given the size constraints on Navy vessels, even on aircraft carriers. One alternative possibility is to have the plasma system (or other treatment systems) placed on a specialized "garbage ship", but tactical problems arise in this option. Other disadvantages to plasma arc pyrolysis include the requirement of a large DC power source needed to create the plasma (500 KW to 1 MW is quoted); this would necessitate having an external generator placed on the ship. This generator would likely be diesel fired, with air emissions problems of its own. In addition, the working gas for an inductively coupled plasma arc system needs to be inert (e.g., argon), thus additional gas storage problems could exist, although some plasma systems are able to run on air or steam. Electromagnetic interference (EMI) by the system could be a problem requiring extensive shielding, and the extremely high temperatures present in the device could significantly alter the vessel's infrared signature. Finally, the buildup of slag on the liners of the pyrolysis chamber (usually from glass) can be excessive, requiring periodic removal, but then, liner replacement is common to most thermal destruction systems.

Despite the above concerns, the Navy is putting a great deal of emphasis on the development of the plasma-based technology for future shipboard waste destruction. In fact, it was Congressionally mandated in 1994 to earmark \$1.8 million during FY95 for the development of this technology for Navy waste problems, and \$14 million has been programmed by the Navy into an Advanced Technology Demonstration scheduled to begin in FY97.

An alternative plasma-based system that may be suitable for shipboard waste treatment is the "silent discharge" or "non-thermal" plasma process[16, 17]. This process is based on the formation of free radicals in a non-equilibrium plasma which is created by electrical discharges from electrodes covered with a dielectric material. The non-thermal plasma contains very energetic electrons, with ions and neutral species at different temperatures so that non-equilibrium conditions exist. The radical species which are formed can thus destroy hazardous compounds at close to ambient gas temperatures. This process has also been used in conjunction with a thermal packed bed reactor to initially volatilize or combust liquid organics. Hazardous waste surrogates as well as trichloroethylene (TCE) and PCBs have been destroyed successfully in such systems; DREs of 99.9999% or six "nines" are obtained from the packed bed reactor, then subsequent non-thermal plasma systems can further remove toxins for an overall destruction efficiency of eight to nine "nines". This approach can have advantages over thermal plasma systems in that the non-thermal plasma component of the system is mostly silent and it destroys waste at ambient temperatures. Its power requirements are typically lower than that for thermal plasma systems for similar volumes of waste processed.

7.2.6 Shipboard Incineration Systems

Incineration is the most widely applied means of thermal treatment of hazardous wastes on land. The basic process utilizes high temperatures which are created via turbulent flame structures formed in an oxidizing environment. There are a variety of different types of incinerators, generally classified on the basis of the combustion chamber configuration, e.g.,

liquid injection, fixed-hearth, fluidized bed, and rotary kiln incinerators. All incinerators share similar fundamental features that enable them to destroy waste efficiently; these rules of thumb include: 1) high temperatures in the combustion zone (of the order 2000°K), 2) sufficient free oxygen in the combustion zone, depending on the waste type, 3) turbulent flow for thorough mixing of waste, oxidant, and, if present, supplemental fuel, and 4) sufficient residence times of the waste in the combustion chamber to allow reactions to proceed to completion (of the order 2 seconds).

On-board incineration of wastes could be an attractive waste treatment alternative from a technological point of view. In addition to providing a significant reduction in waste volume, it is often possible to recover a substantial amount of energy (heat) and material (e.g., acids) through incineration. Many metals are rendered less toxic through promotion to higher oxidation states during incineration[18]. Nevertheless, as a result of public opposition to the notion of incineration systems, only a small fraction ($<5\%$) of *combustible* wastes have historically been treated by incineration on land[19], and the technology has received mixed support within Navy circles. The fact that existing on-board incineration systems are extremely old and inefficient does not help to improve this perception.

Shipboard incineration systems could take the form of comprehensive waste treatment systems, such as the fixed hearth/refractory box or rotary kiln, which are capable of incinerating a wide variety of solid and liquid wastes. Yet most comprehensive thermal treatment systems which are available today and are capable of handling the waste volume of a large aircraft carrier, e.g., a rotary kiln, would probably be too large to be a viable alternative on tight quarters. Alternatively, specialized incineration systems could be used on board to treat specific waste streams, such as sewage concentrate, oily wastes, and/or food and paper slurries. The negative features of an on-board incinerator, of course, include the presence of a high temperature energy source, yielding a detectable IR signature for the ship (either for the primary ship or a "garbage" ship). Thermal insulation and shielding, noise control, and operator training and expertise are essential issues for the development of on-board incineration systems. In addition, air emissions from the incinerator would have to be strictly monitored, and it is expected that standard air pollution control devices may be required for continuous operation.

7.2.7 Advanced Incineration Systems: The Resonant Incinerator

From a technological point of view, incineration of shipboard wastes could be advantageous, especially if the thermal treatment system could be made compact, lightweight, and transportable. This is the basis for the evaluation and prototype development program underway at the Naval Air Warfare Center at China Lake, CA under SERDP 1994, the Strategic Environmental Research and Development Program[20]. The aim of this study is to develop a compact shipboard incinerator for the incineration of fluid wastes, so that the device could be used either as a stand-alone incinerator for shipboard fluid wastes or as an afterburner for a larger thermal treatment system.

One of the technologies under consideration for this SERDP activity is the **Resonant**

Incinerator, a device that was developed at UCLA and which has been studied there since 1988. It is an incinerator which has been demonstrated to achieve high volumetric heat release rates in a compact combustion cavity, with extremely high waste surrogate destruction under natural or forced acoustical excitation.

The resonant incinerator is an adaptation of the dump combustor concept, typically used in ramjet engines. Dump combustors are characterized by the sudden expansion of a fuel-air mixture into a combustion cavity in which large scale recirculation zones can be stabilized. These recirculation zones contain mostly high temperature combustion products, assisting with flame stabilization at the location of sudden expansion or "dump plane". In the current resonant incinerator configuration, fluid waste surrogates are injected into these stable recirculation regions. Hence the basic rules of thumb for efficient incinerator operation can be achieved: 1) high temperatures, near the adiabatic flame temperature, are present in the recirculating regions, 2) sufficient oxygen can be present in these zones for waste destruction, especially with oxygen enrichment in the core flow, 3) high levels of mixing (although not necessarily turbulence) occur in the recirculation zones, and 4) increased residence times exist for the waste trapped in the recirculating zones, roughly 20 times the residence time of the fluid in the inlet core flow.

While stabilized flame structures can be sustained in the resonant incinerator, combustion instabilities frequently arise in which vortical structures coincident with the flame are periodically shed from the dump plane. The periodic heat addition that results from vortex shedding acts to supply energy to specific acoustic modes of the device and can affect waste destruction. These naturally occurring combustion instabilities are observed to play a significant role in the degree to which waste surrogates are destroyed in the device[21, 22, 23]. Surrogates of hazardous wastes are frequently used in fundamental as well as prototype testing of thermal destruction devices in order to minimize hazards to research participants while studying the destruction of species which behave thermodynamically as real hazardous wastes do.

In specific cases, high frequency acoustic modes in the resonant incinerator can produce DREs of the waste surrogate sulfur hexafluoride (SF_6) which are as high as "eight nines" or 99.999999%, four orders of magnitude greater than the EPA minimum requirement of 99.99%. An alternative waste surrogate, acetonitrile, has also been tested in the resonant incinerator[23]; operating conditions are identified here in which acetonitrile destruction is at least two orders of magnitude above the EPA minimum. The behavior of the resonant incinerator has also been examined under externally forced acoustical conditions, where the external forcing acts as a means of active control. Active acoustical forcing can have a pronounced effect on waste destruction[24]; for example, external forcing at specific natural resonances can produce an increase in SF_6 DREs by four to six orders of magnitude.

As noted in the previous section, incineration systems in general have a number of drawbacks for potential implementation on board military vessels; the resonant incinerator has similar disadvantages. Its extremely high efficiencies, however, suggest that the technology holds some promise for development as a compact, highly efficient thermal destruction device for fluid wastes.

8 Observations and Conclusions

From the foregoing discussion, it is clear that there are a number of "off-the-shelf" as well as new, advanced technologies that could be exploited to solve many of the military's shipboard waste problems. It is also clear, at least to this author, that some of the technologies under consideration may be excellent choices for land-based waste treatment but may be inappropriate alternatives for shipboard waste treatment. An elaboration on some of these observations follows.

In terms of commercially available incineration, compacting, and filtering technologies, it appears that the Coast Guard has found a number of appropriate systems. These off-the-shelf technologies can in many cases be implemented more quickly and cheaply than developing specialized treatment units, and while retrofitting legions of existing vessels can be quite expensive and time consuming, the scale at which the systems need to be applied (in terms of numbers of vessels) probably makes this option quite reasonable for the Coast Guard. Ongoing testing of these off-the-shelf treatment systems, e.g., that of the Kvaerner Golar waste incinerator in Alaska, should be carefully followed, especially by Navy waste treatment specialists, for possible implementation on smaller Navy vessels.

The scale at which waste systems need to be applied in the Navy, especially for large carriers, submarines, and assault ships, perhaps makes the option of commercially available systems less viable. Many of these vessels are away from port for months at a time, and the number of personnel present is often in the thousands. As noted above, the Navy is considering two possible overall approaches to its shipboard waste problems: 1) development of specialized systems which each can treat specific classes of waste, and 2) development of one or two larger scale treatment systems for virtually all classes of shipboard wastes. For either approach, it is clear that power requirements, size, IR signature, EMI, throughput, and potential operator expertise need to be given strong consideration, in addition to process efficiencies and emissions.

In making any sort of decision among various treatment technologies, it seems clear that one ought to consider the military vessel as a **technological system as a whole**, and to determine how various alternatives fit within that system. The propulsion system (e.g., the gas turbine engine) which propels the ship can provide a source of power which could be extracted to run a waste treatment device, but this power is limited and is of course already consumed by various other components of the existing shipboard system. Thus, for waste treatment systems to fit appropriately within the **shipboard system**, they should not have sizable energy requirements. Filtration systems, pulpers, and possibly plastics processors apparently can work appropriately within the shipboard system, as they are currently implemented without significant energy cost. The necessity for an external generator to provide power for a waste treatment device disrupts the existing shipboard system, and introduces additional complexities that need to be dealt with, e.g., the need to add air pollution control systems to reduce diesel generator air emissions.

On the other hand, fuel for the ship's propulsion system is *already* a part of the existing shipboard system, so that if a waste treatment system could utilize this same fuel during

waste destruction, the shipboard system as a whole would not be significantly impacted. Thus fuel-fired systems such as incinerations could be a more logical thermal treatment system, say, than some of the more energy-intensive technologies discussed above. Again, other issues which impact the military mission of the shipboard system must be accounted for in deciding among treatment technologies, e.g., IR signature, EMI, and operator expertise.

As for the issue of one overall treatment technology vs. specialized technologies, it seems that very reliable backup systems would be required for the first option. If there is significant "down time" that occurs for the single large system, whether due to routine maintenance or system upset, nearly every waste stream generated on board could be affected and would need to be stored, at least temporarily, until the problem is solved. If multiple, diverse technologies are used, there are perhaps more complex failure pathways that could occur, but they would be much less likely to affect all waste streams at once. This observation, coupled with the investment that has been made already in plastics processors, shredders, pulpers, and liquid waste filters, indicates that specialized treatment systems could work quite well within the "system" that makes up a military vessel. Clearly, shipboard waste treatment and disposal problems are extremely complicated for the U.S. military, requiring objective technical evaluation in the future.

9 Postscript

After nearly completing this document, the author was made aware of a very recent (late October, 1995) series of meetings held at The Hague, Netherlands, of the NATO Special Working Group 12 on Maritime Environmental Protection. Among the presentations given was one dealing with a draft of a pre-feasibility study for the environmentally sound ship of the future, conducted by NIAG SG/50. A number of alternative shipboard waste treatment technologies were considered in this assessment, including many noted in the present study, and preliminary findings concerning differences in ease of retrofit, utilities requirements, reliability, military impact, etc. were presented. An example of a draft "scoring sheet" for plastics waste treatment alternatives is provided in Table 2. Thermal plasma, pyrolysis/combustion, and different shredding and compacting systems are compared here. As noted above, there are vast differences among these technologies, especially in terms of size, power required, cost, throughput, and operability. Studies such as this are precisely what are needed in order to make reasoned, technically-based decisions on future shipboard waste treatment systems.

Table 2. Value analysis "scoring sheet" for different potential shipboard waste treatment systems, with plastic waste destruction as an example.

VALUE ANALYSIS SCORING SHEET
QUANTITATIVE MEASURES BY TECHNOLOGY

WASTE: Plastics

CRITERIA	INITIAL SCORES				
	Thermal Plasma Pact-2	Thermal Plasma Pact-5	Pyrolysis/ Combustion	Shredding and Compaction	Shredding and Compaction Under Heat
EASE OF RETROFIT					
Weight	3500	10500	31000	2500	2500
Center of gravity	No Restrict	No Restrict	No Restrict	No Restrict	No Restrict
Utilities	Extensive	Extensive	Medium	Low	Low
Volume	Large	Large	Large	Medium	Medium
Personnel	45	45	30	14	15
Cost	\$ 4,000,000	\$ 8,000,000	\$ 1,000,000	\$ 85,000	\$ 90,000
EQT CHARACTERISTICS					
Weight	3500	10500	31000	2500	2500
Volume	70	200	60	4	3
Disembarkability	no	no	no	yes	yes
Utilities					
Power (kw)	200	750	50	6.6	9
Air Pres (atm)	5	5	1	0	0
Air Flow (cu m/min)	170	2100	3000	0	0
Steam	0	0	0	0	0
Fresh Water (kg/hr)	20	100	0	0	0
Salt Water (kg/hr)	2200	100000	0	0	0
Chilled Water	0	0	0	0	0
Operating limits	Wide	Wide	Wide	Wide	Wide
Consumables	Low	Low	Low	Low	Low
Secondary waste streams	Low	Low	Low	Medium	Medium
Equipment life	Medium	Medium	Medium	Medium	Medium
Upgrade potential	Low	Low	Low	Low	Low
Modularity	No	No	No	Yes	Yes
MILITARY ISSUES					
Signatures	Medium	Medium	Medium	Low	Low
Mission profile	Low	Low	Low	Medium	Medium
Shock resistance	No	No	No	Yes	Yes

(Continued)

Table 2 (Continued)
VALUE ANALYSIS SCORING SHEET
QUANTITATIVE MEASURES BY TECHNOLOGY
PLASTICS EQUIPMENT EXAMPLE

CRITERIA	INITIAL SCORES				
	Thermal Plasma Pact-2	Thermal Plasma Pact-5	Pyrolysis/ Combustion	Shredding and Compaction	Shredding and Compaction Under Heat
ECONOMICS					
Capital purchase	\$ 4,000,000	\$ 8,000,000	\$ 1,000,000	\$ 85,000	\$ 90,000
Consumables	\$ 9,000	\$ 10,000	\$ 1,000	\$ 1,000	\$ 1,000
Installation costs	High	High	High	Low	Low
Maintenance	Medium	Medium	Medium	Low	Low
Ship to shore/ship to ship interface	Low	Low	Low	Medium	Medium
Shore based disposal costs	Low	Low	Low	Medium	Medium
Development costs	3/10	3/10	6/10	10/10	10/10
Equipment life	Medium	Medium	Medium	Medium	Medium
Ship's personnel cost	High	High	Medium	Low	Low
Scale factor	No	No	No	Yes	Yes
RISK ANALYSIS					
Technical/Cost/Performance	3/10	3/10	6/10	10/10	10/10
PERFORMANCE					
Consistent Quality	Medium	Medium	Medium	High	High
Reliability (MTBF)	500	1000	500	500	900
Compliance with Standards	Yes	Yes	Yes	Yes	Yes
OPERABILITY					
Hazards	Low	Low	Low	Low	Low
Operation	High	High	Medium	Low	Low
Maintenance	Medium	Medium	Medium	Low	Low
Throughput (kg/hr)	10	125	220	1.25	1.25

10 Acknowledgements

I wish to thank numerous individuals who have taken the time to discuss these issues with me in the course of my study, and whose contributions to my understanding of shipboard waste problems have been substantial. These include Admiral (Ret.) Mike McCaffree, Admiral (Ret.) Bob Hilton, Joel Tumarkin, Tom Morehouse, and Susan Clark of the Institute for Defense Analyses; LDCR John Hautala and Craig Corl of the U.S. Coast Guard; Tony Rodriguez and Gene Keating of the Naval Surface Warfare Center Carderock Division; Larry Koss of the Office of the Chief of Naval Operations; Captain Steve Evans of the Naval Sea Systems Command; Lou Rosocha, Steve Buelow, and Jacek Dziewinski of Los Alamos National Laboratory, and Randy Seeker and Jerry Cole of Energy and Environmental Research, Inc.

11 References

References

- [1] Corl, C. A., "International Environmental Policy: MARPOL 73/78 Annex V and Shipboard Solid Waste Handling", U.S. Coast Guard Report, December 12, 1994.
- [2] Koss, L., "SWG/12 Maritime Environmental Protection: An Overview", talk presented at the NATO Naval Armaments Special Working Group 12 Meeting, The Hague, Netherlands, October, 1995.
- [3] McCaffree, M., Internal IDA Report on Shipboard Solid Waste and Oily Liquids, 1995.
- [4] TR-2008-ENV, "Calendar Year 1992 Hazardous Waste Minimization Report", Naval Facilities Engineering Service Center, Port Hueneme, CA, November, 1993.
- [5] Seeker, W. R., "Engineering Support to Naval Air Warfare Center on Compact, Closed-Loop Controlled Waste Incinerator", presentation at SERDP meeting, Energy and Environmental Research Corporation, June 13, 1995.
- [6] Brunner, C. R., *Handbook of Incineration Systems*, McGraw-Hill, 1991.
- [7] Morehouse, Thomas, "Visit Report for Trip to Carderock Laboratory for Project 1249, NATO Pollution Prevention", Institute for Defense Analyses, September 28, 1994.
- [8] Almeda, SN Anna, "Today's Cutter of Tomorrow", *Commandant's Bulletin*, p. 13, February, 1995.
- [9] Manahan, Stanley E., *Hazardous Waste Chemistry, Toxicology, and Treatment*, Lewis Publishers, 1990.

- [10] Alig, C., "Membrane Treatment of Bilge Oily Wastewater Update", talk presented at the NATO Naval Armaments Special Working Group 12 Meeting, The Hague, Netherlands, October, 1995.
- [11] Buelow, S. J., Robinson, J., and Foy, B., "Mixed Waste Focus Area Team Briefing: Hydrothermal Processing of Hazardous Wastes", August 2, 1995.
- [12] Dziewinski, J., Zawodzinski, C., and Smith, W. H., "Electrochemical Treatment of Mixed (Hazardous and Radioactive) Wastes", presented at the 5th International Conference on Radioactive Waste Management and Environmental Remediation, Berlin, Germany, September 3-8, 1995.
- [13] Dziewinski, J., Marczak, S., Smith, W. H., and Purdy, G., "Electrochemical Treatment of Waste", presented at the I&EC Special Symposium, American Chemical Society Meeting, Atlanta, GA, September 17-20, 1995.
- [14] Dziewinski, J., Marczak, S., Smith, W. H., and Nuttall, E., "Electrochemical Treatment of Mixed and Hazardous Waste", to be presented at the Scientific Basis for Nuclear Waste Management XIX, Materials Research Symposia, Fall Meeting, Boston, MA, November 27-December 1, 1995.
- [15] Sartwell, B., Talk on plasma arc pyrolysis presented at the ONR Environmentally Sound Ships meeting, Crystal City, VA, March 28, 1995.
- [16] Rosocha, L. A., "Processing of Hazardous Chemicals Using Silent Electrical Discharge Plasmas", LA-UR-94-4278.
- [17] Rosocha, L. A., Anderson, G. K., Coogan, J.J., Kang, M., Tennant, R. A., and Wantuck, P. J., "Two-Stage Thermal/Nonthermal Waste Treatment Process", LA-UR-93-1538.
- [18] Krieger, J., "Hazardous Waste Management Database Starts to Take Shape", *C&EN*, pp. 19-21, February 6, 1989.
- [19] Oppelt, E. T., "Incineration of Hazardous Waste: A Critical Review", *JAPCA*, 37, pp. 558-556, 1987.
- [20] SERDP 1994 Annual Report and Five-Year(1994-1998) Strategic Investment Plan, September, 1994.
- [21] Smith, O. I., Marchant, R., Willis, J., Lee, L. M., Logan, P. and Karagozian, A. R., "Incineration of Surrogate Wastes in a Low-Speed Dump Combustor", *Comb. Sci. and Tech.*, 74, pp. 199-210, 1990.
- [22] Marchant, R., Hepler, W., Smith, O. I., Willis, J., Cadou, C., Logan, P., and Karagozian, A. R., "Development of a Two-Dimensional Dump Combustor for the Incineration of Hazardous Wastes", *Comb. Sci. and Tech.*, 82, pp. 1-12, 1992.

- [23] Willis, J. W., Cadou, C., Mitchell, M., Karagozian, A. R., and Smith, O. I., "Destruction of Liquid and Gaseous Waste Surrogates in an Acoustically Excited Dump Combustor", *Comb. and Flame*, 99, pp. 280-287, 1994.
- [24] Pont, G., Willis, J. W., Karagozian, A. R., and Smith, O. I., "Effects of External Acoustical Excitation on Waste Surrogate Destruction in a Resonant Incinerator", Paper 95S-071, Proceedings of the Western States/Central States Section/The Combustion Institute Spring Meeting, April, 1995.

**G. NANOCRYSTALLINE MATERIALS FOR ABSORBING
MICROWAVE AND INFRARED RADIATION**

Brent T. Fultz

**Department of Materials Science and Engineering
California Institute of Technology
Pasadena, CA**

S. Lance Cooper

**Department of Physics
University of Illinois
Urbana-Champaign, IL**

SUMMARY

We argue that certain nanocrystalline materials will be powerful absorbers of microwave and far infrared radiation. The materials should have crystallite sizes of 10 nm or smaller, and these crystallites should be interconnected in ways that promote mechanical and electrical interactions between them. Our paper is constructed as follows:

1. *Some Principles of Radar.* This section reviews a few basic principles of how a reduction in radar cross-section influences the range of detectability by, and survivability against, a threat system. We present simple models without and with the use of electronic radar countermeasures.
2. *Design of RAM and RAS.* This section describes electrical features of non-reflective radar absorbing materials or radar absorbing structures. These considerations are basic to how an electrical engineer would design RAM and RAS systems. These concepts do not involve properties intrinsic to the material, however, so these systems design concepts are not central to the present paper.
3. *Nanostructured Materials.* We review recent methods for the synthesis of nanostructured materials and describe some nanophase microstructures. Magnetic properties of nanostructured materials, relevant to magnetic RAM, are discussed.
4. *An Idea.* This section presents a new concept for how mechanical vibrations of nanostructured materials can be tuned to lie in microwave and far-infrared frequencies. Broadband frequency spectra are calculated with a lattice dynamics method, and choices of crystallite masses and intercrystallite force constants are discussed.
5. *Evidence for Nanostructural Vibrations.* We present our experimental results from inelastic neutron scattering experiments on nanocrystalline Ni_3Al and nanocrystalline Fe that show an enhanced vibrational density of states at energies that extend into the infrared and microwave frequencies.
6. *Electromagnetic Absorption by Nanostructured Materials.* This section reviews the physics literature on previous work with microwave- and infrared-active small particles and amorphous materials. Mechanisms for coupling the electromagnetic energy into the material and some experimental results are described. This section ends with some speculative but novel ideas about how electromagnetic absorption could be controlled or modulated in time.

7. *Preliminary Results.* We present our experimental results from a first effort at measuring the microwave absorption spectra of nanophase Fe powder, with a comparison to a similar powder having larger crystallites. These results are encouraging, but they are preliminary and require further verification.

I. SOME PRINCIPLES OF RADAR

A. BASIC CONSIDERATIONS

In the simplest problem where an isolated conducting sphere is placed far from the microwave transmitter/receiver, the detected power from the microwave reflection, P_r , depends on the power of the microwave transmitter, P_t , as:

$$P_r = \frac{P_t G^2 \lambda^2 \sigma}{(4\pi)^3 R^4} , \quad (1)$$

where G is the antenna gain (a dimensionless number that appears squared because the same antenna is used for the transmission and reception), and λ is the microwave wavelength (in the same units as σ and R). Here R is the distance between the sphere and the antenna. R appears to the inverse fourth power because the intensity of the outgoing wave at the sphere follows a R^{-2} law, as does the wave traveling from the sphere back to the receiver. The variable σ is the microwave cross section. With the $(4\pi)^3$ normalization in Eq. 1, a sphere of diameter 1 m has a radar cross section of 1 m². Not all of the surface of the sphere contributes equally to the reflected signal, with the center of the sphere being more important, especially at higher frequencies (and hence the factor λ^2 in Eq. 1).

By calibrating with a reference sphere, radar cross sections can be determined in a test range by measuring the reflected power. It is found that the shape of the reflecting surface has an enormous effect on the reflected power, so different shapes have strongly different microwave cross sections, σ . For 3 cm waves, for example, very strong reflections are obtained from a square plate of 1 m $\times$ 1 m, giving it an impressively large cross section of 14,000 m². On the other hand, a square plate of 0.1 m $\times$ 0.1 m, a reduction in area of a factor of 100, does not have a radar cross section for 3 cm waves of 140 m², but is instead about 1 m². The smaller square plate has a size comparable to the microwave wavelength, and diffraction edge effects are significant.

All services of the military can benefit from reductions in the radar cross section (RCS), σ , for pieces of hardware. This paper does not intend to describe specific ways that low observability (LO) provides tactical advantages. Here we merely illustrate the advantage of LO for an oversimplified case of an aircraft against a threat from a radar-guided surface-

to-air missile (SAM). We begin by assuming that the reflected power at the radar receiver must be greater than P_{rmin} for detection. From Eq. 1, the minimum range that the aircraft or ship can remain undetected is then:

$$P_{\text{rmin}} = \frac{P_t G^2 \lambda^2 \sigma}{(4\pi)^3 R_{\text{min}}^4} \quad (2)$$

$$R_{\text{min}} = \sqrt[4]{\frac{P_t G^2 \lambda^2}{(4\pi)^3 P_{\text{rmin}}}} \sqrt[4]{\sigma} \quad (3)$$

The $\sqrt[4]{\sigma}$ may seem disappointing. This means that, for example, a reduction in RCS of a factor of 625 provides a decrease of only a factor of 5 in R_{min} , the minimum range that the aircraft can remain undetected.

B. AIRCRAFT SURVIVABILITY WITH SIMPLE CONSIDERATIONS

A factor of 5 reduction in R_{min} could be significant, however. Detection and tracking are not the same problem for a SAM battery. The lock-on required before a missile launch will occur some time after detection. We model this delay as a reaction time, t_r , that is required before a missile can be launched successfully. Figure 1 illustrates how a reduction in RCS by a factor of 5 affects survivability. Detection occurs at time t_0 . We assume that the aircraft is detectable for a time of $2R_{\text{min}}/v$, where v is its velocity. If the reduction in R_{min} causes the time window of detectability to fall below the reaction time of the SAM battery, the aircraft will survive.

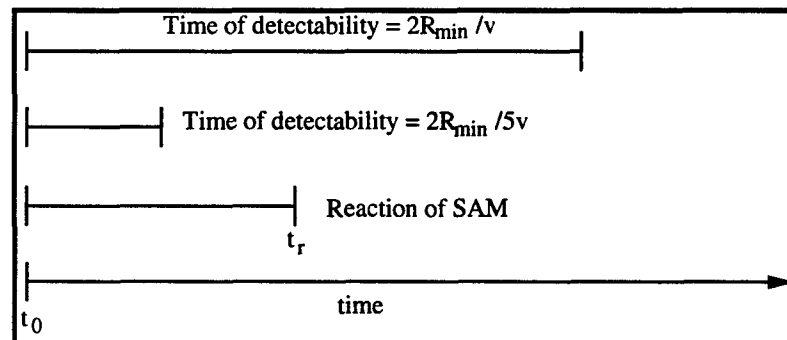


Figure 1. Some Characteristic Times Where Aircraft Is Detectable (top), Detectable With a Reduction in RCS of a Factor of 5 (middle), and the Reaction Time for a Sam Battery (bottom)

If the aircraft (or other hardware) need not approach closely the threat radar system, a reduction in RCS may make it easier to find a path where the aircraft remains entirely

invisible. When adjacent radar sites have overlapping areas of coverage, with a low RCS it may be possible to travel undetected between the radar sites.

C. AIRCRAFT SURVIVABILITY WITH ELECTRONIC COUNTERMEASURES

A low RCS can be even more advantageous when electronic countermeasures (ECM) are in use. Any full analysis of the benefits of low RCS in conjunction with ECM systems is both highly complex and highly classified, so we confine ourselves to a simplified illustration of principle. All ECM techniques fail below some distance defined as the "burn-through radius" (R_{BT}). Assume that R_{BT} decreases in proportion to the power from the ECM system at the threat radar receiver, P_{spoof} . For a fixed power of the ECM system, P_{spoof} decreases as R^{-2} , whereas the useful detection signal, P_r , decreases as R^{-4} as described by Eq. 1. We expect burn-through when P_r exceeds some fraction of P_{spoof} . This condition is:

$$P_r \propto P_{spoof}, \quad (4)$$

$$\frac{P_t G^2 \lambda^2 \sigma}{(4\pi)^3 R_{BT}^4} \propto \frac{\kappa}{R_{BT}^2}, \quad (5)$$

$$R_{BT} \propto \sqrt[2]{\frac{P_t G^2 \lambda^2}{(4\pi)^3 \kappa}} \sqrt[2]{\sigma}. \quad (6)$$

The burn-through radius depends on the square-root of σ , rather than the fourth root as in Eq. 1 for the problem without ECM. For our example of a decrease in σ by a factor of 625, we find a decrease in the burn-through radius of a factor of 25, not the smaller factor of 5 for the detection radius without ECM.

It is sometimes argued that a reduction in RCS can permit a reduction in the size and complexity of the ECM suite. This can be particularly important for aircraft. For example, the ECM suite of the B1 bomber comprises 108 units, weighing a total of 5,200 pounds, with a peak operating power of 120 kW [1].

II. DESIGNS OF RAM AND RAS

A. GENERAL FEATURES OF RAM AND RAS

There have been numerous approaches to the design of radar absorbing materials (RAM) and radar absorbing structures (RAS), ranging from the ridiculous to the sublime. A fine example of the former was the proposal to coat all surfaces of an aircraft with ^{210}Po [2]. With sufficient radioactivity, a plasma would be generated around the aircraft that would allegedly impede microwave reflection.

The design of RAM and RAS must address two criteria. First, there must be no impedance mismatch at the surface of the material. Any impedance mismatch will cause reflected energy, even if the material has good radar absorbing properties. The second design criterion is that the RAM and RAS must be able to convert microwave energy into heat.¹

The RAM and especially the RAS are integral to the airframe. They will likely see service for many years. Although it is often possible to obtain good microwave absorption in a narrow band of frequencies tuned for a particular threat, this approach is considered shortsighted; the nature of the threat is likely to change over the lifetime of the airframe. Broadband absorbers are therefore the best choice.

Two types of radar absorbing materials are allowed by electromagnetics: a lossy dielectric material, and a lossy magnetic material. (Combinations of these two are also possible.) Dielectric and magnetic materials are parameterized by their complex permittivities, $\epsilon = \epsilon' + i\epsilon''$, and permeabilities, $\mu = \mu' + i\mu''$, respectively. Lossy materials are characterized by having large complex components, ϵ'' or μ'' . The condition for impedance matching, however, often requires that the lossy material be embedded as particles in a matrix having different μ or ϵ . This requirement may elevate the importance of the ratios ϵ''/ϵ' or μ''/μ' .

In shielding a conductive surface, there are different requirements for the shapes of dielectric or magnetic RAM. This difference is based on the fact that the electric field is zero

¹ If all materials in the airplane (including the pilot) had the free space impedance of 377Ω , the airplane would be transparent to microwave energy. This ideal is impractical, primarily because conductive metal components (0Ω impedance) are used in many systems in the aircraft.

at a conductive surface, whereas the magnetic field is a maximum. To ensure maximum absorption, a broadband dielectric RAM should extend above the conductive surface to a distance of at least one-fourth of the longest wavelength to be absorbed. For this reason, dielectric RAM is too thick to be used as a surface coating but is often used as an integral part of the airframe structure, i.e., as RAS. On the other hand, magnetic RAM will work well as a thin coating over a conductive surface.

B. MAGNETIC RAM

This paper is directed primarily to the concept of a magnetic RAM (although the basic idea of Section IV could be modified for dielectric RAM). Magnetic RAM materials are much like the "iron ball" coating that was used on the Lockheed SR-71 [3, 4]. The sizes of the iron particles are chosen to match approximately the skin depth of iron at microwave frequencies. If the particles have a high magnetic permeability at microwave frequencies, the induced $\partial B / \partial t$ is large. Large eddy currents will flow in the particle, especially if it has a low resistivity. The ensuing I^2R losses will dissipate the microwave energy in the form of heat. The magnetic fields from the microwave radiation are small (of order mOe for a narrow beam radar of 10^6 W at distances of tens of km), so it is often important for the RAM material to be magnetically "soft" at microwave frequencies. A low coercive field and high permeability are expected to be desirable for magnetic RAM material.

To make a magnetic RAM coating, the magnetic particles are embedded in a matrix that provides for an appropriate impedance for the composite RAM coating. In general, the magnetic RAM coating does not contribute to the mechanical strength of the airframe but does add weight. Many types of magnetic RAM materials are based on iron and iron compounds such as ferrites. The added weight may limit the thickness of the coating and its efficiency.

This paper does not provide a complete design for a RAM coating. Issues of dielectric permittivity and impedance matching are ignored. Instead this paper describes a new idea that should increase μ'' and allow control over the ratio μ''/μ' . This idea is based on a new phenomenon that should provide a high density of vibrational modes at GHz frequencies in nanostructured materials. The idea should be applicable to many classes of materials that are already in service as magnetic RAM but is different in principle from eddy current losses. The idea is to promote efficient photon-to-phonon conversion by creating nanostructures that have a high density of vibrational modes at microwave and far-infrared

frequencies.² If the nanostructure distorts as its magnetization changes in the presence of microwave radiation, dissipation takes place by excitation of the vibrational modes of the nanostructure. These distortions could be driven by magnetostriction or by changes in dipole coupling between the nanocrystallites, for example. More about the physics of loss mechanisms is presented in Section VI.C. Section IV discusses the design of nanocrystalline materials with a high density of vibrational modes in the 10-100 GHz range.

² The mechanism for energy loss could be like that termed the "phonon effect," which has been proposed as a loss mechanism at low frequencies but is poorly understood [5].

III. NANOSTRUCTURED MATERIALS

A. RESEARCH DIRECTIONS IN NANOSTRUCTURED MATERIALS

In the last 5 years there has been a strong interest in "nanostructured materials," materials having internal structures with spatial scales on the order of 10 nm or less. Research is active in nanostructured polymers, metals, and semiconductors. Examples include control of structures in diblock copolymers [6, 7], magnetic properties of nanocrystalline ferromagnets [8, 9], and optical properties of GaAs quantum dots [10]. Some of this interest arises from the challenge of controlling materials structures on small spatial scales and characterizing their microstructures with new methods of microscopy, diffraction, and spectroscopy.

More interest originates with the unconventional properties of nanostructured materials. There are two classes of microstructure-properties relationships. The first is termed a "confinement effect," and the second an "interface effect." An example of a confinement effect is the shift in frequency of optical absorption that occurs in semiconductors when the crystallite size is comparable to the size of the electron-hole exciton. An alleged example of the second interface effect has been reported in ceramic TiO_2 , which has a surprisingly large ductility at low temperature attributed to atom movements in its copious grain boundaries [11].

B. SYNTHESIS OF NANOSTRUCTURED MATERIALS

There are many ways to synthesize nanostructured materials, including methods of wet chemistry, sol-gel reactions, physical vapor deposition, self-assembly in polymer blends, and highly controlled methods such as molecular beam epitaxy. Here we describe two methods that are appropriate for synthesizing nanostructured, ferromagnetic Fe-based alloys. The first method, developed by Buhrman [12], but brought to prominence by Gleiter [13] and Siegel [14], is known as the "gas-condensation method." A metal or alloy is heated above its melting temperature in an inert gas environment. After evaporation, the metal vapor is cooled rapidly by the gas and solidifies into small crystalline particles of sizes typically less than 10 nm. These particles undergo Brownian motion in the gas and coalesce into loosely bound structures as shown in Figure 2.

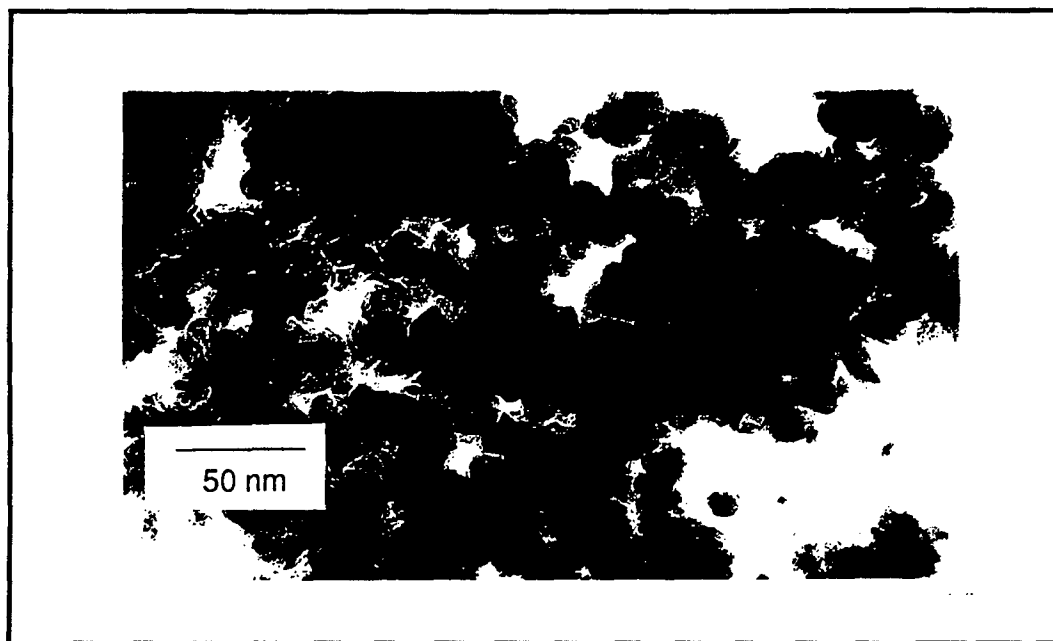


Figure 2. Nanocrystalline Fe Agglomerates Prepared by the Gas Condensation Method

[Image is a transmission electron micrograph of agglomerates collected on a holey carbon microscope grid and examined in a Philips EM 301 transmission electron microscope operated in bright field mode at 100 kV.]

Another method of making nanostructured materials is “high energy ball milling” [15-17]. Metal powders are put together with several hardened steel balls, sealed in a hardened steel vial, and shaken violently by a mechanical device similar to a paint shaker. The mechanism of grain refinement during mechanical alloying is not fully understood, but it involves the generation of crystalline defects and the coalescence of defects into structures that provide misorientation between different regions of crystallite. With several hours of energetic milling, these defect structures take the form of low- and high-angle grain boundaries.

In nanocrystalline materials prepared by high energy ball milling, the shapes of the crystallites are irregular and vary considerably in their size. Although there are few crystallites with sizes below a lower threshold (about 2 nm for the alloys considered here), above this threshold the size distribution is often like that of a decaying exponential. For as-milled material having an average grain size of 10 nm, for example, crystallites of sizes 3 nm and 30 nm are typically present. Figure 3 shows an example of a nanocrystalline material prepared by high energy ball milling. In comparison to the material synthesized by gas condensation, shown in Figure 2, the crystallites in the ball milled material are more angular, less uniform, and packed more densely.

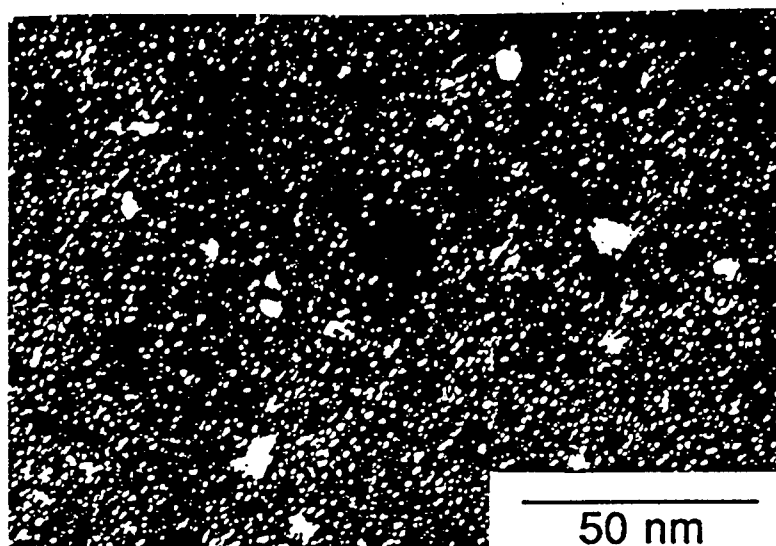


Figure 3. Alloy of Fe-Cr Prepared by High Energy Ball Milling

[Image was made with a Philips EM 301 transmission electron microscope operated in (110) dark field mode at 100 kV. Individual crystallites, suitably oriented, appear bright in this image.]

C. MAGNETIC PROPERTIES OF NANOSTRUCTURED MATERIALS

Magnetic properties of nanostructured metallic alloys are topics of much ongoing fundamental and applied research. An older, well-known phenomena is superparamagnetism, in which the alignment of the atomic magnetic moments in a small particle can be flipped easily with thermal energy. A newer phenomenon is superferromagnetism, where strongly aligned magnetic moments of one particle can couple to the magnetic moments of an adjacent particle. The exchange energy of these particle-particle interactions can be controlled, at least in principle, by tuning the interparticle interface. The tendency for a common magnetic alignment of all particles would be expected to be strong, for example, if the particles were spaced closely together with a narrow interface between them.

Research in nanostructured magnetic materials with especially low coercive fields has been stimulated by the commercial product "Finemet," which comprises nanocrystalline bcc Fe-Si crystallites within an amorphous matrix [8, 9]. Its extremely soft magnetic properties are explained by the "random anisotropy model" [18]. This model predicts that the exchange interaction (A) will dominate the magnetocrystalline anisotropy (K_1) for materials having grain sizes smaller than the ferromagnetic exchange length, $L_{ex} = \sqrt{A/K_1}$. In essence, when a polycrystalline grain size is smaller than the width of a magnetic domain wall, there will be little tendency for crystalline orientations to affect the energy of magnetic orientation in the material. The coercive field, H_c , will be low.

Finemet is produced commercially by first preparing an amorphous metal, which itself has a low coercive field. The coercive field and saturation magnetization are improved further when the amorphous material is heat-treated, and a dense precipitation of nanocrystalline bcc Fe-Si crystallites forms within the amorphous matrix [8,9]. One of the authors is investigating the synthesis of soft magnetic materials by high energy ball milling. The crystallite sizes are frequently in the appropriate range of about 10 nm. Unfortunately, high internal stresses are present in nanocrystalline materials prepared by high energy ball milling. These internal stresses cause the as-milled materials to have coercive fields of 30 Oe or so. Annealing the as-milled powders at low temperatures provides the needed stress relief, and a reduction of coercive field. Unfortunately, annealing at high temperatures causes the crystals to grow in size. The trick is to design a material that will undergo sufficient stress relief before it will undergo significant grain growth. We have performed such studies on Fe-Ge [19], Fe-Si [20], Fe-Si-Nb [21, 22], and Fe-Al-Ge. The motivations for these choices of materials and their magnetic properties are provided in the references.

In short, our best results so far were obtained with Fe-Al-Ge. The alloy composition was close to that required for a magnetostriction coefficient of zero [23], so its coercive field should be particularly insensitive to internal stress. On our first attempt with this composition, we obtained a reduction of coercive field from about 30 Oe in the as-milled material to below 300 mOe in the annealed material (300 mOe was the limit of sensitivity of our equipment). With further annealing we observed grain growth and an increase in H_c . This work shows that soft magnetic properties can be achieved for nanocrystalline materials synthesized by routes other than that of Finemet.

IV. AN IDEA

A. CONCEPT

In solids, heat is the vibrations of atoms. In metals and alloys, the characteristic frequency of these vibrations is around 10^{12} - 10^{13} Hz. Knowing a "spring constant," k_{at} , between an atom and its first-nearest-neighbor atom, and the mass of the atom, m_{at} , the frequency, ω , can be calculated:

$$\omega = \sqrt{k_{at}/m_{at}} = \text{several THz.} \quad (7)$$

The core idea of this section is that nanostructured materials have crystallites that can serve as individual masses (heavy masses), and grain boundaries that can serve as springs (weak springs). It is not unreasonable to expect that the masses of the nanocrystals can be 10^5 atom masses, and the spring constants between crystallites will be weak, only a few times k_{at} . The characteristic frequency could then decrease considerably to become:

$$\omega = \sqrt{k_{at}/(10^5 m_{at})} = \text{several tens GHz.} \quad (8)$$

For loosely bound nanocrystallites, the characteristic vibrations will be in the microwave range, or perhaps in the far infrared. Mechanisms exist for coupling electromagnetic energy into these vibrations, and these are described in Section VI. Some of these mechanisms become more efficient when there is a high density of vibrational modes at the frequency of the electromagnetic wave. In Section IV.B, we show that the vibrational spectrum is more complicated than a single characteristic frequency. This works to advantage in making a broadband RAM. A broadband high density of vibrational modes at microwave or infrared frequencies is our design goal.

Although low frequency vibrational modes are generally expected in nanostructures, it is important to ask if there are more such modes per gram of nanostructured material than in a single crystal of the same material. All solids have vibrational modes at low frequencies. Consider for example a cube, with edge length ℓ , of a simple cubic crystal. Suppose the first

nearest-neighbor interatomic force constant is ϕ . Over a length ℓ , the interatomic springs add in series, so the force constant per atom in cross section is no longer ϕ , but rather ϕ/ℓ . For the same mass of nanostructured material to attain a higher density of states at low frequencies, it is therefore necessary for the interparticle spring constant to decrease more rapidly than the grain size (i.e., more rapidly than ϕ/ℓ).

Since there is a linear relationship between frequency and wavevector at low frequencies (manifested, for example, as a constant speed of sound), in three dimensions a solid has a density of vibrational modes that increases as the square of the frequency:

$$\rho(\omega) = \frac{V \omega^2}{2\pi^2 v^3} \quad (9)$$

Here V is the volume of the crystal, and v is the speed of sound, $v = \omega/k$. The dependence of $\rho(\omega)$ on the inverse cube of v is interesting. In the Debye model, which assumes a constant velocity of sound for all vibrational modes, the characteristic Debye frequency of the solid is proportional to the speed of sound. In other words, a linear decrease in the characteristic Debye frequency leads to a cubic increase in the number of low frequency vibrational modes. To obtain a high density of vibrational modes at “low” microwave frequencies, it is not really necessary to make a solid with a characteristic vibrational frequency in the microwave range. Just some suppression of the characteristic vibrational frequency of the solid will provide a strong enhancement of the density of modes at low frequencies.

B. LATTICE DYNAMICS OF NANOCRYSTALS

We used the Born-von Kármán model to calculate the vibrational dynamics of a coupled array of nanocrystals, and the results are described here. The problem was cast as an array of masses on a bcc lattice, and connected by springs to their first nearest neighbor masses only. This is depicted in Figure 4.

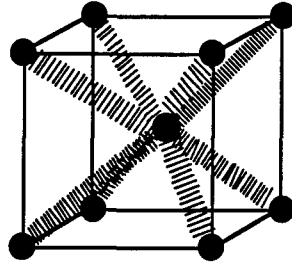


Figure 4. Unit Cell of the Mass-Spring System Modeled by the Born-von Kármán Calculations

The equations of motion of the masses are:

$$m_k \ddot{u}_{\alpha}^{(l)}(k) = - \sum_{l', k', \beta} \phi_{\alpha\beta} \begin{pmatrix} l & l' \\ k & k' \end{pmatrix} u_{\beta}^{(l')}(k') \quad (10)$$

where $u_{\alpha}^{(l)}(k)$ is the displacement vector of the α^{th} cartesian component (x, y, z) for the atom k in the l^{th} unit cell. $\phi_{\alpha\beta} \begin{pmatrix} l & l' \\ k & k' \end{pmatrix}$ is the force constant involving the displacement of the atom k and the neighboring atom k' . The phonon solution for $u_{\alpha}^{(l)}(k)$ is:

$$u_{\alpha}^{(l)}(k) = \frac{1}{\sqrt{m_k}} U_{\alpha}(klq) \exp(i(\mathbf{q} \cdot \mathbf{x}(l) - \omega(\mathbf{q})t)) \quad (11)$$

where $\mathbf{x}(l)$ is the shift of $u_{\alpha}^{(l)}(k)$ from its equilibrium position in a perfect lattice, and $\mathbf{q}$ is the wavevector of the vibrational mode. The frequency for this particular vibrational mode is $\omega(\mathbf{q})$. Substituting Eq. 11 in 10, and rearranging the equations of motion we have:

$$\omega^2 U_{\alpha}(klq) = \sum_{k', \beta} D_{\alpha\beta} \begin{pmatrix} \mathbf{q} \\ kk' \end{pmatrix} U_{\beta}(k'lq) \quad (12)$$

$$D_{\alpha\beta} \begin{pmatrix} \mathbf{q} \\ kk' \end{pmatrix} = \frac{1}{\sqrt{m_k m_{k'}}} \sum_l \phi_{\alpha\beta} \begin{pmatrix} 0 & l \\ i & j \end{pmatrix} \exp(-i \mathbf{q} \cdot \mathbf{x}(l)) \quad (13)$$

We have defined the dynamical matrix $\mathbf{D}$ in Eq. 13 in terms of the Fourier transform of the force constant matrix. Equation 12 provides an eigenvalue equation for the ω^2 , and these frequencies are obtained by diagonalizing the dynamical matrix for a particular choice of phonon wavevector, $\mathbf{q}$:

$$\|\mathbf{D}(\mathbf{q}) - \omega^2(\mathbf{q}) \delta_{ij}\| = 0 \quad (14)$$

To get the full vibrational spectrum, we must diagonalize $\mathbf{D}$ for all wavevectors within the first Brillouin zone of the crystal, and collect a histogram of vibrational frequencies. A characteristic vibrational spectrum is presented in Figure 5. The first nearest neighbor force constants used in calculating the spectrum in Figure 5 were chosen to match the phonon dispersion curves for bcc Fe, and the vibrational spectrum of Figure 5 is much the same as the phonon density of states of bcc Fe [24].

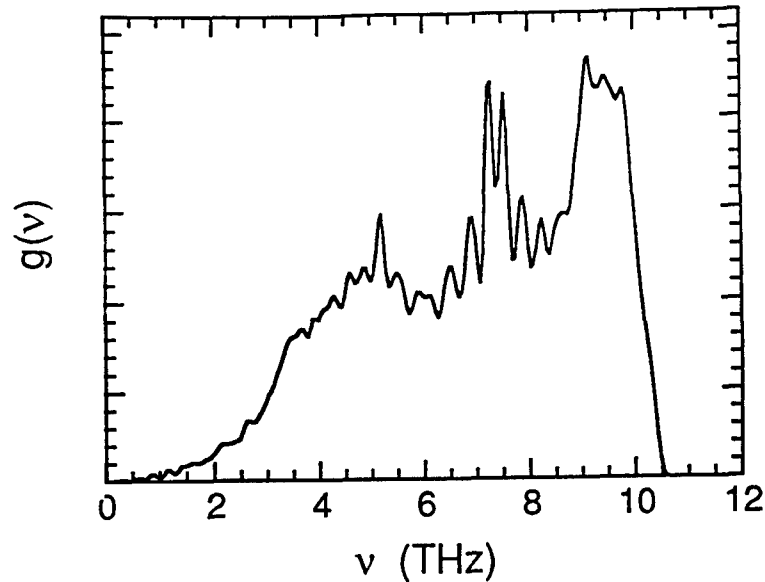


Figure 5. Vibrational Spectrum for a bcc Lattice of Mass 55.847 amu and Force Constants $C1XX = 25.0$ and $C1XY = 20.0$ N/m

To model the vibrational dynamics of nanocrystals, similar calculations were performed for cases where the masses were much heavier, and force constants were somewhat larger. Because we scaled all force constants by the same factor, the vibrational spectra all had the same shape, but were rescaled in energy. Results from these calculations are presented on a double logarithmic scale in Figure 6. The data set in the upper right hand corner is the same vibrational spectrum of pure Fe shown in Figure 5. A straight line with slope 2 is drawn through this data set. We expect this line to provide the density of modes in normal materials where the velocity of sound is constant at low frequencies. If nanostructured materials are to offer more vibrational modes at GHz frequencies than conventional materials, it is important for them to have a higher density of modes than the straight line in Figure 6 with slope 2. In favorable cases, we see that the vibrational spectrum of nanostructured materials lies two orders of magnitude above the line.

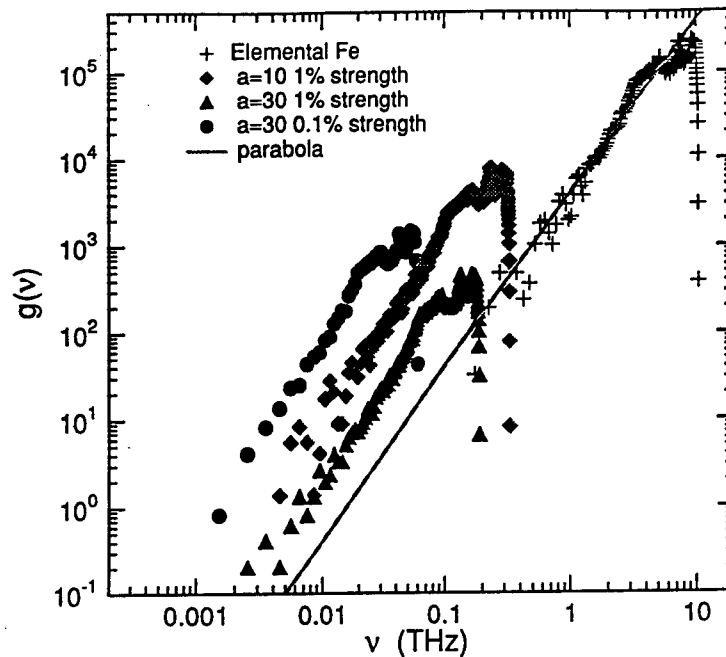


Figure 6. Results of Born-von Kármán Calculations for bcc Lattices with Various Combinations of Force Constants and Masses

[The force constants for each atom at the interface were the percentages of those given in Figure 5; the masses were a^3 times 55.847 amu.]

C. ASSUMPTIONS AND APPROXIMATIONS

There are two reasons to temper the optimism provided by the previous section. First, it was assumed that the modes of the nanostructure were decoupled from the modes of the crystallites themselves. This assumption is most successful when the two types of modes have very different frequencies. The strong elevation of some of the simulated results above the straight line in Figure 6 gives credence to this assumption. A second, riskier assumption is that the interparticle stiffnesses are low. This depends strongly on the nature of the intercrystallite interface. It is clear that this interface is highly disordered [13, 25] and probably has a low characteristic stiffness. It is also possible to tune the stiffness of this intercrystallite interface by the method of synthesis. However, it is not possible to predict *a priori* the precise stiffness of the interface. Determining the stiffness of these intercrystallite interfaces is perhaps best performed by experimental work to confirm the existence of low frequency vibrational modes of the nanostructure.

V. EVIDENCE OF NANOSTRUCTURAL VIBRATIONS

One of us has been studying the vibrational entropy of materials, and this work has involved inelastic neutron scattering measurements [26]. The intent of this work was to compare the bulk vibrational modes of disordered metallic alloys to the vibrational modes of ordered alloys of the same chemical composition. To prepare sufficient quantities of material for the neutron scattering experiments, we used high energy ball milling. The as-milled materials were indeed chemically disordered but they also had very small crystallites. This was unintentional and unwanted at the time but has subsequently provided an interesting observation. This section describes our experimental neutron scattering studies on Ni_3Al and Fe. In both cases the powders were studied in the as-milled state and again after an annealing treatment that induced crystallite growth.

Powders of Ni_3Al were made by mechanical alloying [27, 28]. Measured amounts of elemental nickel and aluminum powders were milled in a Spex 8000 mixer/mill with hardened steel vials and stainless steel balls and a ball-to-powder weight ratio of 2:1. With hexane added to the vial, nearly complete alloying occurred in 3-4 hours, and several batches of Ni_3Al powder were prepared by milling for 8 hours at room temperature. X-ray diffractometry was performed with an Inel CPS-120 diffractometer using $\text{Co K}\alpha$ radiation. The x-ray diffraction patterns of the as-milled Ni_3Al powders showed: (1) that the as-milled Ni_3Al was a chemically disordered fcc solid solution, and (2) that the as-milled Ni_3Al had crystallite sizes of 6-7 nm. Chemical ordering and grain growth were obtained by heating the powder at 450 °C for 10 hours [29-31]. X-ray diffractometry showed that the grain size after annealing was 20 nm or more, and a high degree of L_{12} chemical order had developed.

Nanocrystalline bcc Fe powders were also prepared in steel vials with a ball-to-powder weight ratio of 5:1 (without hexane surfactant) and milled for 12 hours. X-ray diffractometry showed that the grain size in the as-milled powders was about 11 nm. Grain growth was achieved by heating the powder at 500 °C for 1 hour. X-ray diffractometry showed that the grain size after annealing was 28 nm.

Samples of the as-milled and the annealed powders, each about 50 grams, were placed in thin-walled aluminum cans and mounted on the goniometer of the HB3 triple axis spectrometer at the High Flux Isotope Reactor at the Oak Ridge National Laboratory. The spectrometer was operated in constant-Q mode with the fixed final energy, E_f , being 14.8

meV. The energy loss spectra were made by scanning the incident energy from 14.8-64.8 meV. The neutron flux from the monochromator was monitored with a fission detector, which was used to control the counting time for each data point. The incident beam on the pyrolytic graphite monochromator crystal was collimated with 40' slits, and 40' slits were also used between the monochromator and the sample. Pyrolytic graphite filters placed after the sample were used to remove the $\lambda/2$ contamination. The filtered beam passed through 80' slits before the pyrolytic graphite analyzer crystal. Following the analyzer, 2° slits were used before the ^3He detector. With this arrangement, the energy resolution varied from about 2 meV at low energy transfer to 5 meV at 40 meV energy transfer. Four values of Q were chosen for each specimen, ranging from 3.23 to 4.23 $\text{\AA}^{-1}$.

Energy loss spectra from the ordered and disordered Ni_3Al are presented in Figure 7. Individual runs were highly reproducible, as shown by the two independent (but nearly coincident) data sets from the ordered alloy with $Q = 4.23 \text{ \AA}^{-1}$. The energy loss spectra from the material with $L1_2$ order shows a peak around 39 meV, and the measuring the change in intensity of vibrational modes at 39 meV was the main goal of this study on Ni_3Al [26].

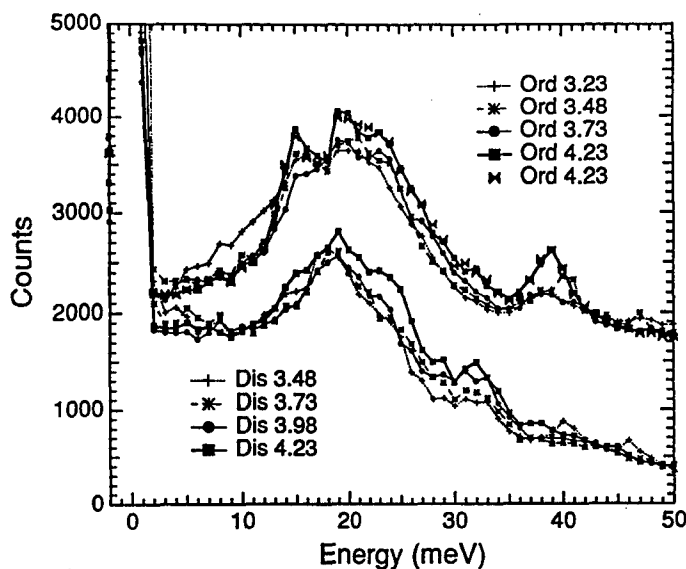


Figure 7. Raw Neutron Energy Loss Spectra for the Ordered and Disordered Ni_3Al Powders, Obtained at Four Values of Q
[Data for the ordered alloy are offset vertically by 1500 counts.]

The feature of interest in Figure 7 for the present paper is the stronger intensity from the disordered alloy at energy losses below 15 meV or so. This intensity could well originate with vibrational modes of the nanostructure; modes that are not found when the grain size has

increased after annealing. We point out that the spectra are not the density-of-states functions of the material. Converting the data to density-of-states functions requires, among other things, a correction for thermal occupancy. This correction factor will suppress considerably the intensity of the low energy modes. It will do so by the same factor for both specimens, however, so a comparison of the raw data in Figure 7 is appropriate.

A similar set of experiments was performed on the as-milled and annealed samples of bcc Fe. Data were obtained for two values of Q (4.0 and 4.6 $\text{\AA}^{-1}$), and the data were summed for presentation in Figure 8. Again, as for the case of as-milled and annealed Ni_3Al , there is enhanced absorption below 15-20 meV. The effect is not so large for Fe as for Ni_3Al , however. This could be attributable to the smaller crystallite size of the Ni_3Al or to stiffer interfaces in the Fe powders. (There is some controversy over the nature of the nanocrystalline structure in as-milled Fe, and there is some evidence that the as-milled Fe is better described as crystallites with defects rather than as individual nanocrystals [33].)

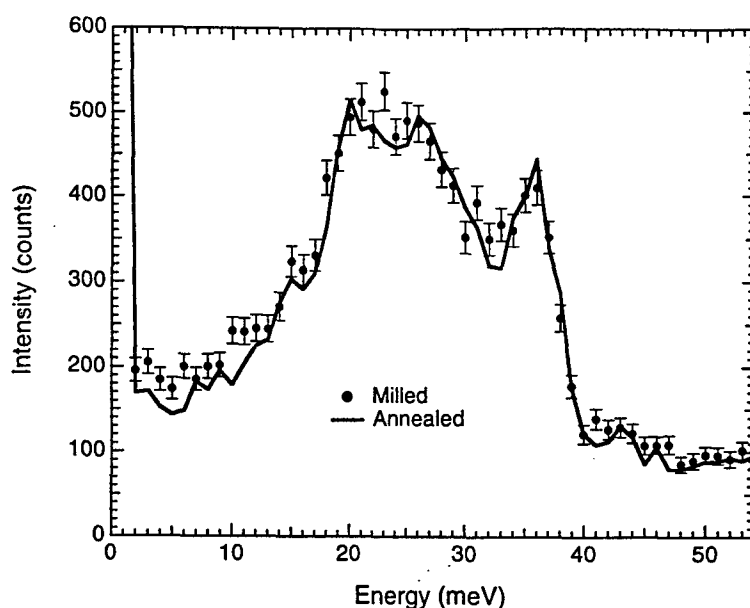


Figure 8. Raw Neutron Energy Loss Spectra for the As-Milled and Annealed Fe Powders, Each Summed for $Q = 4.0$ and $Q = 4.6 \text{ \AA}^{-1}$

Another problem that we considered is whether the low energy part of the spectrum could have originated with quasi-elastic scattering from hydrogen contamination of the specimen. Not much hydrogen would be needed to explain the present results, perhaps on the order of 1 atomic percent. We have performed an analysis of interstitial gases in the samples of Fe and Ni_3Al by heating these powders to 400 $^{\circ}\text{C}$ in a gas chromatograph coupled

to a mass spectrometer. To our surprise, more hydrogen was evolved from the specimens that were annealed than the specimens that were disordered. In other words, the differences in the phonon DOS at low energies cannot be explained by hydrogen contamination.

VI. ELECTROMAGNETIC ABSORPTION BY NANOSTRUCTURED MATERIALS

Nanophase absorbers with composite structures, comprising an optically active material embedded in a dielectric medium, are potentially useful not only because the absorption characteristics of these materials can, in principle, be carefully controlled, but also because composite absorbers can, in certain cases, have enhanced absorptivity across wide regions of the electromagnetic spectrum. For example, a number of experiments in the past 20 years have demonstrated that various materials having nanoscale (i.e., characteristic length scales $\ell \sim 10 - 1000 \text{ \AA}$) structural components, such as small metal particle clusters and amorphous materials, can manifest low frequency (microwave and far infrared) absorption coefficients that are enhanced relative to homogeneous dielectric or magnetic materials. Generally speaking, there appear to be two chief mechanisms through which such enhancements occur in these materials. First, the large far infrared (FIR) and microwave absorption coefficients in metal-particle clusters embedded in various dielectric media, $\alpha = 5 - 50 \text{ cm}^{-1}$ [34-36], are likely because of clustering of the small metal particles, which causes radiative dissipation through the induced eddy currents in the clusters; in this sense, nanocrystalline clusters are similar to conventional "circuit analog" absorbers, in which arrays of tiny RCL circuits provide a tunable band of low frequency electromagnetic absorption. Second, the inherent nanoscale disorder associated with amorphous materials results in low energy Debye-like and non-Debye-like vibrational modes. Light can couple to these modes, providing an additional source of low frequency electromagnetic absorption. The low frequency absorption coefficients of certain silicate and chalcogenide glasses, for example, are several times larger than those of corresponding crystalline phases [37]. Some features of microwave absorption with this second mechanism are similar to features expected from the vibrational modes of nanocrystalline materials, as described in Section IV.

Significantly, the potential technological utility of amorphous and composite materials derives not simply from their large low frequency absorption coefficients, but also from the ease with which the absorption coefficients of these materials can be engineered by controlling the cluster or domain size. Further, the multi-component nature of these composite structures, involving an optically active element and a dielectric host, may also provide a means by which the optical response of a protective "coating" can be remotely and

dynamically controlled. The final part of this section will discuss potentially useful design schemes for nanocrystalline microwave absorbers, which naturally exploit the properties of these materials, and will describe potential applications, which are particularly suited to composite absorber and reflective materials.

In the following, we elaborate on the underlying physics of microwave and far infrared absorption in metal particle clusters and amorphous dielectrics.

A. METAL PARTICLE CLUSTERS

It was discovered in the late 1970s that clusters of 60 - 400 Å Cu, Al, Ag, Au, Sn, or Pb particles embedded in a dielectric matrix exhibit large absorption coefficients ($\alpha = 5\text{-}50\text{ cm}^{-1}$) in the far infrared. The absorption scales linearly with the metal volume fraction and has a roughly quadratic frequency-dependence [34-36] (see Figure 9). The large absorption coefficients in these materials are not explained by the classical electromagnetic theory of isolated metal particles, which predicts absorption coefficients on the order of $\sim 0.004\text{ cm}^{-1}$ [38]. Consideration of surface phonon contributions [39], non-local dielectric response [40], and quantum size effects [41], have also been unsuccessful in accounting for the large measured absorption of small metal particles.

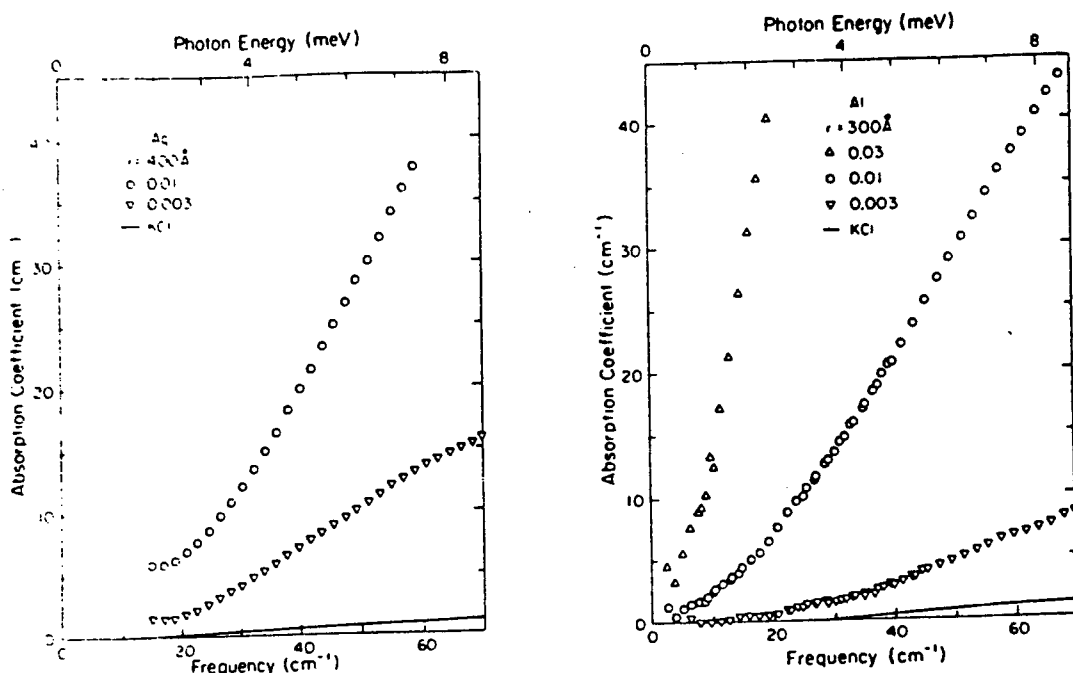


Figure 9. Far-Infrared Absorption Coefficient of (left) 400 Å Radius Ag Particles in KCl for Metal Volume Fractions of 0.003 and 0.01, and (right) 300 Å Radius Al Particles in KCl for Metal Volume Fractions of 0.003, 0.01, and 0.03 [35]

Importantly, recent experiments on Ag particles, in which the degree of clustering was carefully controlled, showed that particle clustering is crucial to the enhanced low frequency absorption in these materials [42]. In order to gain some insight into this behavior, it is useful to consider the absorption coefficient of metal particle clusters for the case of low frequencies and small volume fractions:

$$\alpha = \frac{\eta\omega^2}{c^2} \left(\frac{9c}{\omega_p^2\tau} + \frac{r^2\omega_p^2\tau}{10c} \right), \quad (15)$$

where η is the metal volume fraction, c is the speed of light, ω_p is the plasma frequency of the metal, τ is the carrier scattering time, and r is the median particle or composite size. The first term in Eq. 15 represents the contribution from electric dipole absorption, and the second term represents the contribution from magnetic dipole absorption.

Equation 15 shows that for large particle or composite sizes (typically $r > 100$ Å), electromagnetic absorption is dominated by magnetic dipole absorption, or more particularly, by the dissipation of circulating eddy currents induced in the clusters by the applied magnetic field. This regime is most probably realized in small-particle systems that have been

heat-treated so as to form fused clusters of metal particles with large median composite sizes, $r = R \sim 2000 - 3000 \text{ \AA}$ [34-36]. In this case the absorption coefficient may be written as [43]:

$$\alpha = \frac{\eta}{10c} \left(\frac{r\omega}{c} \right)^2 \epsilon_h^{1/2} \sigma_1 \quad (16)$$

where ϵ_h is the dielectric response of the insulating host, σ_1 is the real part of the optical conductivity (which is essentially constant for metals in the FIR and microwave regions), and r again represents the median cluster size, $r = R \sim 2000 - 3000 \text{ \AA}$. The resulting absorption has the experimentally observed quadratic frequency dependence and increases with cluster size. Most importantly, Eq. 16 predicts a factor of R/a ($\sim 10^1 - 10^2$) enhancement in the absorption compared to that expected in isolated particles ($r = a$), in agreement with experimental results.

Small metal particle systems that have not been appropriately heat-treated do not form "fused clusters," as described above, but are instead better described as individual particles in an insulating matrix. (The particles come into electrical contact to form chain- or cluster-like structures for sufficiently large metal volume fractions.) Absorption in this case is not dominated by magnetic dipole absorption (see Eq. 15), but rather by the electric dipole absorption term in Eq. 15,

$$\alpha = \frac{9\eta c}{\omega_p^2 \tau} \left(\frac{\omega}{c} \right)^2 \quad (17)$$

Curtin and Ashcroft [43] have performed a more detailed calculation of the dielectric response for metal clusters in this (non-heat-treated) regime, in which they model the clusters as simple cubic lattices with spacing $2a$ (where a is the individual particle size $\sim 100 \text{ \AA}$) and edge length (i.e., cluster size) L . They obtain absorption coefficients, $\alpha = 0.49 \text{ cm}^{-1}$ (at 1 meV), that are comparable to experimentally determined values and exhibit a linear frequency dependence in the far infrared and microwave regimes that increases with decreasing cluster size.

B. AMORPHOUS DIELECTRICS

The low frequency (FIR and microwave) absorption coefficients of amorphous dielectrics are also quite large, i.e., roughly an order of magnitude larger than those of comparable crystalline dielectrics (see, for example, Figure 10) [37,44,45], due most probably to the presence of low-energy vibrational modes in these systems. In this respect, microwave absorption in glassy materials shares important similarities with that proposed for nanocrystalline materials (see Section IV). This section describes mechanisms through which photons couple to low energy vibrational modes in glassy materials.

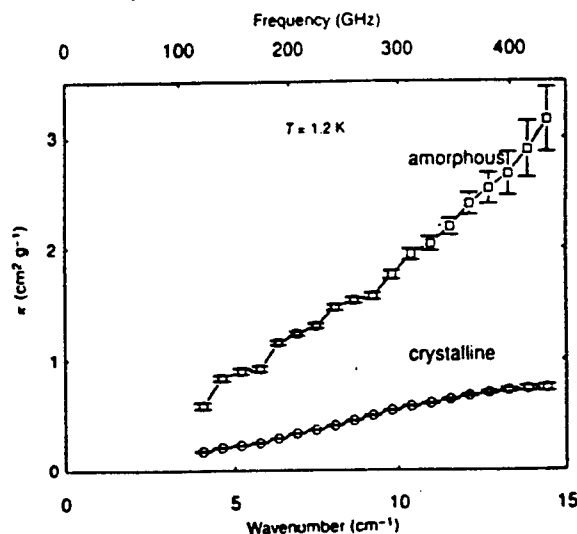


Figure 10. Absorption Coefficient per Unit Mass (κ) for Amorphous (top curve) and Crystalline (lower curve) Forsterite Powders [45]

Most chalcogenide glasses exhibit an absorption coefficient in the FIR and microwave regions that depends quadratically on frequency, $\alpha \propto \omega^2$, while some oxide glasses appear to have additional optical modes at low frequencies that obscure this characteristic quadratic frequency dependence. Absorption in these materials is rather temperature-independent in the FIR ($\omega > 5 \text{ cm}^{-1}$) but has been found to increase with temperature in the microwave region ($\omega < 5 \text{ cm}^{-1}$) [44].

The enhancement of low frequency absorption in amorphous materials has been variously attributed to disorder-induced optical coupling to Debye-like modes [46], disorder-induced charge defects or local dipole moments [47, 48], two-level tunneling states, and damped low-frequency lattice modes [44]. The effect of these mechanisms on absorption in amorphous solids can be reasonably modeled by assuming that the presence of anomalous low-frequency lattice modes induce charge fluctuations or dipole moments in the material. The absorption coefficient can in this case be written as:

$$\alpha(\omega, T) \propto \sum_q g(q) \left(\frac{1}{\omega^2 - \omega_0^2(q) + i\gamma(q)\omega} \right), \quad (18)$$

where the sum is taken over all phonon wavevectors q , $\gamma(q)$ and $\omega_0(q)$ are the damping rate and frequency of phonon mode q , respectively, and $g(q)$ is an envelope (or correlation) function representing the Fourier transform of the charge distribution associated with a

particular charge fluctuation. For example, for an exponential charge distribution, $\exp(-r/\ell)$, where ℓ is the correlation length, the corresponding correlation function is given by:

$$g(q) = 1 - \left(\frac{1}{1+q^2\ell^2} \right)^2 \quad (19)$$

For the case of Debye-like low frequency modes in amorphous solids [49], the sum over q in Eq. 18 gives:

$$\alpha(\omega) \propto K_0 \omega^2 g(q_D) \quad , \quad (20)$$

where K_0 is a constant that is independent of temperature, provided that the optical dipole moment and index of refraction are independent of temperature in the FIR and microwave frequency range (a good assumption in glasses). Equation 20 (with 21) is an extension of Eq. 9 to include the photon-phonon coupling through an electric field. Equation 20 describes an absorption coefficient that varies quadratically with frequency at high frequencies, $\alpha \propto \omega^2$, but has a quartic frequency dependence at low frequencies, $\alpha \propto \omega^4$. The additional factor of ω^2 in the latter case arises from the q^2 dependence of the photon-phonon coupling constant in the hydrodynamic regime [50].

The chief parameters determining the strength of absorption in amorphous glasses can be examined by considering the following explicit form for the constant K_0 in Eq. 20:

$$K_0 = \kappa^2 \frac{N l e^* l^2}{\rho c v^3} \quad , \quad (21)$$

where ρ is the mass density, v is the Debye sound velocity ($\sim 10^5$ cm/sec), κ is associated with local-field corrections, N is the density of charge fluctuations having magnitude $l e^* l$. The factor $N l e^* l^2$ in Eq. 21, which is related to the net effective polarization induced by the low-frequency lattice modes, is the chief parameter controlling the strength of absorption in these amorphous glass materials in this model. One expects, however, that the addition of highly ionic constituents, such as H_2O or alkali ions like Na_2O , should strongly enhance the optical absorption in amorphous materials.

Many glasses also exhibit a contribution to the optical absorption in the microwave region ($\omega < 10$ cm⁻¹) which increases substantially with increasing temperature [44], suggesting that these materials may be particularly useful as high temperature microwave absorbers. The temperature-dependent contribution to the low-frequency absorption is most likely attributable either to the presence of low-frequency, non-Debye-like modes, whose lifetimes are expected to decrease strongly with increasing temperature [51], or to relaxational processes involving the coupling of incident photons to anharmonic tunneling modes [52]. In both cases, details related to materials processing would be expected to influence the low frequency absorption characteristics, offering some promise that the

electromagnetic response of these materials can be controlled. We note finally that two proposed sources of this temperature-dependent component that can likely be ruled out are the multiphonon processes [53] and the presence of two-level systems at low frequencies, which resonantly couple to incident photons. For example, in contrast to the increase in absorption commonly observed in glassy systems with increasing temperature, it is expected that optical excitation of a distribution of two-level systems will result in a decreasing optical absorption as a function of increasing temperature.

C. APPLICATIONS AND DESIGN CONSIDERATIONS

The large variety of ways in which the absorption characteristics of composite absorbers can be modified, for example by altering the grain size, the connectivity of clusters, the dielectric medium in which the nanocrystals are suspended, or even the optically active element in the dielectric medium, suggests that these materials provide new flexibility in the engineering design of RAM and RAS. Possibilities include (1) "decoupling" the absorption characteristics of different spectral ranges; (2) improving the impedance matching to "protected" materials; and perhaps even (3) creating "responsive" structures whose dielectric response can be remotely modified and tuned.

1. Control of the Dielectric Medium

The proposed nanocrystalline absorber materials are essentially two-component materials, comprising metal nanophase particles suspended in a dielectric medium. This feature allows flexibility in controlling, to some extent independently, the transmission/absorption characteristics in widely separated spectral regions. For example, one could use a dielectric "matrix," which optimizes transmission in the visible region and yet maintain significant microwave/far-infrared absorption by the dilute suspension of metallic nanocrystals. Such a design would be of use for microwave/infrared protection of optical components, airplane canopy coatings, etc. By contrast, the current use of thin layers of single-component gold or other metal films in protecting optical parts suffers from the direct competition between optimization of visible transmission (decreased film thickness) and infrared absorption (increased film thickness).

2. Graded Dielectric Designs

Similar to conventional "Salisbury screen" and "Dallenbach layer" absorber designs, the magnetic absorption associated with a single layer of comparably sized magnetic nanocrystalline clusters has a narrow absorption bandwidth that is given roughly by $\mu_{\max} \Delta\omega$

$\approx 4\pi\gamma M_s$, where $\mu''_{\max}$ is the maximum value of the imaginary part of the permeability (representing the peak in the dissipation), $\Delta\omega$ is the absorption bandwidth, and M_s is the magnetization. However, by creating a continuous distribution of nanocrystal cluster sizes as a function of layer thickness (for example, by physically layering crystallites of different sizes in a manner similar to layering using standard MBE techniques, or by introducing thermal gradients during processing to control the degree of "fusing" of particles at different depths), one may be able to obtain a "graded absorption" design which can (1) greatly expand the absorption bandwidth (similar to that provided by the conventional Jaumann absorbers), and (2) improve the impedance match between the "coating" and the protected material.

3. Dynamically Adjustable/Field-Responsive Designs

An ideal absorber material would incorporate a "responsive" element whose absorption characteristics could be either tuned dynamically in real time or could self-adjust in response to changes in the incident field from a threat. For example, one would like to be able to modulate the optical response of the material at the incident frequency of the "threat" field in order to achieve signal suppression via optical interference. While extensive tests are needed to determine the extent to which one can achieve responsivity and dynamical control with the composite absorber materials proposed here, the multi-component structure of these materials clearly provides a means by which remote and dynamic control of optical materials can be obtained. For example, viewed as an array of dipoles embedded in some medium, the absorption coefficients of metal-cluster absorbers can in principle be tuned dynamically (1) by changing the dielectric response of the medium (electro- or acousto-optically, for example) or (2) by modulating the orientation of the dipoles through the application of an electric or magnetic field.

One way to allow small reorientations of magnetic particles is to embed them in a viscous medium. Under an applied field, the nanocrystals would migrate slightly, and the energy of their magnetic alignment would develop greater anisotropy. Altering the clustering of the particles would allow for control over the eddy current losses and control over the vibrational modes of the particles. It might be possible to change the RAM characteristics at a frequency comparable to the modulation of the radar signal, although it is unlikely that these mechanical motions could occur with the nanosecond time constants required for interaction with the radar carrier frequency.

The vibrational properties of nanophase particles are likely to be highly nonlinear. It is possible that in higher amplitude electromagnetic fields the vibrational spectrum could change considerably. It is also possible for the local arrangements of nanoparticles to change

when driven by high amplitude electromagnetic fields. The time constants for these changes could be quite short and may allow for interaction with the radar carrier frequency.

Finally, the introduction of certain "phase change" materials as the optically active element may be of use as remotely tunable absorber materials. In particular, a number of materials undergo metal-insulator transitions as a function of free-carrier density, temperature, or magnetic field. These materials may be useful as optical switches operated by different external control parameters (see Figure 11). For example, materials such as Mott insulators (e.g., VO_x and V_2O_3) and doped semiconductors, which exhibit metal-insulator transitions as a function of free-carrier density, can, in principle, be the basis of optical switches whose control parameter is the incident light power (which controls the density of excited electron-hole pairs). An obvious application for such an optical switch would be laser protection, in which low-power ambient light is transmitted below the optical gap in the insulating phase, but high laser intensities drive a insulator-metal transition causing strong reflection of the beam. It is likely that the responsivity of this transition is too slow for critical applications in which the detector can be damaged by high light power, but such a "switch" may be useful for non-critical applications such as protecting CCD detectors. A class of materials in which either temperature or magnetic field can be used as the control parameter for inducing a metal-insulator transition are Giant Magnetoresistance (GMR) materials such as $\text{La}_{1-x}\text{Sr}_x\text{MnO}_3$ [54]. An example of a metal-insulator transition in $\text{La}_{1-x}\text{Sr}_x\text{MnO}_3$ is illustrated on the right side of Figure 11, showing the dramatic increase in the resistivity of the material above a critical temperature, $T_c \sim 300$ K. These materials are potentially useful as optical switches in which temperature or magnetic field are convenient (remote) control parameters. Although the metal-insulator transitions in these materials involve large changes in optical density, as required of a good optical switch, the usefulness of these materials is limited by the slow switching speeds and the restricted spectral regions for high optical transmission. The metal-insulator transition could, however, provide control over the eddy current losses in RAM.

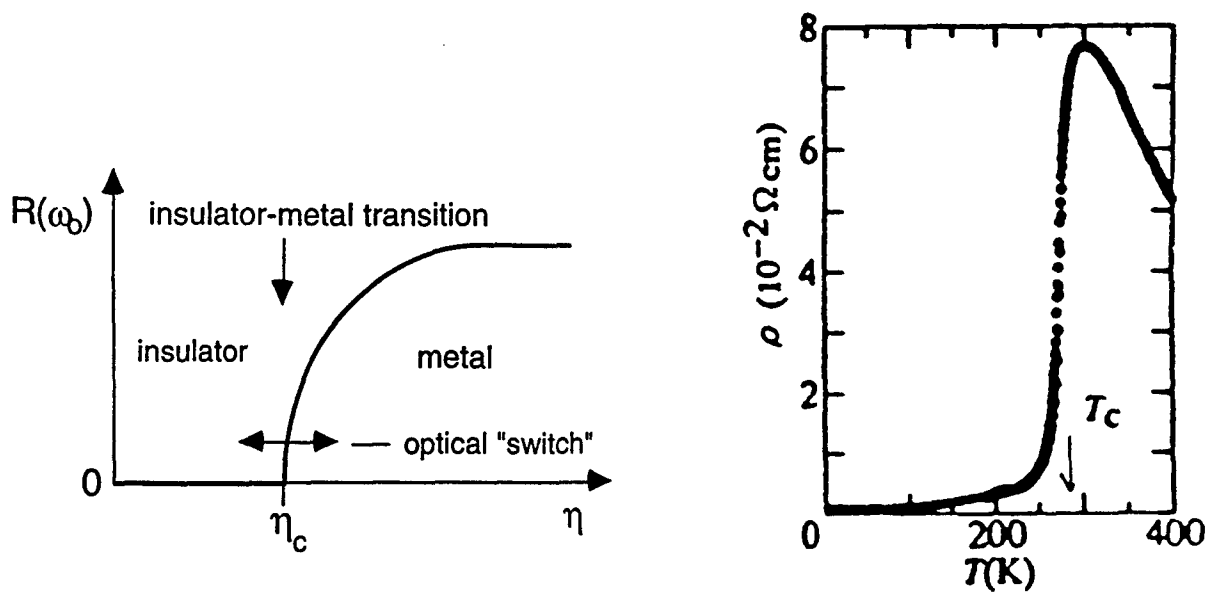


Figure 11. (left) Schematic Illustration of the Use of a Metal-Insulator Transition as an “Optical Switch,” Showing the Optical Reflectance, R , as a Function of a Control Parameter, η ; (right) An example of a Temperature-Induced Metal-Insulator Transition in $\text{La}_{1-x}\text{Sr}_x\text{MnO}_3$ [54]

VII. PRELIMINARY RESULTS

We have made some preliminary measurements of microwave absorption and phase shifts over the frequency range of 1- 20 GHz. The measurements were performed on nanophase Fe powder prepared by high energy ball milling. The powdered metal had a characteristic crystallite size of 10 nm, as determined by the broadening of x-ray diffraction lines [55-58]. Annealing the material at 500° C for 2 hours caused significant growth of the crystallites to a characteristic size of greater than 30 nm. Small amounts of powder were mounted in a coaxial waveguide located between the transmitter and receiver of a Hewlett Packard 8510C microwave network analyzer. The microwave power density in the sample region, shown in Figure 12, was about 10 W/cm².

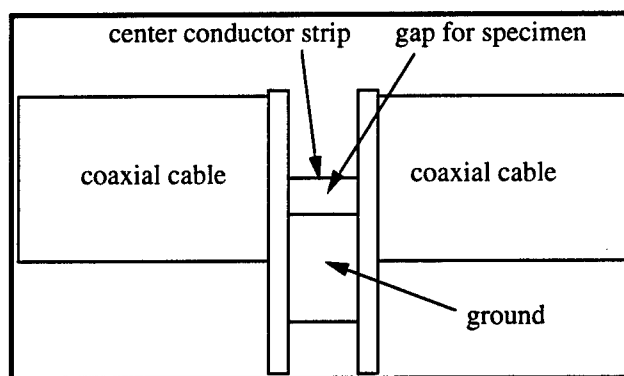


Figure 12. Specimen Region for Microwave Absorption Measurements

The network analyzer was calibrated with an empty specimen region; we term this absorption spectrum the instrument response. Powders of the two samples (nanophase and large-grained) were placed in the specimen cavity between calibration runs. The spectrum of the voltage detected at the receiver versus frequency was normalized by the instrument response. The absorbed power, obtained as the square of the reduction in normalized voltage, is presented in Figure 13.

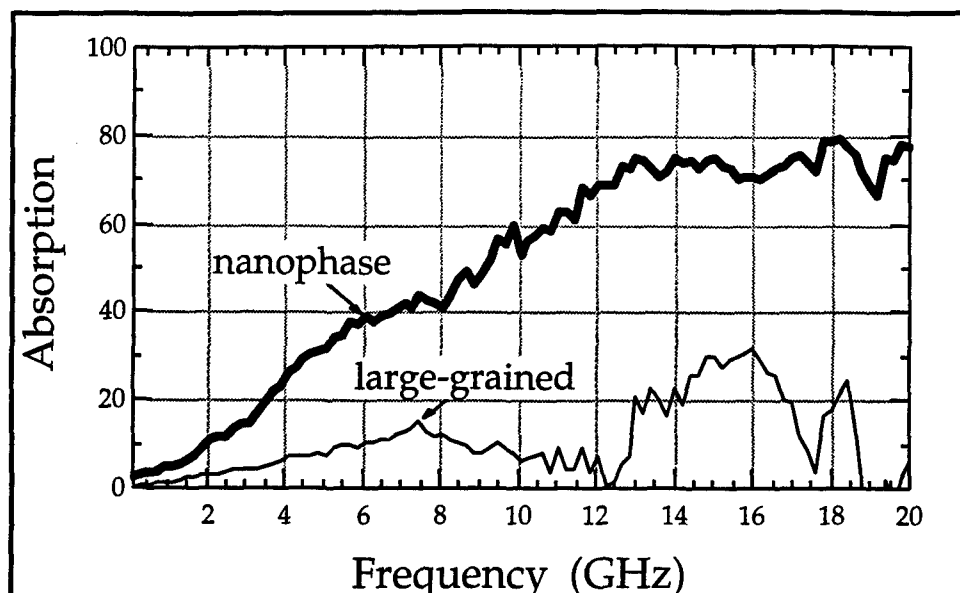


Figure 13. Microwave Power Absorption by Nanophase Fe with 11 nm Crystallites (dark curve), and Annealed Fe with >30 nm Crystallites (light curve)

The microwave absorption spectra of both the nanophase and large-grained Fe powder are broadband. No strong resonances are seen in the nanophase material, although there may be some resonances in the large-grained material at the higher frequencies. More interestingly, the power absorbed by the nanophase material is significantly greater than the power absorbed by the large grained material. The data indicate that an equivalent mass of nanophase Fe may have an absorption strength three or four times larger than for large-grained Fe.

We caution the reader of some uncertainties in our measurements. The sample region was quite small, and we had difficulty in ensuring that equal amounts of powder were used. Nevertheless, the measurements were repeated with four sets of samples, and the nanophase material showed a consistently stronger absorption. (The data of Figure 13 are an average of these different data sets.)

We are not sure if the reasons for the microwave activity of the nanophase powder are related to the vibrational mechanism proposed in Section IV. The electrical conductivity of the Fe powders is probably changed upon annealing, for example. A decrease in conductivity after annealing could be responsible for reduced eddy current losses in the large-grained material. We expect, however, that the electrical conductivity should decrease after annealing, in disagreement with the data of Figure 13.

Further studies on the microwave absorption characteristics of nanophase magnetic materials seem appropriate. A redesign of the specimen cavity could be a good next step. Further work could also involve applying a large static magnetic field to the sample to suppress its high frequency magnetic response. Work with powders that have soft magnetic properties superior to iron also seems appropriate. One of the authors is interested in these further investigations.

REFERENCES

1. *Jane's Radar and Electronic Warfare Systems*, ed. Bernard Blake (Surrey, UK: Jane's Information Group), p. 591.
2. O. Sandus, confidential study at the Univ. of Michigan (1961).
3. J. Jones, *Stealth Technology: The Art of Black Magic*, ed. M. Thurber (Blue Ridge Summit, PA: Aero Division of TAB Books, 1989).
4. B. Sweetman, *Stealth Bomber – Invisible Warplane, Black Budget* (Osceola, WI: Motorbooks International, 1989).
5. B.D. Cullity, *Introduction to Magnetic Materials*, (Reading, MA: Addison-Wesley, 1972).
6. Y. Ng et al., *Chem. Mater.*, 4 (1992), p. 885.
7. S.I. Stupp et al., *Science*, 59 (1993), p. 259.
8. Y. Yoshizawa et al., *J. Appl. Phys.*, 64 (1988), p. 6044 .
9. G. Herzer, *IEEE Trans. Magnetics*, 26 (1990), p. 1397.
10. S. Farad et al., *Phys. Rev. B*, 52 (1995), p. 5752.
11. J. Eastman et al., *Research and Development*, 57 (January 1989).
12. C.G. Granqvist and R. A. Burhman, *J. Appl. Phys.*, 47 (1976), p. 2200.
13. H. Gleiter, *Prog. Mater. Sci.*, 33 (1989), p. 223.
14. R.W. Siegel, *Mater. Res. Soc. Bull.*, 15 (1990), p. 60.
15. E. Hellstern et al., *J. Mater. Res.*, 4 (1989), p. 1292.
16. H.J. Fecht et al., *Adv. Powder Metall.*, 1-2 (1989), p. 111.
17. B. Fultz et al., *J. Mater. Res.*, 4 (1989), p. 1450.
18. R. Alben et al., *J. Appl. Phys.*, 49 (1978), p. 1653.
19. H.H. Hamdeh et al., "Structure and Magnetic Properties of Sputtered Thin Films of $\text{Fe}_{0.79}\text{Ge}_{0.21}$," *J. Appl. Phys.*, 74 (1993), p. 5117-23.
20. Z. Gao and B. Fultz, *Hyperfine Interactions*, 94 (1994), pp. 2213-18.
21. Z. Gao and B. Fultz, *Nanostructured Materials*, 2 (1993), pp. 231-40.
22. Z. Q. Gao and B. Fultz, *NanoStructured Materials*, 4, (1994), pp. 939-47.
23. F. Sato et al., *Phys. Stat. Solidi*, 107 (1988), p. 355.
24. K.-H. Hellwege, ed. *Landolt-Börnstein: Numerical Data and Functional Relationships in Science and Technology*, (Springer-Verlag, Berlin, 1981), Volume III/13a.

25. B. Fultz et al., *J. Appl. Phys.*, 76 (1994), p. 5961.
26. B. Fultz et al., *Phys. Rev. B*, 52 (1995), p. 3315.
27. J. S. C. Jang and C. C. Koch, *J. Mater. Res.*, 5 (1990), p. 498.
28. T. Nasu et al., *J. Non-Crystalline Solids*, 150 (1992), p. 491.
29. S. R. Harris et al., *J. Mater. Res.*, 6 (1991), p. 2019.
30. A. R. Yavari, *Acta Metall. Mater.*, 41 (1993), p. 1391.
31. M. D. Baró et al., *Acta Metall. Mater.*, 41 (1993), p. 1065.
32. C. Stassis et al., *Phys. Rev. B*, 24 (1981), p. 3048.
33. C.N.J. Wagner, TMS Annual Meeting, Las Vegas, NV, February 1995.
34. D.B. Tanner et al., *Phys. Rev. B*, 11 (1975), p. 1330.
35. G.L. Carr, et al., *Phys. Rev. B*, 24 (1981), p. 777.
36. N. E. Russell et al., *Phys. Rev. B*, 23 (1981), p. 632.
37. K.K. Mon et al., *Phys. Rev. Lett.*, 35 (1975), p. 1352.
38. J.C. Maxwell-Garnett, *Philos. Trans. Royal Soc. London*, 203 (1904), p. 385.
39. A.J. Glick and E.D. Yorke, *Phys. Rev. B*, 18 (1978), p. 2490.
40. P. Apell, *Phys. Scr.*, 29 (1984), p. 146.
41. P.N. Sen and D.B. Tanner, *Phys. Rev. B*, 26 (1982), p. 3582.
42. R.P. Devaty and A.J. Sievers, *Phys. Rev. Lett.*, 52 (1984), p. 1344.
43. W.A. Curtin and N.W. Ashcroft, *Phys. Rev. B*, 31 (1985), p. 3287.
44. U. Strom and P. C. Taylor, *Phys. Rev. B*, 16 (1977), p. 5512.
45. N.I. Agladze et al., *Nature*, 372 (1994), p. 243.
46. A. Hadni et al., *C. R. Acad. Sci.*, 260 (1965), p. 4973.
47. W. Bagdade and R. H. Stolen, *J. Phys. Chem. Solids*, 29 (1968), p. 2001.
48. E. Whalley, *Trans. Faraday Soc.*, 68 (1972), p. 662.
49. A Debye-like density of states, and the elastic continuum model of solids, is appropriate for a glassy system if the phonon wavelength is much larger than the correlation length, i.e., $\lambda \gg l$; see, for example, Ref. 11.
50. A. J. Martin and W. Brenig, *Phys. Status Solidi B*, 64 (1974), p. 163.
51. G. Winterling, *Phys. Rev. B*, 12 (1975), p. 2432.
52. N. Theodorakopoulos and J. Jäckle, *Phys. Rev. B*, 14 (1976), p. 2637.
53. E.M. Amrhein and F.H. Müller, *J. Am. Chem. Soc.*, 90 (1968), p. 3146.
54. Y. Okimoto et al., *Phys. Rev. Lett.*, 75 (1995), p. 109.
55. G.K. Williamson and W.H. Hall, *Acta Metall.*, 1 (1953), p. 22.

56. H.P. Klug and L.E. Alexander, *X-Ray Diffraction Procedures* (Wiley-Interscience, New York, 1974) p. 664.
57. P. Scherrer, *Gött. Nachr.*, 2 (1918), p. 98.
58. H.P. Klug and L.E. Alexander, *X-Ray Diffraction Procedures* (Wiley-Interscience, New York, 1974) p. 656.

**H. FROM CHIPS TO SHIPS: APPLYING VLSI CAD
TECHNIQUES TO NAVAL VESSEL DESIGN**

**Gabriel Robins
Department of Computer Science
University of Virginia
Charlottesville, VA**

ABSTRACT

We address optimization issues in ship and submarine construction, using techniques from integrated circuit design. We begin by observing that much like chip design, naval vessel design entails *placing* numerous components (engines, controls, sensors, actuators, weapon systems, etc.) inside a hull, and then *interconnecting* them (with cables, pipes, ducts, etc.) into a single integrated system. This process mirrors the methodology of circuit design, where we place electrical components (gates, registers, adders, memories, etc.) on a chip and interconnect them with wires. Fundamental objectives in both domains include minimizing the total interconnect used (i.e., wirelength in circuit design, vs. cablelength, ductlength, and pipelength in ship/sub design), as well as optimizing performance, reliability, and fault-tolerance. This basic analogy enables methods that address both circuit design and naval vessel construction. With the growing size and complexity of modern naval vessels, such methods could improve procurement cost, overall weight, reliability, survivability, fuel consumption, and ease-of-maintenance.

From Chips to Ships: Applying VLSI CAD Techniques to Naval Vessel Design

Abstract

We address optimization issues in ship and submarine construction, using techniques from integrated circuit design. We begin by observing that much like chip design, naval vessel design entails *placing* numerous components (engines, controls, sensors, actuators, weapon systems, etc.) inside a hull, and then *interconnecting* them (with cables, pipes, ducts, etc.) into a single integrated system. This process mirrors the methodology of circuit design, where we place electrical components (gates, registers, adders, memories, etc.) on a chip and interconnect them with wires. Fundamental objectives in both domains include minimizing the total interconnect used (i.e., wirelength in circuit design, vs. cablelength, ductlength, and pipelength in ship/sub design), as well as optimizing performance, reliability, and fault-tolerance. This basic analogy enables methods that address both circuit design and naval vessel construction. With the growing size and complexity of modern naval vessels, such methods could improve procurement cost, overall weight, reliability, survivability, fuel consumption, and ease-of-maintenance.

1 Motivation and Scope

The members of the Defense Science Study Group (DSSG) have recently visited a number of U.S. Navy ships and submarines.¹ During our visits we have observed that one striking and ubiquitous feature common to all these naval vessels is the sheer number, extent, and complexity of the pipes/cables/wires/ducts that permeate the length of the ship/sub's interior. These pipes, cables, wires, and ducts, collectively referred to as *interconnect* in this paper, while crucial to the successful operation of the ship's subsystems, nevertheless seem to exact an enormous toll in terms of the ship's overall weight, performance, complexity, reliability, maintainability, and procurement cost.

The main goal of this study is to develop combinatorial algorithms and techniques in order to optimize the overall interconnect structures used in ship construction. This in turn would help reduce the ship's complexity and cost, while increasing performance, reliability, survivability, and maintainability. Since all the methods proposed here apply equally well to surface ships as well as to submarines, we shall use the term "ship" in this paper to include both surface vessels and submarines. Indeed, our combinatorial optimization techniques are also applicable to other types of structures/domains altogether, including buildings and even entire installations/bases.

¹Including the aircraft carrier USS Eisenhower, the Ohio-class Trident missile submarine USS Nebraska, the Los Angeles-class attack submarine USS Minneapolis Saint Paul, the supply ship USS Supply, and an Aegis missile cruiser.

The basic premise and motivation of this work is the observation that there exists a very natural and compelling set of analogies between routing electrical components on a VLSI (Very Large Scale Integrated) chip, and interconnecting subsystems aboard a naval vessel:

- **Placement and routing:** both chip design and ship design entail the *placement* of various components (gates, adders, memories, etc. in the case of chips, vs. engines, controls, sensors, actuators, weapons systems, etc. in the case of ships) about a given area (i.e., a VLSI chip, vs. a ship's hull), and then *interconnecting* them into a single integrated system (using wires and buses in chips, vs. pipes, cables, and ducts in ships). Moreover, in both domains there is substantial interplay/synergy between the placement and routing phases [35];
- **Interconnect minimization:** in both chip design and in ship design, we typically seek to minimize the overall interconnect length: on a chip, wirelength savings translates into reduced area, higher operating speed, and increased manufacturability; similarly, onboard a ship, a reduction in overall interconnect length yields weight and space savings, which in turn has good implications with respect to increased cargo/payload capability, improved operating range, higher speed and maneuverability, etc.;
- **Numerous subsystems/nets:** both problems assume that multiple interconnections must be routed without interfering with each other, and without producing overly-congested "hotspots". On a chip for example, multiple nets such as power, ground, clock, buses, and various other signals must remain electrically disconnected from each other, while a ship similarly contains multiple separate interconnect subsystems, such as power, water, drainage, air-conditioning, communications, fuel, hydraulics, etc.;
- **Problem size:** the problem size in both domains tends to be very large and steadily increasing in size: many tens of thousands of modules/devices must be interconnected. Consequently, designers are unable to *manually* produce effective solutions in either domain;
- **Obstacle avoidance:** both problem domains require an obstacle-avoidance capability during the design process (e.g., wires on a chip must pass around active electronic components, while pipes on a ship must avoid equipment, bulkheads, certain restricted spaces, etc.);
- **Rectilinear metric:** in both domains the underlying routing space is implicitly rectilinear; that is, all wires on a chip may extend only in the horizontal and vertical directions², while the standard principles of ship design dictate an analogous constraint [15] [46];
- **Multiple layers:** both problems are essentially three dimensional (i.e., multiple metal layers of a chip correspond to decks on a ship, etc.);
- **Reliability and fault-tolerance:** chips and ships are both prone to component failure, thus fault-tolerance is an important issue (e.g., heat stress or electro-migration can disconnect VLSI wires, while accidents, corrosion, or battle damage can sever shipboard interconnect).

²This restriction induces the *Manhattan*, or *rectilinear* metric. Using only rectilinear wiring in VLSI circuits facilitates routability and manufacturability. In practice, each wiring layer has a preferred wiring direction (e.g., horizontal layers and vertical layers might alternate), with structures called *vias* used to connect orthogonal wires on adjacent layers [35].

These basic analogies enable the application of VLSI Computer-Aided Design (CAD) formulations and techniques in order to afford naval vessel designers the following potential advantages:

- **Procurement cost:** our optimizations will enable the reduction of total interconnect length (i.e., cables, ducts, and pipes), which translates into savings in building materials and construction time, which in turn reduces the overall manufacturing cost;
- **Overall weight:** less materials used in the construction of a ship imply an overall reduction in the ship's gross weight, yielding a corresponding improvement in payload capacity and/or performance (i.e., speed, range, maneuverability, and fuel economy);
- **Reliability:** reduced interconnect length implies a smaller number of potential points-of-failure (due to corrosion, metal fatigue, defective welds, etc.) and therefore would contribute to a higher overall fault-tolerance;
- **Maintenance:** less overall interconnect would require correspondingly less maintenance; moreover, in any given shipboard workspace/quarters, a lower density of cables, pipes, and ducts implies improved accessibility during routine maintenance (or during emergency / battle operations); and
- **Survivability:** some of our formulations enable resilient interconnect topologies that utilize redundancy to minimize the damage sustained by, e.g., a missile impact (Admiral Ike Kidd pointed out that "graceful degradation" is a crucial and much-needed capability for the Navy).

Our approach to shipboard interconnect optimization will leverage off our research in algorithms for VLSI routing, and many of our proposed techniques also apply to VLSI CAD as well as to ship design. While offering some preliminary solutions, this paper does not claim to be comprehensive. Rather, we focus on formalizing the various issues and suggest possible methods for attack. We plan to continue to meet with ship designers and key Navy personnel for feedback and "reality-checks", in order to ensure that our approaches and solutions are viable/practical. Also, we would like to continue to implement the algorithms discussed here into a unified workbench for shipboard interconnect design, and coordinate with other research groups (and defense contractors) in order to embed our solutions/code into larger, existing systems for ship/submarine design.

2 Interconnect Optimization: The Steiner Problem

A major phase of naval vessel design entails the *placement* of various components (e.g., engines, controls, sensors, actuators, weapon systems, etc.) about the hull, and then *interconnecting* them into a single integrated system. We noted above that this process resembles the basic methodology of circuit design, where we place electrical components (e.g., gates, adders, memories, etc.) on a chip and then interconnect them with wires. In both of these domains, a fundamental objective is the minimization of the total interconnect used (i.e., wirelength in circuit design, vs. cablelength,

ductlength, and pipelength in ship design). We start our technical discussion with a formal statement of the interconnect minimization problem (i.e., the Steiner problem) as it applies to naval vessel design, and we then describe some of our proposed algorithms and solutions.³

Given a set of shipboard "components" to interconnect (e.g., electrical outlets, fire sprinklers, water faucets, sinks/drains, computers, air-conditioning outlets, intercoms/speakers, controls, actuators, etc.) we model the objects to be interconnected as a set of points in three-dimensional space, with the coordinates of the points corresponding to the components' positions on the ship (in order to streamline the exposition, we will for now ignore such issues as obstacles and congestion - these considerations will be addressed later in Section 5).

If the interconnection topology is only allowed to make *direct* connections between pairs of the original/input points, then minimizing the overall length of the topology yields the well-known *minimum spanning tree* (MST) problem [37]. However, if intermediate junctions, or *Steiner points*, are allowed, this additional flexibility enables further savings in total interconnect cost by effective overlapping, as shown in Figure 1. The cost of an edge is defined as the (rectilinear) distance between its endpoints, and the cost of a whole tree is the sum of its individual edge costs (the term "tree" itself is borrowed from graph theory and is used to denote an acyclic interconnect topology; in a Section 6.1 below we will address non-tree topologies). The term *points* in this discussion denotes the set of components that are to be interconnected. The Steiner minimal tree problem, which seeks to minimize the overall interconnect length by the judicious addition of Steiner points, may now be formally defined as follows:

The Steiner Minimal Tree (SMT) problem: Given a set P of n points, find a set S of *Steiner points* such that the minimum spanning tree (MST) over $P \cup S$ has minimum cost.

Due to constraints imposed by both VLSI CAD and naval vessel design, the underlying space is assumed to have a so-called Manhattan geometry, where the distance between two points (x_1, y_1) and (x_2, y_2) is defined as the $\Delta x + \Delta y = |x_1 - x_2| + |y_1 - y_2|$. Arguably, the Manhattan metric is valid not only for naval vessel design [15] [46], but also in many other domains, such as designing buildings (due to the orthogonality of walls / floors / corridors), where the interconnect (i.e., pipes, cables, ethernet, ducts) are all aligned parallel to one of the main coordinate axis.⁴ Therefore, our discussion will focus on the Minimum Rectilinear Steiner Tree (MRST) version of the SMT problem. Figure 1 gives an example of an MST and an MRST for the same set of four points (for clarity, we shall often use two-dimensional illustrations, but all of our methods apply to three dimensions as well).

Research on the MRST problem has been historically guided by several fundamental results. First, Hanan [19] has shown that there always exists an optimal MRST with Steiner points chosen from the

³We focus primarily on the *routing* aspects of the ship's design, rather than the *placement* considerations, since the latter is much more constrained/predetermined by external factors. For example, some of the ship's main systems/components (e.g., engines, reactors, radars, weapons, and bridge) must be placed in specific locations on the ship, as dictated by physical/mechanical constraints, or by military doctrine, etc.

⁴All of our techniques are generalize to "restricted-angle" geometries, where only a small set of allowable routing directions can be used - e.g., 45-90 or 0-30-60 -degree geometries.

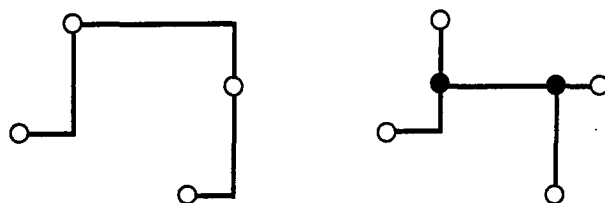


Figure 1: A minimum spanning tree (left) and MRST (right) for a fixed pointset in the Manhattan plane; hollow dots represent the original pointset P , while solid dots represent the set S of added Steiner points. Note how the judicious introduction of Steiner points reduces the total amount of interconnect required.

intersections of all the horizontal and vertical lines passing through the input points, and this result was generalized by Snyder to all higher dimensions [45]. However, a second major result by Garey and Johnson [12] establishes that despite this restriction on the solution space, the MRST problem remains NP-complete⁵; this has given rise to numerous heuristics, as surveyed in [27].

In solving intractable problems, we often seek provably-good heuristics having bounded worst-case error from optimal. Thus, a third important result is the discovery by Hwang [24] that the rectilinear MST is a fairly good approximation to the MRST, with a worst-case performance ratio⁶ of $\frac{\text{cost}(MST)}{\text{cost}(MRST)} \leq \frac{3}{2}$. One of the implications of this is that any MST-based strategy that operates by improving upon an initial MST topology will also enjoy a performance ratio of at most $\frac{3}{2}$. Moreover, such algorithms tend to be computationally efficient, since the MST for a planar pointset is easy to compute [17] [36]. These considerations have prompted a large number of Steiner tree heuristics that resemble classic MST construction methods [3] [4] [11] [21] [22] [26] [32] [38] [41] [43].

3 Inefficacy of Previous Methods

Minimum spanning tree (MST) -based Steiner tree heuristics produce solutions with average cost 7% to 9% smaller than MST cost [27]. Unfortunately, all MST-based MRST constructions have a worst-case performance ratio as poor as $\frac{3}{2}$ times optimal, as shown below [29]. We restrict our initial discussion to two-dimensional Manhattan space, although later we will generalize our analysis to three (and higher) dimensions.

Consider the following two prototypical heuristic approaches, called *MST-Overlap* and *Kruskal-Steiner*, for the rectilinear SMT problem.

⁵ *NP-complete* is a technical term which means that it is not known whether the given problem can be solved in time polynomial in the input size. Moreover, if the problem at hand can be solved efficiently, many other problems with previously unknown time complexities can be solved efficiently as well. Thus, NP-completeness captures our intuitive notion of computational intractability [13].

⁶ A *performance ratio* is an upper bound, given as the worst-case ratio of how much costlier the heuristic solution can be with respect to the optimal solution. Naturally, the lower the performance ratio, the better the heuristic, with an *exact* algorithm having a “performance ratio” of precisely 1.

- **MST-Overlap:** this approach starts with a rectilinear MST and obtains a Steiner tree by overlapping edge embeddings. In other words, a monotone (staircase) embedding is selected for each MST edge, and then all superposed segments are merged since they represent redundant wiring.
- **Kruskal-Steiner:** this strategy emulates the standard MST constructions of Kruskal [31] and Prim [37]. It begins with a spanning forest of n isolated components (the points of P) and repeatedly connects the closest pair of components in the spanning forest until only one component (i.e., the final Steiner tree) remains. Typically, the actual geometric embeddings of edges within their bounding boxes are left unresolved for as long as possible during the construction, which allows the greatest possible freedom in minimizing the edge lengths.

We now show that MST-Overlap and Kruskal-Steiner belong to a general class of greedy Steiner tree heuristics, all of which perform no better than the simple MST (i.e., $\frac{3}{2}$ times the optimal). Recall that a Steiner tree may be viewed as a minimum spanning tree over $P \cup S$, where P is the input pointset and S is the added set of Steiner points. We are interested in Steiner tree constructions which induce new edges, and possibly new Steiner points, using the following types of *connections* within an existing spanning forest over $P \cup S$: (i) *point-point* connections between two points of P ; (ii) *point-edge* connections between a point of P and an edge, which may induce up to one new Steiner point in S ; and (iii) *edge-edge* connections between two edges, which may induce up to two new Steiner points in S . To reflect the fact that the embedding of a given edge is indeterminate, we allow any edge between two points of $P \cup S$ to be arbitrarily *re-embedded* by the Steiner tree construction. Figure 2 defines this class C of Steiner tree heuristics.

Class C of greedy Steiner Heuristic:
Input: n isolated components (points of P)
Output: Rectilinear Steiner tree over P
While there is more than one connected component Do
Select a connection type $\tau \in \{ \text{point-point, point-edge, edge-edge} \}$
Connect the <i>closest</i> pair of components greedily with respect to τ
Optionally at any time, Re-embed any edge within its bounding box
Optionally at any time, Remove redundant (overlapped) edge segments
Output the single remaining component

Figure 2: The class C of greedy Steiner tree heuristics.

Theorem 3.1 *Every heuristic $H \in C$ has performance ratio arbitrarily close to $\frac{3}{2}$.*

Proof: The MST of the pointset depicted in Figure 3(b) is unique since all interpoint distances < 3 are unique. Thus, all connections in the MST are horizontal point-point connections except for exactly two connections, one from the top row to the middle row and one from the middle row to the

bottom row. The greedy routing of every edge but these two is unique since all edges except these two have degenerate bounding boxes. No improvement is possible by edge re-embedding within these degenerate bounding boxes. Therefore, no heuristic in C can do better than the result in Figure 3(d). The optional re-embedding within the two non-degenerate bounding boxes is negligible as n grows large, hence the performance ratio is arbitrarily close to $\frac{3}{2}$. $\square$

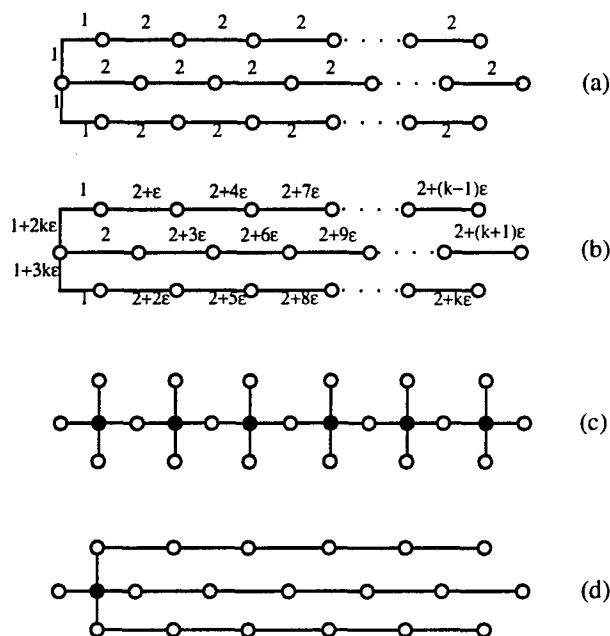


Figure 3: An MST for which $\frac{\text{cost}(\text{MST-Overlap})}{\text{cost}(\text{SMT})}$ is arbitrarily close to $\frac{3}{2}$. For n points, any Steiner tree derivable from the MSTs of (a) or (b) will have cost $2(n-2)$, while the SMT (c) has cost $\frac{4}{3}(n-1)$, yielding a performance ratio arbitrarily close to $\frac{3}{2}$ for large enough n . In (d), we show the best possible rectilinear Steiner tree that can be produced by any MST-Overlap or Kruskal-Steiner heuristic.

There are many heuristics in the literature with previously unknown performance ratio, which by Theorem 3.1 have performance ratio arbitrarily close to $\frac{3}{2}$. Greedy Kruskal-like constructions include the methods of [3] [4] [11] [25] [33] [38] [41]. Algorithms which start with an initial MST and then overlap edges within their bounding boxes, such as those of [21] and [22], also belong to C : an MST can be constructed using only point-point connections, and the optional re-embedding is then used to induce edge overlaps. Exponential-time methods can also belong to the class C , notably the suboptimal branch-and-bound method of [47]. Theorem 3.1 implies that all of these methods have the same worst-case error bound as the simple MST.

The counterexample of Figure 3 also establishes lower bounds arbitrarily close to $\frac{3}{2}$ for the performance ratios of several heuristics not in C , such as the three-point connection methods of [26] [32], and the Delaunay triangulation-based method of [43]. This is easy to verify using the pointset in Figure 3(b): as with the heuristics in C , these latter methods are severely constrained by the nature

of the unique minimum spanning tree. The works of [8] and [42] give a less general construction with the goal of showing a $\frac{3}{2}$ performance ratio for MST-like heuristics.

The rectilinear SMT problem remains well-defined when the points of P are located in d -dimensional Manhattan space with $d > 2$. Many heuristics, including those in the class C defined above, readily extend to higher dimensions. However, the construction of Figure 3 also extends to d dimensions, where it again provides a lower bound for the performance ratio of such heuristics. In d dimensions, the Figure 3 construction generalizes to $n = k(2d - 1) + 1$ points, for any given k . As Figure 4 illustrates for $d = 3$, the cost of the optimal Steiner tree is at most $\frac{2d(n-1)}{2d-1}$; the cost of the (unique) MST is $2(n - 1)$; and the cost of the best Steiner tree obtainable from the MST by edge-overlapping is $2(n - d)$. Thus, in d dimensions the performance ratio of a heuristic in C will be arbitrarily close to $\frac{2d-1}{d}$, which converges to 2 for large d (our analysis improves on the bound given in [9]).

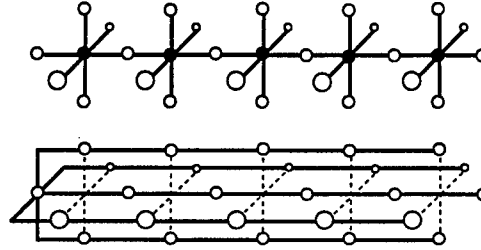


Figure 4: For $d = 3$, the SMT (top) has cost $\frac{6}{5}(n - 1)$, while any Steiner tree derivable from the MST by re-embedding edges (bottom) has cost $2(n - 3)$, yielding performance ratio arbitrarily close to $\frac{5}{3}$ as n grows large.

4 An Effective Steiner Heuristic

The negative results discussed above motivate research into alternate schemes for MRST approximation. Thus, we now describe an effective heuristic for the Steiner problem, which we call the Iterated 1-Steiner method [28]; this heuristic, appropriately extended and generalized as described below, can be directly applied to the shipboard interconnect optimization problem.

4.1 The Iterated 1-Steiner Method

For two pointsets P and S we define the MST savings of S with respect to P as $\Delta MST(P, S) = \text{cost}(MST(P)) - \text{cost}(MST(P \cup S))$. We use $H(P)$ to denote the set of Hanan Steiner point candidates (i.e., the intersections of all horizontal and vertical lines passing through points of P). For a pointset P , a 1-Steiner point $x \in H(P)$ maximizes $\Delta MST(P, \{x\}) > 0$. The Iterated 1-Steiner (IIS) algorithm repeatedly finds 1-Steiner points and includes them into S . The cost of the MST over $P \cup S$ will decrease with each added point, and the construction terminates when there is no x with $\Delta MST(P \cup S, \{x\}) > 0$. Figure 5 illustrates the execution of IIS, and Figure 6 describes the algorithm formally.

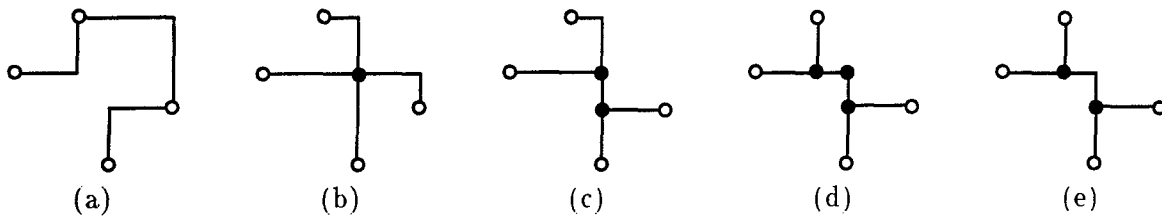


Figure 5: Execution of Iterated 1-Steiner (I1S) on a set of 4 points. Note that in step (d) a degree-2 Steiner point is formed and is thus eliminated from the topology.

Algorithm Iterated 1-Steiner (I1S)
Input: A set P of n points
Output: A rectilinear Steiner tree over P
$S = \emptyset$
While $T = \{x \in H(P) \Delta MST(P \cup S, \{x\}) > 0\} \neq \emptyset$ Do
Find $x \in T$ with maximum $\Delta MST(P \cup S, \{x\})$
$S = S \cup \{x\}$
Remove from S points with degree ≤ 2 in $MST(P \cup S)$
Output $MST(P \cup S)$

Figure 6: The Iterated 1-Steiner (I1S) algorithm.

Although a single 1-Steiner point may be found in $O(n^2)$ time using complicated techniques from computational geometry [14] [28], such methods suffer from large constants in their time complexities, and are notoriously difficult to implement. Thus, a *batched* variant of I1S is preferable (called *batched 1-Steiner* or B1S), which efficiently adds an entire set of “independent” Steiner points in a single *round*, thereby affording both practicality and reduced time complexity [28]. In Section 7 we present experimental results which indicate that Iterated 1-Steiner performs very well in practice.

4.2 Performance Ratio

For small pointsets, we can show that the performance ratio of the Iterated 1-Steiner method is quite good (recall that the performance ratio of a heuristic is defined as the worst-case ratio of the cost of the heuristic solution vs. the optimal solution, taken over all possible pointsets).

Theorem 4.1 *IIS is optimal for $|P| \leq 4$ points.*

Proof: When the optimal Steiner tree for a given pointset P has one or less Steiner points, IIS clearly produces the optimal solution since it examines all Steiner point candidates. For $|P| = 3$, there can be at most one Steiner point. For $|P| = 4$ and when the number of Steiner points in the optimal Steiner tree is two, it is known that only two distinct topologies are possible [24], as shown in Figure 7. A simple case analysis shows that in both of these cases, IIS always selects both of the optimal Steiner points, yielding an optimal solution in either case. $\square$

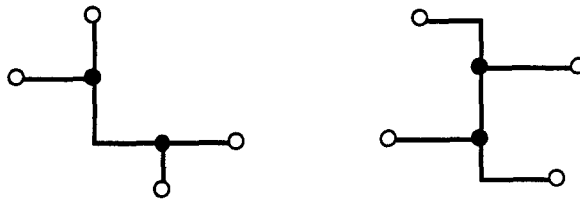


Figure 7: The only possible optimal Steiner tree topologies over 4 points.

In contrast to IIS, many previous MST-based methods are generally not optimal even for $|P| = 4$. Figure 8 shows that a performance ratio approaching $\frac{4}{3}$ is possible (note that IIS performs optimally on this example). On the negative side, the worst-case performance ratio of IIS can reach $\frac{7}{6}$ for 5 points and $\frac{13}{11}$ for 9 points (Figure 9).

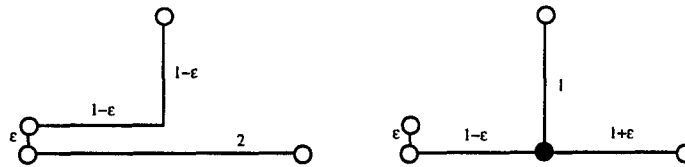


Figure 8: A 4-point instance on which MST-based heuristics perform arbitrarily close to $\frac{4}{3}$ times optimal (left); the (optimal) IIS solution is also shown (right).

4.3 Generalization to Three Dimensions

The routing problem in naval vessel design is inherently three-dimensional, since shipboard interconnect may pass vertically from one deck to another; moreover, the problem is not completely isotropic in the three dimensions, since the cost of going vertically through decks may be higher than routing horizontally along a deck (e.g., routing interconnect from one deck to another may require additional cutting and welding of superstructure, etc.). Fortunately, the Iterated 1-Steiner method generalizes to three dimensions, as well as to the intermediate case where all points lie on L parallel planes (i.e., decks).⁷ Note that in three-dimensional rectilinear space the Hanan candidate set is the com-

⁷In three-dimensional Manhattan space the distance between two points (x_1, y_1, z_1) and (x_2, y_2, z_2) is defined to be $\Delta x + \Delta y + \Delta z = |x_1 - x_2| + |y_1 - y_2| + |z_1 - z_2|$.

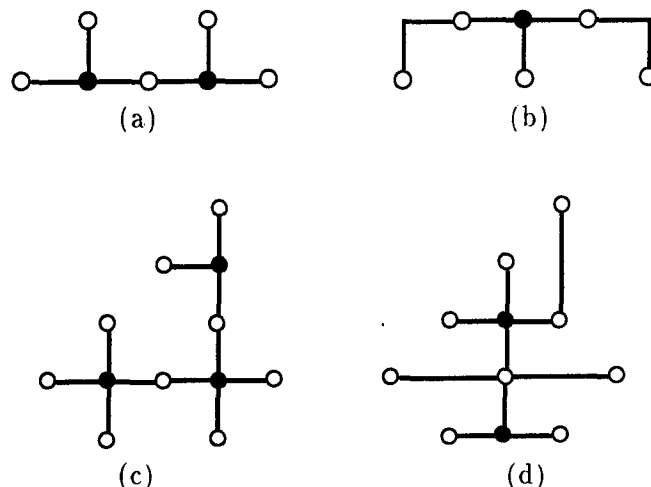


Figure 9: Top: a 5-point example where the IIS performance ratio is $\frac{7}{6}$: the optimal SMT (a) has cost 6, while the (possible) heuristic output (b) has cost 7. At the bottom we see a 9-point example where the IIS performance ratio is $\frac{13}{11}$: the optimal SMT (c) has cost 11, while the (possible) heuristic output (d) has cost 13.

mon intersection of the orthogonal planes passing through each of the input points [45]. Aside from our primary goal of shipboard interconnect optimization, these methods have applications in the multi-layer VLSI routing problem [5] [18] [20], as well as in the design of buildings [44]. Preliminary empirical simulations suggest that our Iterated 1-Steiner algorithm described above is highly effective for three-dimensional Steiner routing, yielding up to 15% average improvement over MST cost in three-dimensional rectilinear space (see Section 7).

We conjecture that our three-dimensional 1-Steiner heuristic has a performance bound of $\frac{5}{3}$ or less. Proving this will solve an important problem open since the 1970's, namely, identifying a 3D rectilinear Steiner heuristic with performance ratio better than that of the best general graph-based Steiner heuristic [49] (which has a performance ratio of $\frac{11}{6}$). In order to prove a performance bound on our 1-Steiner heuristic in three dimensions, it may be possible to generalize Hwang's classification scheme of optimal Steiner topologies to three dimensions, and then apply Zelikovsky's triplet-based approach [49] [50].

4.4 A Fast Implementation

Since the Iterated 1-Steiner method consists of mostly calls to a minimum spanning tree routine, IIS can be implemented very easily (i.e., the algorithm template of Figure 6 can be coded in less than 100 lines of code). However, such a succinct implementation, while running in $O(n^4 \log n)$ -time, may still

not be efficient enough to handle large problem instances. We therefore outline a practical technique for substantially speeding up the execution of the Iterated/Batched 1-Steiner algorithm.

A key observation toward achieving this speedup is that once we have computed an MST over a pointset P , the addition of a single new point x into P can only induce a small *constant* number of changes between the topologies of $MST(P)$ and $MST(P \cup \{x\})$. This follows from the observation that each point can have at most 8 neighbors in a rectilinear planar MST, i.e. at most one per octant [22]. Thus, to update an MST with respect to a newly added point x , it suffices to consider only the closest point to x in each of the 8 plane octants with respect to x . We can further refine this analysis by observing that for dynamic MST maintenance it suffices to search at most 4 regions around each point [16].

Thus we have the following linear-time algorithm for dynamic MST maintenance: connect the new point x to each of its potential neighbors (i.e, the closest point to x in each of the 4 tilted quadrants around x); then, delete the longest edge on any resulting cycle. An execution example of this scheme is shown in Figure 10, and Figure 11 gives a more formal description.

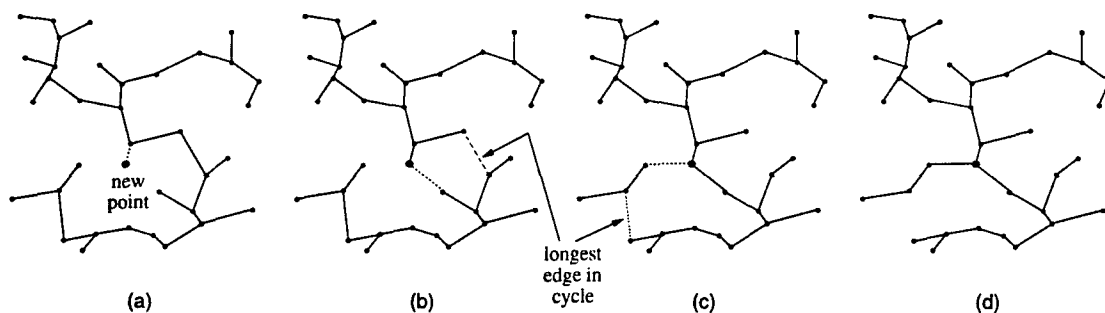


Figure 10: Dynamic MST maintenance: adding a point to an existing MST can be accomplished by connecting the point in turn to its closest neighbor in each octant, and then deleting the longest edge on each resulting cycle (the Euclidean metric has been used for clarity in this example).

For the B1S variant, we perform the following steps in each round: (1) compute in $O(n^3)$ time the MST savings of all $O(n^2)$ Hanan candidates, (2) sort them by decreasing MST savings in time $O(n^2 \log n)$, and (3) march down the sorted list, adding into the input pointset those candidates with “non-interfering” MST savings (at linear time per candidate, using our dynamic MST maintenance scheme described above). This scheme reduces the time complexity of each round of B1S from $O(n^4 \log n)$ to $O(n^3)$, a substantial savings (see Section 7 for specific sample run times). This approach is also quite effective in *three dimensions*, since it is known that every pointset in three-dimensional rectilinear space has an MST with maximum degree of at most 14 or less [39]. Finally, we note that our dynamic MST maintenance method is more practical and easier to implement than similar schemes described in, e.g. [48] and [10].

Dynamic MST Maintenance
Input: A set P of n points, $MST(P)$, a new point x
Output: $MST(P \cup \{x\})$
$T = MST(P)$ For $i = 1$ to $\#regions$ Do Find in region $R_i(x)$ the point $p \in P$ closest to x Add edge (p, x) to T If T contains a cycle Then remove from T the longest edge on the cycle Output T

Figure 11: Linear-time dynamic MST maintenance.

5 Obstacles and Congestion

The shipboard space in which we route is not purely geometric, in the sense that various obstacles such as walls, bulkheads, and fixed equipment may have to be “sidestepped” when deciding on the “best” routes for pipes/ducts/cables. In addition, to keep pipes/ducts/cables accessible for maintenance and repair, areas already congested with numerous previously routed interconnect should not become overcrowded with new routes, provided that such “congestion spreading” can be accomplished at modest cost. The previous discussion ignored such issues for the sake of exposition clarity, but we now address these important practical considerations.

5.1 The Graph Steiner Problem

In order to model congestion and to enable the routing of interconnect around obstacles, we employ a graph version of the Steiner problem. In this graph-based formulation, the cost/difficulty of routing between pairs of point/nodes is represented by edge weights in an underlying graph [27]. Such a graph framework provides a very robust model and can also be used to (1) minimize “jogs” (i.e., 90-degree turns), (2) avoid deck-hopping (i.e., routing the interconnect from one ship’s deck and another), (3) improve accessibility of critical interconnect to routine repair/maintenance, and even (4) enforce safety-related or engineering constraints (e.g., fuel-carrying pipes should not pass through crew berthing quarters, hot and cold water pipes should not be in close proximity to each another, etc.)

Given a weighted graph $G = (V, E)$ with node set V , edge set $E \subseteq V \times V$, and a set of nodes $N \subseteq V$ to be spanned, the graph version of the Steiner Minimal Tree problem seeks a minimum-cost tree in G that spans N . Any node in $V - N$ is a potential Steiner point, and these may be used as intermediate junctions in order to optimize the overall interconnect length. Each graph edge e_{ij} has a weight w_{ij} , and as before, the *cost* of a tree (or any subgraph) is the sum of the weights of its edges. The graph Steiner problem is thus formally stated as follows:

The Graph Steiner Minimal Tree (GSMT) problem: Given a weighted graph $G = (V, E)$, and a set of nodes $N \subseteq V$, find a minimum-cost tree $T = (V', E')$ with $N \subseteq V' \subseteq V$ and $E' \subseteq E$.

The GSMT problem is NP-complete, since the geometric SMT problem is a special case of GSMT (i.e., each point induces a node, and the edge weights are assigned the interpoint geometric distances). In the next subsection we offer an effective heuristic for this difficult problem.

5.2 A Graph Steiner Heuristic

The Iterated 1-Steiner approach can be generalized to solve the Steiner problem in arbitrary weighted graphs, by combining the geometric I1S heuristic with, e.g., the heuristic of Kou, Markowsky and Berman (KMB) [30]. The resulting hybrid method inherits the good empirical performance of the Iterated 1-Steiner method, while also enjoying the error-bounded performance of the KMB algorithm, namely, no worse than twice the optimal. We refer to this hybrid method as the *Graph Iterated 1-Steiner* (GI1S) algorithm. The GI1S template is essentially an adaptation of I1S to graphs, where the “MST” in the inner loop is replaced with the KMB construction to determine the “savings” of each candidate Steiner point/node. Thus, given a graph $G = (V, E)$, a set $N \subseteq V$ of nodes to span, and a set S of potential Steiner points, we define the following:

$$\Delta\text{KMB}(N, S) = \text{cost}(\text{KMB}(N)) - \text{cost}(\text{KMB}(N \cup S))$$

Intuitively, the greedy GI1S heuristic (Figure 12) repeatedly finds Steiner node candidates that reduce the overall KMB cost and includes them into the growing set of Steiner nodes S . The cost of the KMB tree over $N \cup S$ will decrease with each added node, and the construction terminates when there is no $x \in V$ with $\Delta\text{KMB}(N \cup S, \{x\}) > 0$ (i.e., when no improving Steiner nodes remain).

Graph Iterated 1-Steiner (GI1S) Algorithm
Input: A weighted graph $G = (V, E)$ and a set $N \subseteq V$
Output: A low-cost tree $T' = (V', E')$ spanning N (i.e. $N \subseteq V' \subseteq V$ and $E' \subseteq E$)
$S = \emptyset$
While $T = \{x \in V - N \mid \Delta\text{KMB}(N \cup S, \{x\}) > 0\} \neq \emptyset$ Do
Find $x \in T$ with maximum $\Delta\text{KMB}(N \cup S, \{x\})$
$S = S \cup \{x\}$
Return $\text{KMB}(N \cup S)$

Figure 12: The Graph Iterated 1-Steiner algorithm (GI1S).

Given a weighted graph and an arbitrary set of nodes N , a performance ratio of 2 follows for GI1S directly from the KMB bound and the fact that the cost of the GI1S construction can never exceed that of the KMB construction itself (since if no improving Steiner nodes are found, the KMB tree is output by GI1S). If $|N| \leq 3$ (e.g., for three or fewer nodes), GI1S is always guaranteed to find an

optimal solution. Although the worst-case performance ratio of GIIS is the same as that of KMB, in practice GIIS significantly outperforms KMB and may be implemented efficiently [1].

Note that the GIIS template above can be viewed as an *Iterated KMB* (IKMB) construction, and that KMB inside the inner loop may be replaced with any other graph Steiner heuristic, such as that of Zelikovsky (ZEL) [49], yielding an *Iterated Zelikovsky* (IZEL) heuristic. IZEL enjoys the same theoretical performance bound as ZEL, namely $\frac{11}{6}$ times optimal. Preliminary experiments indicate that the average *empirical* performance of these various heuristics is (in order of worst to best): KMB, ZEL, IKMB, and IZEL [1]. Thus, iterating a given Steiner heuristic appears to be an effective general mechanism for improving empirical performance without sacrificing the component heuristic's theoretical performance bounds. Indeed, we conjecture that the performance bound of an iterated heuristic (e.g., KMB or ZEL) in general is strictly smaller than that of the heuristic itself (e.g., we believe that IKMB has a performance bound of *strictly less than 2*).

6 Additional Issues and Formulations

We plan to extend our basic formulations and algorithms to model a number of additional issues that enhance the realism and practicality of our methods. In this section we describe some of these topics.

6.1 Survivability Issues

In order to address fault-tolerance and redundancy issues in shipboard interconnect design, we will explore non-tree interconnect topologies (i.e., where cycles are allowed). While a tree topology suffices to span a set of nodes using a minimum number of edges, non-tree interconnect can offer important safety-critical advantages. For example, if the shipboard power grid is designed so that the deletion of any single edge *does not* disconnect the network, such an interconnect topology can withstand, say, a missile hit to any of its links without disconnecting any part of the ship from the power supply.

Thus, we propose to investigate techniques for producing topologies that correspond to graphs having higher connectivity than do trees (a tree topology becomes disconnected with the removal of any single edge). In particular, a graph is said to be k -connected if the deletion of any $k - 1$ graph edges does not disconnect the graph (although deleting k edges may disconnect the graph). This definition leads to the following natural formulation of the interconnect design problem:

Steiner Minimal k -connected Subgraph (SM k S) Problem: Given a weighted graph $G = (V, E)$, a node set $N \subseteq V$, and an integer k , find a minimum-weight k -connected subgraph $G' = (V', E')$, with $N \subseteq V' \subseteq V$ and $E' \subseteq E$ (or report that none exists).

Note that the SM k S formulation above is different from previous graph augmentation approaches since it allows Steiner nodes to be used to increase the connectivity of the subgraph that we seek (in our ship design applications, the underlying graph G can be, e.g., the three-dimensional Hanan grid). One of our heuristics for minimum-weight 2-connectivity interconnect design is as follows [7]:

- Compute a heuristic Steiner tree T_1 over the given pointset;
- Compute a second Steiner tree T_2 over the leaves of T_1 .
- Output the topology induced by the union of T_1 and T_2 .

This heuristic clearly produces a 2-connected topology, since each pair of the original points is connected by at least two disjoint paths, one in each of the two trees (T_1 and T_2). For this heuristic operating in the rectilinear plane, we can prove a performance bound of $\frac{11}{4}$, and in three-dimensional rectilinear space we can prove a performance bound of $\frac{11}{3}$. We conjecture that both these bounds can be substantially improved. We plan to investigate this heuristic, generalize it to arbitrary k -connectivity, and explore additional heuristics and variants which exploit the freedom inherent in the rectilinear edge-embedding in order to reduce overall interconnect length and/or increase connectivity.

Note that the multiple-connectivity approach and formulation outlined above is general enough to accommodate cases where only a subset of the nodes (the critical nodes) must be multiply-connected, with single-connectivity sufficing for the remaining (less crucial) nodes. The motivation here is to provide the required redundancy while saving interconnect where possible. The concept of non-tree routing is quite useful in the VLSI CAD domain, where adding extra wires to an existing routing tree topology can improve signal propagation delay, reduce signal skew, and increases tolerance to faults due to manufacturing defects or electro-migration [34].

In order to further enhance the survivability and fault-tolerance of shipboard interconnect networks, different links/edges of the topology should be spread as far apart as possible in physical space. In addition, it would be advantageous to avoid positioning long runs of interconnect parallel to one another, since doing so increases the probability that a single bullet/shell following a straight-line trajectory will damage/disconnect multiple adjacent interconnect edges.

This goal of balancing the occurrence of horizontal and vertical interconnect throughout the ship's volume is captured by the following *minimum density* objective, where the *density* of an interconnect topology is defined as the maximum number of edges that are intersected by any line/plane:

Minimum Density Routing (MDR) Problem: Given a set of points, construct an interconnect topology along with a minimum-density physical embedding of the topology into rectilinear space (and minimize cost as a secondary optimization criterion).

Minimizing the "density" corresponds to minimizing parallel runs of interconnect, which also tends to encourage a more balanced utilization of the routing space (i.e., evenly "spreading" the use of horizontal and vertical interconnect segments among different decks can reduce congestion and improve access to interconnect). Thus, our minimum-density formulation constructs interconnect topologies that minimize the number of edges that can be severed by any single cutting plane/line; this enables, e.g., ship-wide power or communication networks that would sustain relatively smaller connectivity damage due to, e.g., an artillery shell or missile impact (in the VLSI domain minimizing the density reduces "crosstalk" between adjacent / parallel wires).

We have developed efficient heuristic constructions for reducing the *average* density of interconnect

trees [2], but much work remains to be done with regard to minimizing the *worst-case* density, and also in producing interconnect embeddings with both cost and density simultaneously bounded by constant factors from optimal. In addition, to further enhance a ship's survivability, we can refine our density formulation above to penalize more heavily parallel interconnect segments which are closer together, while charging only a slight penalty for parallel wires which are farther apart (and are thus less likely to be simultaneously damaged by a single accident or attack).

6.2 Variable-Width Interconnect

Another important practical extension to our basic formulations is to synthesize variable-width interconnect, motivated by the intuition that, e.g., water pipes or air-conditioning ducts that are closer to the main valve/air-blower must be wider in order for the system to operate efficiently.

Conceptually, such a variable-width -edge formulation is similar to the SM&S formulation above, except that edges now also have a *width* associated with them (as well as length). While the edge lengths are given as part of the input, the edge widths are computed during the interconnect design, based on the relevant physical considerations that depend on the interconnect type and the viscosity/flow/consumption requirements (i.e., while an ethernet can be considered to be a "constant-width" network, interconnect which carries fluids/gases such as water, air, fuel, oil, coolant, or steam, must be "tapered" along the paths to the sinks, with the edge widths (i.e., pipe diameters) growing larger as one approaches the source).

The range of the interconnect width may either be continuous or else may be discretized/restricted to a small set of values. The significance of wiresizing in the VLSI domain is now well-established, especially with regard to wiresizing *during* the interconnect topology design [23] (as opposed to wiresizing only *after* the interconnect topology has already been fixed). Such techniques can also be adapted to shipboard interconnect design.

6.3 Multiple Sources

So far in our discussion we have made the implicit assumption that there is only a single "source" node in each interconnect topology. However, this is not always the case. For example, an air-conditioning duct system may have several air blowers attached to it at different locations; similarly, an electrical grid may be connected to a number of generators, located at different positions around the ship (the reasons for such a scheme may be based on safety, redundancy, graceful-degradation considerations, and even military doctrine).

Having multiple sources distributed around the interconnect will affect a number of the issues / formulations discussed above. For example, higher-width interconnect may now have to be designed in the proximity of *all* the sources (in the case, say, of an electrical grid); in addition, higher-connectivity may be required in the vicinity of each source, in order to reduce the probability of inadvertent disconnection of a source from the remaining interconnect topology. Multiple-source optimizations have very recently begun to receive attention in the VLSI CAD literature [6].

7 Experimental Results

We have implemented the Iterated 1-Steiner (IIS) and Batched 1-Steiner (BIS) algorithms, using C in the Sun/UNIX environment. We compared these with the fastest known optimal Steiner tree algorithm (OPT) of [40] on up to 10000 pointsets of various sizes. Since we had difficulty obtaining actual ship data, we tested our methods on random pointsets which were generated by choosing the coordinates of each point independently from a uniform distribution in a 10000×10000 grid (such instances are a standard testbed for Steiner tree heuristics [27]).

Both IIS and BIS have very similar average performance, approaching 11% improvement over MST cost (Figure 13(a)), which is a substantial improvement over previous methods which yield only 9% average improvement over MST. The average number of rounds for BIS is 2.5 for sets of 300 points, and was never observed to be more than 5 on any instance (Figure 13(b)); over 95% of the total improvement occurs in the first round, and over 99% of the improvement in interconnect length occurs in the first two rounds. The average number of Steiner points generated by BIS grows linearly with the number of input points (Figure 13(c)). Figure 14(a) shows the performance comparison of BIS and OPT on small pointsets. We observe that the average performance of BIS is nearly optimal (e.g., for $n = 8$, BIS is on average only about 0.11% away from optimal, and solutions are optimal in about 90% of the cases (Figure 14(b))).

We timed the execution of BIS, using both a naive $O(n^4 \log n)$ implementation and the enhanced $O(n^3)$ implementation which incorporates the efficient, dynamic MST maintenance as described in Section 4.4. Even for relatively small pointsets, the enhanced implementation is considerably faster than the naive one: for $n = 5$, BIS is on average more than twice as fast as the naive implementation, while for $n = 10$ the serial speedup factor approaches 7 (with the serial speedup increasing with problem size). The average running times for various pointset sizes are compared in Figure 13(d). The most time-efficient of the heuristics is BIS, requiring an average of 0.009 CPU seconds for $n = 8$, and an average of 0.375 seconds for $n = 30$. An example of the output of BIS on a random 300-point set is shown in Figure 15.

In three dimensions, we observed that in the limit when the number of planes L approaches ∞ , the average performance of BIS approaches 15% improvement over MST cost, and this performance improves with larger L (Figure 14(c)). Recall that the average savings over MST cost in three dimensions is expected to be higher than in two dimensions, since the worst-case performance ratio is higher as well (i.e., $\frac{5}{3}$ for three dimensions vs. $\frac{3}{2}$ for two dimensions). The number of added Steiner points in three dimensions grows linearly with the pointset cardinality (Figure 14(d)). In all cases, the L parallel planes were uniformly spaced in the unit cube (i.e., they were separated by $\frac{G}{L}$ units apart, where $G = 10000$ is the gridsize). The OPT algorithm of [40] does not generalize to higher dimensions, and thus we were not able to compare the three-dimensional version of BIS against optimal solutions. As in two dimensions, the average number of rounds for BIS is quite small.

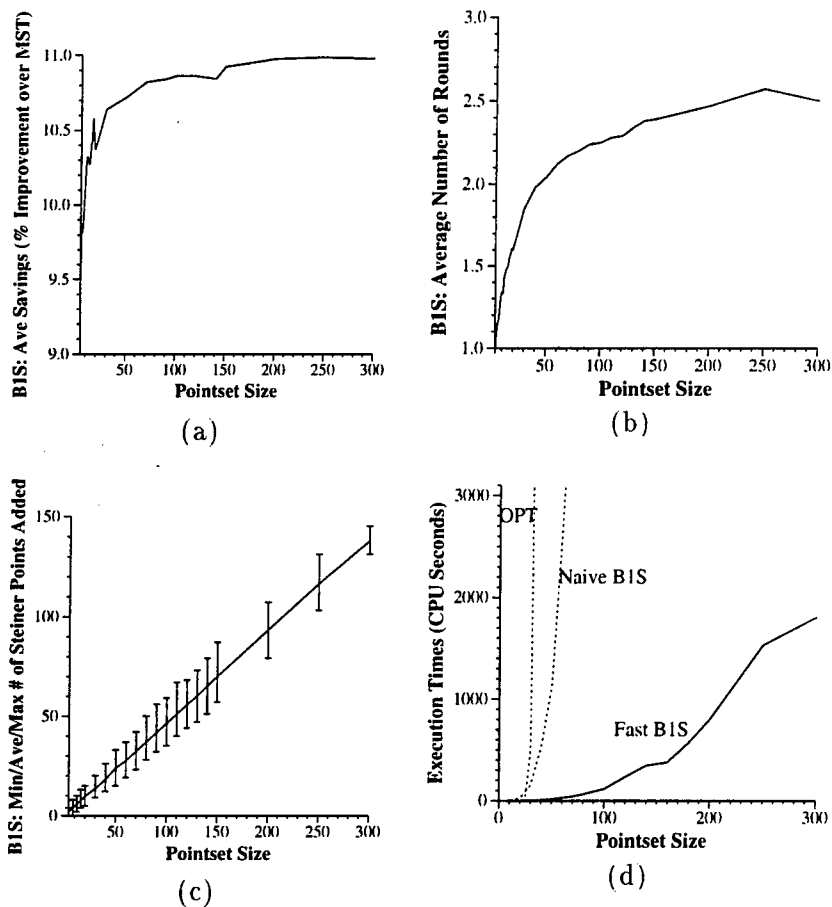


Figure 13: (a) Average performance of B1S, shown as percent improvement over MST cost. (b) Average number of rounds for B1S. (c) Average number of Steiner points induced by B1S (vertical bars indicate the range of the minimum and maximum number of Steiner points added). (d) Average execution times (in CPU seconds) for B1S, for both the naive implementation, as well as the “fast” B1S which uses the incremental MST maintenance scheme (also shown is OPT, the optimal Steiner algorithm of [40]).

8 Conclusions

We started by noting the analogues between chip design and naval vessel design (e.g., both entail *placing* numerous components around a specified region and then *interconnecting* them into a single integrated system). We then addressed interconnect optimization issues in naval vessel construction, using techniques from integrated circuit design. The minimization of total interconnect length is a fundamental objectives in both of these domains, and is effectively modeled by the classical Steiner problem; this enables methods from VLSI CAD to be used in optimizing shipboard interconnect.

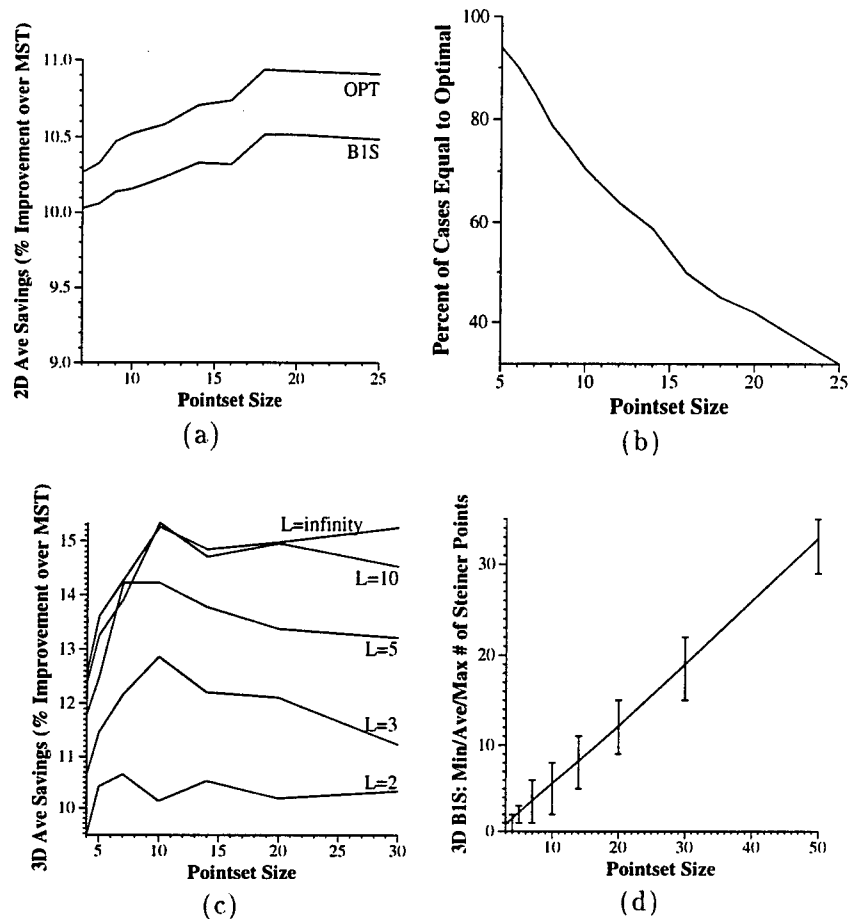


Figure 14: (a) Average performance in two dimensions of B1S, and OPT; note that B1S is less than half a percent away from optimal. (b) Percentage of all cases when the heuristics find the optimal solution (note that B1S yields optimal solutions a large percentage of the time). (c) Average performance of B1S in three dimensions for various values of L = number of parallel planes. (d) Average number of Steiner points added by B1S in three dimensions for $L = \infty$.

We have presented a number of problem formulations which address interconnect optimization as well as other design criteria. With the growing size and complexity of modern naval vessels, such techniques have good implications for ship complexity, weight, procurement cost, reliability, survivability, and ease-of-maintenance. The Steiner heuristic presented here is quite effective in practice, and applies equally well to other domains, including the design of buildings, and even entire installations or bases. Finally, we recommend that these issues and their potential benefits be studied further.

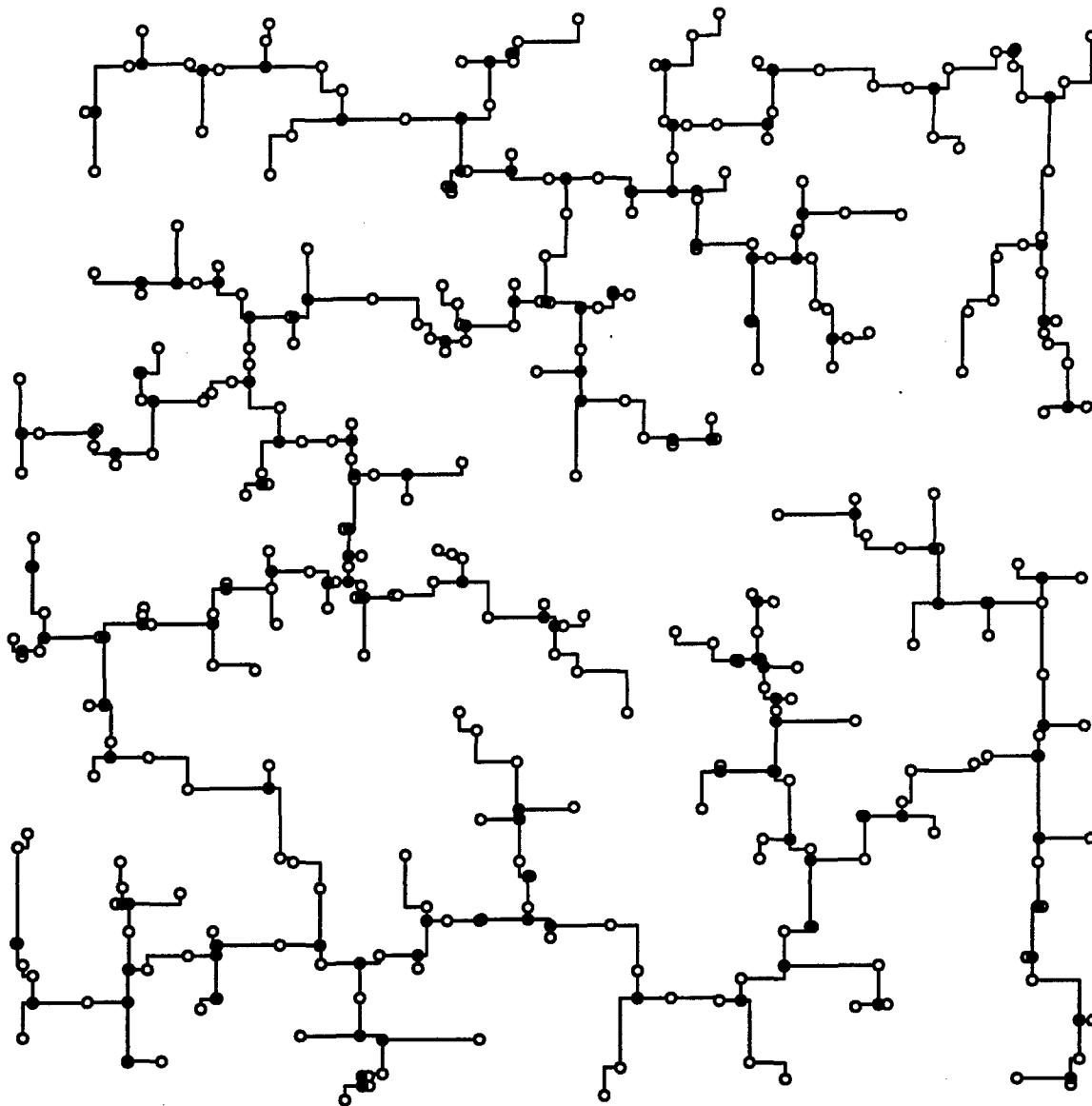


Figure 15: An example of the output of B1S on a random set of 300 points (hollow dots). The Steiner points produced by B1S are denoted by solid dots.

9 Acknowledgements

The author is grateful to Julian Nall of the Institute for Defense Analyses, whom as Director of the Defense Science Study Group (DSSG) has done a superb job in assembling together a fascinating agenda of visits to numerous DoD installations, headquarters, and contractors. Bill Hurley and Nancy Licato also deserve much credit for their excellent help in organizing the various DSSG trips. The generous help, support, and encouragement of General Larry Welch and General Bill Smith of IDA is much appreciated.

The author is also indebted to a number of people who through numerous meetings and discussions helped contribute to the motivations that led to the research described here. These individuals include: Robert Roberts, Colin Hammon and Kevin Saeger of IDA; Eugene Silva, Marc Lipman, Don Wagner, and Ralph Wachter of ONR; Ed Feigenbaum of Stanford University and the Air Force; Admiral Harry Train of SAIC; Paul Kozemchak of ARPA; Admiral Ike Kidd; and Herb York. A special thanks goes to all the members and mentors of the DSSG, for many interesting and enjoyable conversations.

References

- [1] M. J. ALEXANDER AND G. ROBINS, *New Performance-Driven FPGA Routing Algorithms*, in Proc. ACM/IEEE Design Automation Conf., San Francisco, CA, June 1995, pp. 562–567.
- [2] C. ALPERT, J. CONG, A. B. KAHNG, G. ROBINS, AND M. SARRAFZADEH, *Minimum Density Interconnection Trees*, in Proc. IEEE Intl. Symp. Circuits and Systems, Chicago, May 1993, pp. 1865–1868.
- [3] M. W. BERN, *Two Probabilistic Results on Rectilinear Steiner Trees*, Algorithmica, 3 (1988), pp. 191–204.
- [4] M. W. BERN AND M. DE CARVALHO, *A Greedy Heuristic for the Rectilinear Steiner Tree Problem*, Tech. Rep. UCB/CSD 87/306, Computer Science Division (EECS), UCB, 1986.
- [5] D. BRAUN, J. BURNS, S. DEVADAS, H. K. MA, K. M. F. ROMEO, AND A. SANGIOVANNI-VINCENTELLI, *Chameleon: A New Multi-Layer Channel Router*, in Proc. ACM/IEEE Design Automation Conf., 1986, pp. 495–502.
- [6] J. CONG AND P. H. MADDEN, *Performance Driven Routing with Multiple Sources*, in Proc. IEEE Intl. Symp. Circuits and Systems, vol. 1, 1995, pp. 203–206.
- [7] E. COTA-ROBLES AND G. ROBINS, *Graph Connectivity Augmentation Problems: the Steiner Case*. unpublished manuscript, September 1995.
- [8] C. C. DESOUSA AND C. C. RIBIERO, *A Tight Worst Case Bound for the Performance Ratio of Heuristics for the Minimum Rectilinear Steiner Tree Problem*, OR Spektrum, 12 (1990), pp. 109–111.
- [9] L. R. FOULDS, *Maximum Savings in the Steiner Problem in Phylogeny*, J. Theo. Bio., 107 (1984), pp. 471–474.
- [10] G. N. FREDRICKSON, *Data Structures for On-Line Updating of Minimum Spanning Trees*, SIAM J. Comput., 14 (1985), pp. 781–798.
- [11] S. C. GADRE, R. VAIDYANATHAN, AND S. Q. ZHENG, *A Potential-Driven Approach to Constructing Rectilinear Steiner Trees*, in Proc. Great Lakes Symp. VLSI, Kalamazoo, MI, March 1993, pp. 95–99.
- [12] M. GAREY AND D. S. JOHNSON, *The Rectilinear Steiner Problem is NP-Complete*, SIAM J. Applied Math., 32 (1977), pp. 826–834.
- [13] M. R. GAREY AND D. S. JOHNSON, *Computers and Intractability: a Guide to the Theory of NP Completeness*, W. H. Freeman, San Francisco, 1979.
- [14] G. GEORGAKOPOULOS AND C. H. PAPADIMITRIOU, *The 1-Steiner Tree Problem*, J. Algorithms, 8 (1987), pp. 122–130.
- [15] T. C. GILLMER AND B. JOHNSON, *Introduction to Naval Architecture*, Naval Institute Press, Annapolis, Maryland, 1982.
- [16] J. GRIFFITH, G. ROBINS, J. S. SALOWE, AND T. ZHANG, *Closing the Gap: Near-Optimal Steiner Trees in Polynomial Time*, IEEE Trans. Computer-Aided Design, 13 (1994), pp. 1351–1365.
- [17] L. J. GUIBAS AND J. STOLFI, *On Computing all North-East Nearest Neighbors in the L1 Metric*, Information Processing Letters, 17 (1983), pp. 219–223.
- [18] A. HANAFUSA, Y. YAMASHITA, AND M. YASUDA, *Three-Dimensional Routing for Multilayer Ceramic Printed Circuit Boards*, in Proc. IEEE Intl. Conf. Computer-Aided Design, Santa Clara, CA, November 1990, pp. 386–389.

- [19] M. HANAN, *On Steiner's Problem With Rectilinear Distance*, SIAM J. Applied Math., 14 (1966), pp. 255-265.
- [20] A. C. HARTER, *Three-Dimensional Integrated Circuit Layout*, Cambridge University Press, New York, 1991.
- [21] N. HASAN, G. VIJAYAN, AND C. K. WONG, *A Neighborhood Improvement Algorithm for Rectilinear Steiner Trees*, in Proc. IEEE Intl. Symp. Circuits and Systems, New Orleans, LA, 1990.
- [22] J. M. HO, G. VIJAYAN, AND C. K. WONG, *New Algorithms for the Rectilinear Steiner Tree Problem*, IEEE Trans. Computer-Aided Design, 9 (1990), pp. 185-193.
- [23] T. D. HODES, B. A. MCCOY, AND G. ROBINS, *Dynamically-Wiresized Elmore-Based Routing Constructions*, in Proc. IEEE Intl. Symp. Circuits and Systems, London, England, May 1994, pp. 463-466 (Vol. I).
- [24] F. K. HWANG, *On Steiner Minimal Trees with Rectilinear Distance*, SIAM J. Applied Math., 30 (1976), pp. 104-114.
- [25] F. K. HWANG, *An $O(n \log n)$ Algorithm for Rectilinear Minimal Spanning Trees*, J. ACM, 26 (1979), pp. 177-182.
- [26] F. K. HWANG, *An $O(n \log n)$ Algorithm for Suboptimal Rectilinear Steiner Trees*, IEEE Trans. Circuits and Systems, 26 (1979), pp. 75-77.
- [27] F. K. HWANG, D. S. RICHARDS, AND P. WINTER, *The Steiner Tree Problem*, North-Holland, 1992.
- [28] A. B. KAHNG AND G. ROBINS, *A New Class of Iterative Steiner Tree Heuristics With Good Performance*, IEEE Trans. Computer-Aided Design, 11 (1992), pp. 893-902.
- [29] A. B. KAHNG AND G. ROBINS, *On Performance Bounds for a Class of Rectilinear Steiner Tree Heuristics in Arbitrary Dimension*, IEEE Trans. Computer-Aided Design, 11 (1992), pp. 1462-1465.
- [30] L. KOU, G. MARKOWSKY, AND L. BERMAN, *A Fast Algorithm for Steiner Trees*, Acta Informatica, 15 (1981), pp. 141-145.
- [31] M. KRUSKAL, *On the Shortest Spanning Subtree of a Graph, and the Traveling Salesman Problem*, Proc. Amer. Math. Soc., 7 (1956), pp. 48-50.
- [32] J. H. LEE, N. K. BOSE, AND F. K. HWANG, *Use of Steiner's Problem in Sub-Optimal Routing in Rectilinear Metric*, IEEE Trans. Circuits and Systems, 23 (1976), pp. 470-476.
- [33] K. W. LEE AND C. SECHEN, *A New Global Router for Row-Based Layout*, in Proc. IEEE Intl. Conf. Computer-Aided Design, Santa Clara, CA, November 1990, pp. 180-183.
- [34] B. A. MCCOY AND G. ROBINS, *Non-Tree Routing*, IEEE Trans. Computer-Aided Design, 14 (1995), pp. 790-784.
- [35] B. T. PREAS AND M. J. LORENZETTI, *Physical Design Automation of VLSI Systems*, Benjamin/Cummings, Menlo Park, CA, 1988.
- [36] F. P. PREPARATA AND M. I. SHAMOS, *Computational Geometry: An Introduction*, Springer-Verlag, New York, 1985.
- [37] A. PRIM, *Shortest Connecting Networks and Some Generalizations*, Bell Syst. Tech J., 36 (1957), pp. 1389-1401.
- [38] D. S. RICHARDS, *Fast Heuristic Algorithms for Rectilinear Steiner Trees*, Algorithmica, 4 (1989), pp. 191-207.
- [39] G. ROBINS AND J. S. SALOWE, *Low-Degree Minimum Spanning Trees*, Discrete and Computational Geometry, 14 (1995), pp. 151-165.
- [40] J. S. SALOWE AND D. M. WARME, *An Exact Rectilinear Steiner Tree Algorithm*, in Proc. IEEE Intl. Conf. Computer Design, Cambridge, MA, October 1993, pp. 472-475.
- [41] M. SERVIT, *Heuristic Algorithms for Rectilinear Steiner Trees*, Digital Process., 7 (1981), pp. 21-31.
- [42] G. SHUTE, *Worse-Case Length Ratios for Various Heuristics for Rectilinear and Euclidean Steiner Minimal Trees*, Tech. Rep. 90-10, University of Minnesota, Duluth, 1990.
- [43] J. M. SMITH, D. T. LEE, AND J. S. LIEBMAN, *An $O(N \log N)$ Heuristic Algorithm for the Rectilinear Steiner Minimal Tree Problem*, Engineering Optimization, 4 (1980), pp. 179-192.

- [44] J. M. SMITH AND J. S. LIEBMAN, *Steiner Trees. Steiner Circuits and the Interference Problem in Building Design*, Engineering Optimization, 4 (1979), pp. 15–36.
- [45] T. L. SNYDER, *On the Exact Location of Steiner Points in General Dimension*, SIAM J. Comput., 21 (1992), pp. 163–180.
- [46] R. L. STORCH, C. P. HAMMON, AND H. M. BUNCH, *Ship Construction*, Cornell Maritime Press, Centerville, MA, 1988.
- [47] Y. Y. YANG AND O. WING, *Suboptimal Algorithm for a Wire Routing Problem*, IEEE Trans. on Circuit Theory, 19 (1972), pp. 508–511.
- [48] A. C. C. YAO, *On Constructing Minimum Spanning Trees in k -Dimensional Spaces and Related Problems*, SIAM J. Comput., 11 (1982), pp. 721–736.
- [49] A. Z. ZELIKOVSKY, *An $11/6$ Approximation Algorithm for the Network Steiner Problem*, Algorithmica, 9 (1993), pp. 463–470.
- [50] A. Z. ZELIKOVSKY, *A Faster Approximation Algorithm for the Steiner Tree Problem in Graphs*, Information Processing Letters, 46 (1993), pp. 79–83.

I. LASER-ASSISTED FRIEND OR FOE IDENTIFICATION

Clifford Pollock
School of Electrical Engineering
Cornell University
Ithaca, New York

I. INTRODUCTION

The 1991 war in the Persian Gulf brought new attention to an old combat problem: fratricide or "friendly fire" casualties from U.S. or Allied weapons. Fratricide has been noted in every major conflict recorded in history. Famous examples of fratricide include the death of Confederate General Stonewall Jackson during the American Civil War, Operation Cobra in July 1944 where 111 American service men were killed by Allied bombing, and the recent loss of 28 lives in the shooting down of two U.S. helicopters by U.S. forces over Iraq in 1994. After careful evaluation of casualties in all American conflicts, it is now believed that true fratricide rates are significantly higher than the nominal 2 percent rate often cited in military literature. A review by the Office of Technology Assessment of data from World War II, the Korean and Vietnam Wars, and the Gulf War of 1991 shows that actual casualty rates due to fratricide range from 10 percent in Vietnam to 24 percent in the Gulf War [1]. The high fratricide rate of the Persian Gulf War is attributed to the overall low losses in the conflict, the short duration of the conflict, which did not allow troops to gain experience, and the near-absolute dominance of the U.S. forces, which meant that only U.S. rounds were flying through the air. Perhaps if the war had lasted longer, the realized fratricide rate would have declined.

Beyond the needless loss of lives, fratricide is disruptive because it leads to hesitation on the combat field, loss of confidence in leadership, and general degradation of cohesion and morale. Reducing fratricide is a desirable and feasible goal, but totally eliminating it is probably impossible. With our current technology and methods of combat, setting a goal of eliminating fratricide entirely is considered by many to be counterproductive. The restrictive rules or the highly redundant technology that would be required to achieve such a goal would reduce combat effectiveness and increase casualty losses due to enemy action. A balance between no antifratricide measures and too stringent antifratricide measures must be found to minimize total casualties in conflict.

Fratricide results from multiple causes. Based on data from World War II, Korea, and Vietnam, approximately 45 percent of fratricide occurs because of poor coordination, 26 percent as a result of misidentification, 19 percent from inexperience, and 10 percent is attributed to unknown causes [1]. With multiple causes, there are multiple solutions. Improved communication and navigation will help improve deficiencies in coordination.

In our travels as DSSG members to the various Services, we saw many examples of modern communication and navigation gear being deployed with units of all sizes, both to provide better command and control of battle, and to improve coordination for reducing chances that units fire on friendly forces. This is a step in the right direction. In our travels we also saw examples of combat training. One has only to experience one of these exercises to understand and appreciate the sheer chaos of battle. Clearly training is critical to improving fire discipline. To improve identification, better sensors are needed for rapid and foolproof identification of friend or foe in the confusion of battle. In this paper, we will focus on the problem of identification.

II. CURRENT STATUS OF IFF

Present Identify Friend or Foe (IFF) systems rely on a variety of techniques. Fixed and rotary wing aircraft presently use radio transponders called the Mark XII. When queried by a radio frequency signal, an aircraft will respond with an encoded radio signal revealing its identity and position. Although susceptible to enemy corruption via code breaking, spoofing, or jamming, the system has evolved to handle most situations. Next-generation IFF transponders will utilize spread-spectrum techniques to reduce jamming and electronic keys to inhibit code breaking.

Ground-based vehicles to date have relied on visual identification, close coordination of movement, and recently, infrared beacons. There is extensive DoD effort based on millimeter wave technology aimed at developing a ground-based Combat Identification System (CIS). This system will require communication and data processing equipment onboard each vehicle in order to create a mobile wireless network of vehicles [2]. This system is not yet implemented. During the Gulf War, ground vehicles resorted to using infrared beacons (the so-called "DARPA Light" and the "Budd Light") which were mounted to Allied vehicles. These beacons were developed during the war and rushed into service. Although they do not represent secure identification, the short duration of the Gulf War ensured there was no chance for the enemy to imitate these devices. Long-term use of these beacons is probably not a viable option for IFF systems.

Ships now use electronic transponders like the Mark XII devices used in aircraft, which allows identification of friendly forces. Fratricide with naval ships has occurred in the past, but with the Mark XII system, advanced radar and sonar techniques, and fused data from satellites and other intelligence inputs positive identification of most warships is possible under most circumstances.

IFF with submarines is presently not a great concern. Because submarines survive by being unobserved, cooperative IFF schemes are not viable options. IFF with submarines requires passive observation of acoustic signatures, coupled with fused intelligence data from satellites and other sources, and coordination with other friendly submarines. Submarines now patrol well-defined areas. Any other submarine entering the area can be considered to be an enemy and could be open to hostile action by Allied subs or surface ships. In times of conflict, when flexibility of movement is critical,

coordination will be more difficult to maintain. Contact with a coordinating surface ship several times a day is the current plan for IFF among submarines in times of conflict.

III. TYPES OF IFF

A viable IFF scheme should work under all weather conditions and through the smoke and dust of actual battle conditions. Ideally, the IFF system will provide perfect identification (no false alarms), however, as discussed below, this involves tradeoffs that make a perfect system difficult. The IFF system ideally will maintain covertness of friendly forces and not infringe on their effectiveness.

Friend or foe identification requires that the shooter have information about the target. There are many types of IFF. An IFF system can be active or passive, cooperative or non-cooperative. Direct systems are those in which the shooter collects the information himself, for example, by observation. Indirect systems rely on having an observer collect information and relay it to the shooter. An artillery forward spotter is an example of indirect IFF. In this study, we will focus on the direct system.

The shooter may use active or passive means to observe the target. Passive techniques are the most desirable, for example, using energy reflected from the sun or electromagnetic emissions emanating from the target itself to track and identify the target. The passive observer does not transmit any energy and is better able to remain concealed from the target. But passive systems must exploit subtle differences between the signature of friendly and enemy emissions. It is not enough to simply identify a tank; one must identify one type of tank from another under very complex battle conditions. Thus, purely passive systems may require sensitive and complex (i.e., expensive) sensors. It is generally unlikely that one single passive sensor will be adequate for most battlefield scenarios.

Active techniques rely on directing energy at the target and observing a return signal. This can be done with or without the cooperation of the target. Radar is the obvious example of active probing. Even if the target remains passive, the reflected energy will contain information that may allow the shooter to identify the target. Active probing has the disadvantage of being detectable by the enemy. Such risks can be minimized by using intermittent probing, minimal power, and encoding.

A cooperative target will transmit energy back to the shooter in response to the probe. In such a system, the shooter sends a special signal inquiring about the target; if

the target wishes and is properly equipped, it can respond with a signal that identifies it as a friend. This is called cooperative IFF and is the basis of the current Mark XII IFF system used by Allied aircraft. Active systems have potentially longer ranges and can provide relevant information such as speed, location, and identity.

Cooperative systems also provide new dimensions for failure, error, and corruption. To be effective, IFF transponders must be installed and operational on all friendly vehicles during a conflict. It is expensive to equip an entire coordinated force, so this inhibits changes in equipment. Failure can occur through hardware problems or human error. Since proper identification requires recognizing and replying with the proper codes, daily IFF codes must be properly installed in all transponders simultaneously. One reason for the 1994 shoot-down tragedy of two U.S. helicopters over Iraq by U.S. Air Force fighters is that the helicopter IFF transponders were not set for the correct codes. This was a human error. Finally, the codes used by the transponders must be encrypted such that the enemy cannot simply respond with the proper answer and penetrate defenses easily.

Cooperative systems are essentially friend identifiers. They do not positively identify the enemy. If no reply is received to an inquiry, the shooter might assume that the target is an enemy, but he must consider that the target is either neutral or a friend with a malfunctioning transponder. The final classification of such a target will depend on other inputs.

In this paper, we will study the impact of using a laser to effect an identification scheme for ground and air vehicles. The laser provides many degrees of freedom that can be exploited to serve as an IFF transponder and also can provide range finding, standoff biochemical threat detection, and laser radar for aircraft. The key element for this cooperative IFF scheme is the use of narrowband fluorescent materials mounted on the outside surface of vehicles. The fluorescent materials respond when optically excited by a laser at the proper wavelength. Using four different materials with different wavelengths, one could create codes with over 300 combinations, so even if the enemy duplicates the materials, there will still be a first-order coding barrier. The input signal from the laser probe could be evaluated by a neural network, which allows the shooter to train the detection system on friendly and enemy vehicles from different aspects, conditions, and as configurations change.

IV. LASER-BASED COOPERATIVE IFF

An IFF system based on active probing with a tunable laser is described. A shooter would query a target using a laser pulse and, based on the spectral signature of the return from the target, would classify the target as friendly, enemy, or unknown. In the cooperative mode, friendly forces could be tagged with special fluorescent materials to aid in positive identifications. In a noncooperative mode, it may be possible for a tunable laser to determine the hyperspectral content of a vehicle cover and use this information for IFF. Additionally, a tunable laser could be applied to lidar applications, such as standoff detection of chemical agents and contrail detection of stealth aircraft. Data from the laser probing could be analyzed by a neural network that would be trained on both friendly and enemy targets. In this way, the IFF system could gain experience (effectively train itself) and better assist the shooter.

A. BASIC ARCHITECTURE

The basic architecture of a laser IFF sensor is shown in Figure 1. A tunable infrared laser would generate a pulse of approximately 5 ns duration. The pulse would be directed into the optics of a Cassegrainian telescope, where through small adjustments in the directing optics, the exit beam could be made to produce the desired spot size on the target. For initial inquiry, a spot size that illuminates most of the vehicle would probably be desired. This corresponds to a laser beam approximately 10 meters in diameter on target. An electro-optic shutter would direct the laser pulse to the telescope and then switch to direct any returning light to a detector.

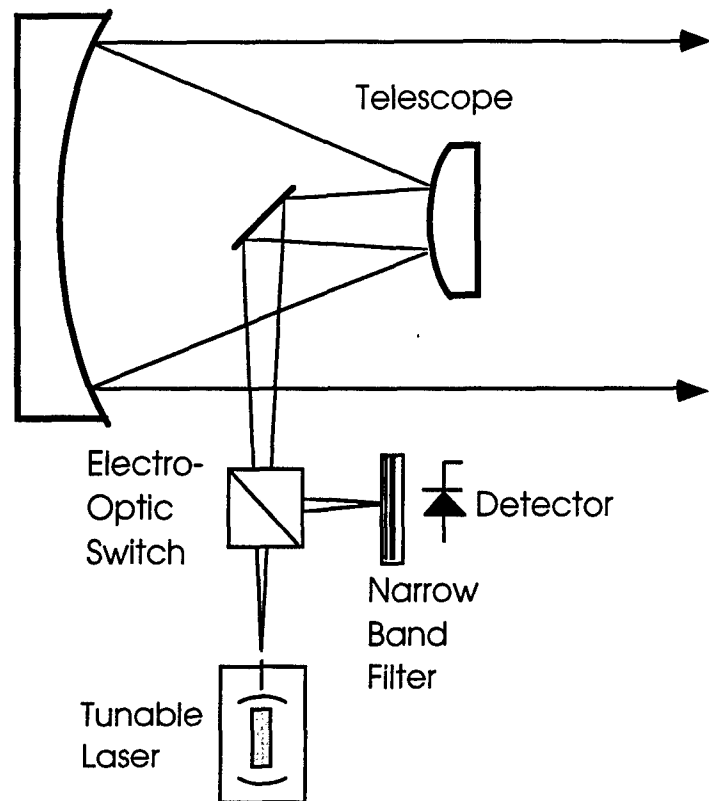


Figure 1. Transmitting and Collecting Optics for Laser Remote Tagging

[The beam from a pulsed laser is expanded and used to illuminate a target. Fluorescence from the target is collected by the telescope and directed by an EO switch through a bandpass filter onto a detector.]

The choice of wavelength for probing is driven by several constraints. Because the laser will be probing friendly forces, it is absolutely essential that it does no harm to the vision of the troops in the vehicle. Light in the 1.4-1.6 μm region is strongly absorbed by water and is considered "eye-safe." Personnel who are illuminated by a 1.5 μm laser beam will not suffer damage to their eyes. But water vapor in the atmosphere will also absorb this wavelength, which could defeat the purpose of this remote sensing technique. The degree to which infrared wavelengths suffer attenuation depends on the relative humidity [3]. For reasons that will become apparent when we calculate background radiation from the sun, a tunable source in the 1.4-1.6 μm region would become a good choice. Tunability allows the operator to optimize the transmission of the beam through the atmosphere and introduces a new degree of freedom for activating spectrally sensitive tagants.

All calculations are based on a laser operating in the 1.45 μm wavelength region. This wavelength will be detectable only with specialized receivers (the Gen1, Gen2, and Gen3 nightvision goggles, will not respond to this wavelength). Furthermore, longer wavelengths suffer less scattering attenuation due to smoke, haze, and other particulates.

For identification of friendly vehicles, a retro-reflecting surface, such as used to paint traffic signs, doped with a special fluorescent materials will be attached to the friendly vehicle. The size of the doped plate would be on the order of 30 x 30 cm (one square foot). Several plates may be needed so that they are visible from any aspect. The plate could be coated with microspheres of doped glass, which serve both to host the fluorescent particles and to reflect the incident light back toward the source. Figure 2 shows a possible construction of such a plate. Small microspheres of doped glass would be suspended in a polymer host, forming an optical composite. Each plate could be coated with spheres containing different dopants to create a unique signature, indicating either a code for the day or a vehicle identification. Calculations on scattering from microspheres indicate that the diameter of the microsphere should be the same as the incident wavelength to maximize backscatter.

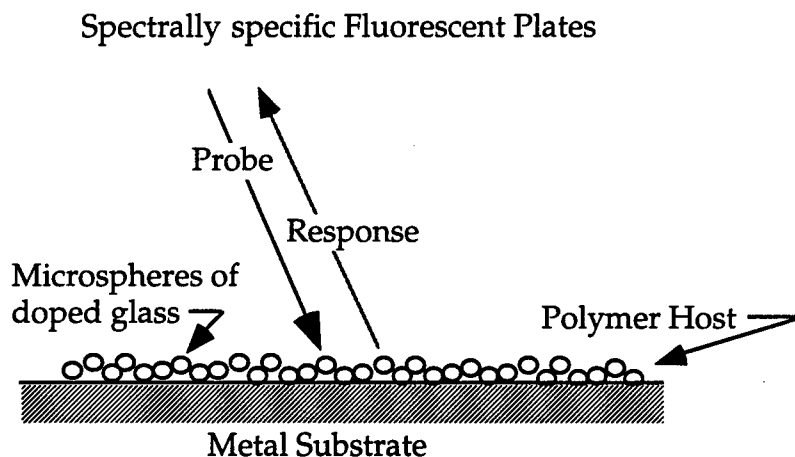


Figure 2. Possible Construction of Spectrally Sensitive Reflecting Plate Used for Laser Tagging

In selecting a dopant, it is desirable to pick an ion which (1) absorbs energy in the spectral range of the probe laser, (2) has a large absorption cross section, and (3) fluorescently radiates the absorbed energy at a longer wavelength, also with a large emission cross-section and a high quantum efficiency. Transition-metal dopants will tend to have large absorption bandwidths and broad spectral emissions. Broad absorption reduces the critical need for precise tuning of the source. Unfortunately, the broad

emission from the plate requires broad pass band filters that will let significant background light in as well. Figure 3 shows the absorption and emission spectrum of two possible broadband materials: $\text{Cr}^{2+}:\text{ZnSe}$, and F_2^+ color centers in KBr.

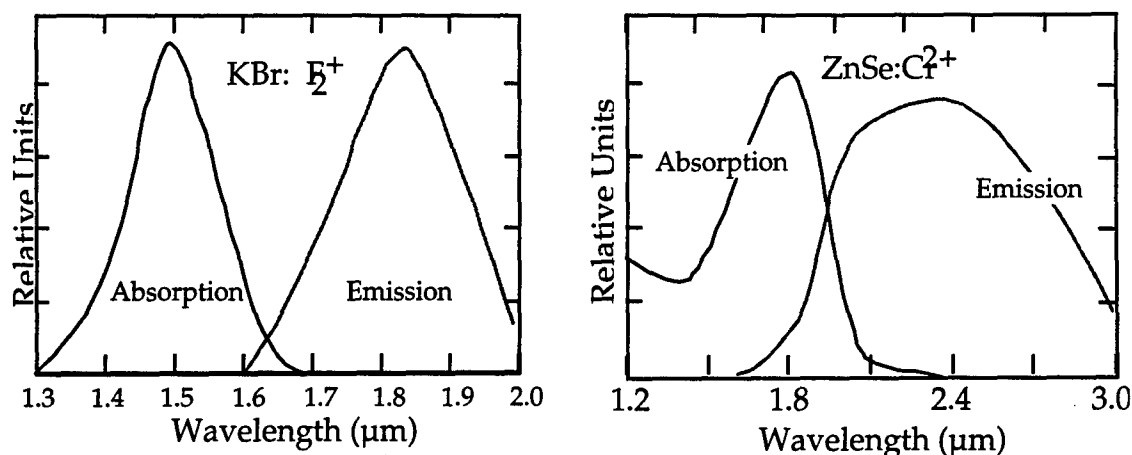


Figure 3. Examples of Two Efficient Fluorescent Sources

Rare-earth dopants, such as Nd or Pr offer much narrower absorption and emission lines. The narrow line requires that the probe laser be tuned to precise wavelengths in order to excite a given transition. However, the spectral emission is restricted to a much narrower linewidth, so spectral filtering at the receiver can be used to reject background light. For use in bright daylight conditions, this would be the desirable situation. All subsequent calculations will assume that a rare-earth ion is used as the active ion ("taggant") on the plates. The emission from the ion can be assumed to fall within a 1 nm spectral bandwidth.

B. SPECTRAL CONSIDERATIONS

The amount of collected signal depends on the reflectivity of the target, the fluorescent efficiency of the plates, the attenuation of the atmosphere, and the distance of the object. In all calculations, we assumed that the size of the probe beam on the target was 10 meters in diameter. This can be accomplished with focusing optics in the telescope. We further assume that all scattered light is emitted in 2π sterads from the surface.

To be operational in daylight, the signal from the target must be greater than the solar background. For the purpose of calculation, we assume the surroundings reflect

approximately 20 percent of incident solar light. This is a slight over-estimation for open sky background, but is realistic for most ground return. Using the same telescope to both send the laser light and to collect the scattered signal, there is no light collected from outside the volume illuminated with the laser. This minimizes the background noise and allows calculations to be made based strictly on the divergence of the probe beam and without concern for aperture size or sensitivity. Using data from The Infrared Handbook [3], the solar irradiance in the spectral region between 1.4 and 1.55 μm is plotted in Figure 4.

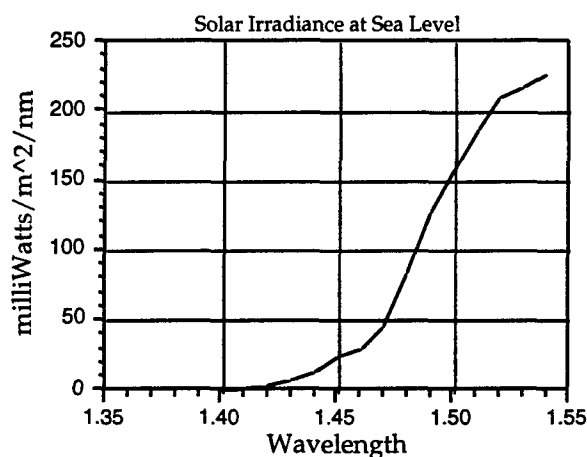


Figure 4. Solar Irradiance at the Equator at High Noon
 [The absorption of the radiation by atmospheric CO_2 and H_2O is responsible for the extinction at 1.4 μm .]

Using a bandwidth of 1 $\text{\AA}$, we need to ensure that the received signal exceeds the background light. Under worst case conditions (high noon on the equator), at 1.45 μm the solar irradiance is approximately 23 $\text{mW}/\text{nm}/\text{m}^2$. The fluorescent signal from the doped panels is proportional to the incident laser pulse energy. As a benchmark, we calculated how much power is needed to provide a signal which is 10 times that due to background radiation from the sun under the worst conditions. Table 1 shows the parameters used to calculate the expected return signal due to fluorescence from the doped panels.

Table 1: Parameters Used to Calculate Spectral Signal due to Laser Probing

Characteristic	Parameters
Size of target	20 m ²
Laser beam diameter	10 m
Photon wavelength	1.45 μ m
Photon energy	0.85 eV
Distance to target	1--20 km
Atmospheric attenuation	10 ⁻² - 10 ⁻³ m ⁻¹
Collection optics (diameter)	10 cm
Fluorescent efficiency	0.9
Reflectivity	0.2
Emitter dome diameter	30 cm
Emitter dome area	0.035 m ²
Background radiation (W/m ² /nm)	23 mJ
Necessary fluence	10 times background

We assumed that a laser pulse with diameter of 10 meters strikes the target. The target carries a reflecting plate that is about the size and shape of a basketball. No matter from which direction the pulse strikes the target, it will see about the same size reflective target. The luminescence from the target is assumed to scatter in a lambertian manner over 2¹ sterads. The telescope used to launch the pulse is used to collect the signal. The field of view is assumed to be brightly illuminated by sunlight, and this leads to a significant solar flux being collected in addition to the fluorescent signal. For reasonable operation, we have calculated the fluorescent signal is 10 times the solar flux. We have included atmospheric attenuation under various circumstances, ranging from "standard clear" to "hazy" [4]. The atmospheric attenuation both reduces the laser probe energy on target and reduces the amount of light that returns to the collection optics. A spreadsheet was used to calculate the collected light under a variety of conditions.

Results of the calculation are plotted in Figure 5. Under standard clear conditions, it is possible to expect reasonable signals under the brightest daylight conditions out to 20 km using laser powers of less than 100 mJ. We choose 100 mJ as a cutoff for pulse energy simply because this is a relatively easy pulse energy to generate using current laser technology, and it does not stress the optics or power supplies to a great extent.

C. DISCUSSION

Experts we spoke to suggested that a distance of 7 km would be a good range. From the data listed in Figure 5, we see that the laser IFF system would satisfy this distance only under clear and medium haze conditions. However, under hazy conditions, standard visibility is only 3 km, which is approximately the same as this IFF system, so in practice the operational limit may not be serious.

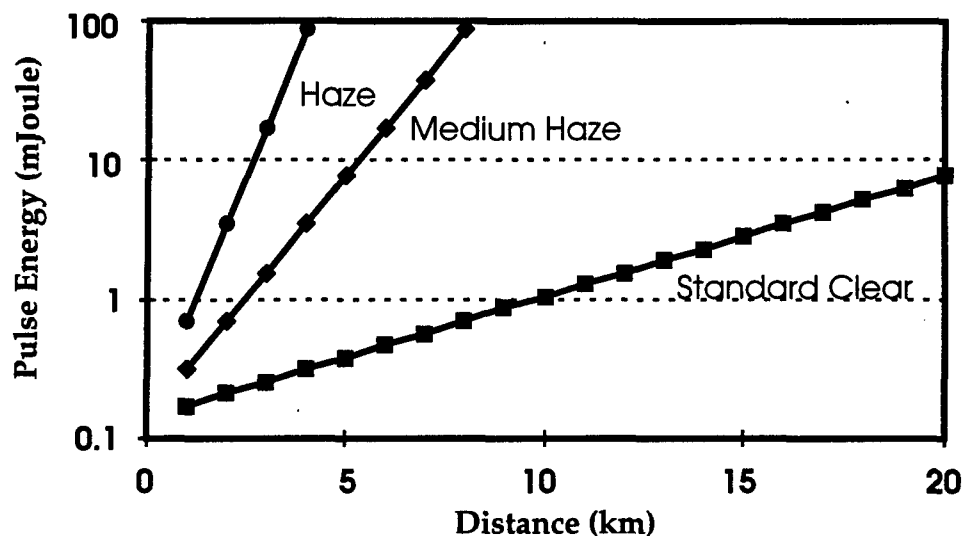


Figure 5. Calculated Energy Required To Achieve an Acceptable Signal Under Bright Daylight Conditions

Such a system could naturally be exploited to extract more information from the probe than simply IFF. By using plates at the front and back of the vehicle, for example, one could determine the vehicle's orientation if the plates had different absorbing and fluorescent wavelengths. Tuning the laser as it emits pulses, first one plate would respond, and then the other. Knowledge of which plate was on the front of the vehicle would be necessary to determine orientation.

In addition to the fluorescent response of the plates, the overall vehicle will reflect light as well. So there will be three signatures to a friendly vehicle: the wavelength-specific response from the two doped plates and the achromatic reflection from the structure of the vehicle. The collected light from the inquiry will be spectrally analyzed to look for this signature. Using neural nets, the system could be trained using friendly vehicles in times of peace, adding more security to the overall signal.

This system has several potential drawbacks. First, it would not be difficult for an enemy to place fluorescent plates on their vehicles as well, making the system ineffective.

To combat this eventuality, a number of different dopants could be used, and daily codes could be sent out to the crews of vehicles to put the desired plates on the front and back of the vehicle. With 5 different dopants and 2 plates on each vehicle, 10 different combinations would provide a unique signature. Adding a third plate would increase this number to 50.

Second, the fluorescent signal will be degraded as the plates become covered with dirt, smoke, and other debris from battle. This would tend to reduce the active distance of the system.

Finally, as we saw at Fort Irwin and at Fort Bragg, one of the first activities in combat is to release smoke between one's own troops and the enemy. This will dramatically reduce the transmission distance of the optical probe but should have little effect on techniques based on radio communication. Therefore, the system could be defeated if the smoke were thick enough. But it must be stated that the smoke screens we saw deployed on battlefields were dense but relatively thin. They will present significant attenuation only over a short range. Smoke will certainly reduce the effective range of the laser probe, but it will still be possible to probe through the smoke if it is not spatially too thick.

V. THE LASER AS A SYSTEM COMPONENT ON A VEHICLE

We have discussed one potential application of the laser as an IFF tool on military vehicles. If the laser is tunable, it could serve several other functions as well. Therefore, it is worth looking at all the potential applications of the laser besides simply IFF.

A. LASER RADAR

Lasers are a source of coherent electromagnetic radiation. In a manner analogous to microwave radar a laser can be used to remotely probe for objects. However, unlike microwave radar, which reflects off of a conducting surface, laser light can respond to changes in scattering of a medium. It is not necessary to have a conducting object in the beam in order to get a signal. Research at Los Alamos [5] has shown that jet aircraft leave aerosols in air up to 1 hour after a flight. A one-dimensional scan of sky can find the remnants of a contrail. Several one-dimensional scans will find direction of the contrail and can point the way to the location of a plane, even a stealth plane.

Theoretically it is possible to use the short wavelength of the laser light to do radar imaging of objects. Techniques have been developed with microwave radar to extract information such as vibrational frequencies. Practically speaking, it is difficult to keep a laser coherent with itself for the times required for radar imaging. It can be done in the laboratory, but there is still much to be done in laser frequency stabilization before this is put in practice in the field.

B. STANDOFF DETECTION

Chemical agents have unique spectral signatures that can be exploited to remotely detect the presence of even trace amounts of these agents. Using a laser as a remote spectrometer, one can analyze the composition of the atmosphere many kilometers away. The laser must be tuned to the known resonance bands of the chemical agent. The major limitation today in implementing this process is that most chemical agents have strong spectral signatures in the ultraviolet region or in the mid-infrared region. Convenient tunable laser sources in these spectral ranges simply do not exist at this time. The laser most suitable for the IFF probing described above operates over the 1.35-1.55 μm region, where there is not a strong unique signature from currently known chemical agents.

VI. CONCLUSIONS

We have described an Identify Friend or Foe system based on optical tagging using a tuned laser to excite fluorescent panels on vehicles. Using different dopants in the panels, a variety of spectral signatures could be created, both to identify type of vehicle or to encode the IFF signature for security. The laser probe is a noncooperative scheme, requiring no action from the targeted vehicle.

The signal derived from the IFF fluorescent panels could be used as one input of several to determine the identity of the target. The probe laser will also reflect off the vehicle, which could provide some hyperspectral information in addition to the fluorescent panels. Effluents such as exhaust gas have been shown to be identifiable to the place of origin of the fuel and to the type of engine. Laser radar processing of the entire return signal could possibly extract information on several dimensions.

All the collected data would be fused to enable a definite IFF determination. The laser response provides input on range, velocity, direction, effluent, and spectral characteristics of the target. Input from a laser tag would further help identify the platform. Fusing this data would be a good task for neural net analysis. Neural nets are in their infancy but have been shown to excel in pattern recognition when exposed to a large variety of information. The laser probe could serve as a very versatile remote sensing tool to feed the neural net. The IFF system could be trained on friendly vehicles, and the data base could be shared between vehicles as they learn to identify vehicles.

The proposed system is feasible today. We have based all laser parameters on currently available systems. It adds one more dimension to the current IFF abilities and, if properly configured, could contribute to the overall information platform via laser radar and remote sensing.

REFERENCES

1. *Who Goes There: Friend or Foe?*, U.S. Congress, Office of Technology Assessment, OTA-ISC-537 (Washington, DC: U. S. Government Printing Office, June 1993).
2. Woody Fox, briefing presented to DSSG, June 18, 1995, Army Night Vision Laboratory.
3. *The Infrared Handbook*, eds. W. Wolfe and G. Zissis, Environmental Research Institute of Michigan, 1978, Section 3-34.
4. C.F. Borhen and D.R. Huffman, *Absorption and Scattering of Light by Small Particles*, New York: John Wiley and Sons, 1983.
5. Nigel Cockroft, et al., *Combat Identification Using Lidar*, Los Alamos National Laboratory.

**J. MILITARY SIMULATIONS USING NEW TOOLS
INVOLVING COMPLEX, ADAPTIVE SYSTEMS**

**Jean M. Carlson
Department of Physics
University of California, Santa Barbara
Santa Barbara, CA**

I. Introduction

The use of computer simulations plays an important role in military decision making. Some of the many applications include: (1) decisions regarding procurement and cutbacks, (2) equipment design, (3) estimation of enemy capabilities, (4) training for and analysis of combat situations, and to a lesser extent (5) developing a better understanding of the origins and patterns of war.^{1,2} Furthermore, the use of simulation by the military is not restricted to war-related objectives. For example, there is also widespread interest in simulating traffic flow³ and the supply and demand of illicit drugs.⁴

In these kinds of simulations, one is necessarily dealing with highly complex, nonlinear, nonequilibrium systems with a large number of degrees of freedom. Our knowledge of other kinds of physical systems in this general category suggests that this is uncertain territory from which features such as chaotic behavior (with the implications of strong sensitivity to initial conditions, and unpredictability), bifurcations (qualitative changes in the macroscopic behavior of the system, which may emerge as a consequence of small changes in one or more parameters), and self-organization (the emergence of spatial structures and temporal fluctuations extending over a broad range of scales), are likely to emerge. In fact, recent examination of the Richardson Equations¹ which were proposed by L.F. Richardson to model arms races between two or more countries, have shown that these equations subject to discrete time dynamics (e.g., if budget decisions are made on an annual basis) exhibit chaotic behavior even in the case of two competing countries.⁵ Furthermore, Richardson also compiled evidence supporting the notion of structure over a broad range of scales in combat by showing that the distribution of the number of people who died as a result of a given war followed a power law:

$$N(s) \sim 1/s^\alpha, \quad (1)$$

where $N(s)$ is the number of wars in which s people were killed, and $\alpha \approx 0.5$ (see Fig. 1). The power law extends from the battles in which roughly 10^3 were killed to the largest battles in which of order 10^7 were killed. A steeper power law describes the smallest events—e.g. murders—in which fewer than 10^3 are killed. Such relationships are reminiscent of other complex systems—earthquake faults, population dynamics and extinctions, and economic markets—in which fluctuations are known to be anomalously large. More generally, when chaos, bifurcations, and self-organization occur, the most meaningful information which one obtains from simulation is statistical in nature, i.e. results are compiled for a large number of runs. However, in some of the largest military simulations, the code is so long and so expensive to run that it is not feasible to run the simulations more than once, and at best individual components are tested for stability with respect to parameter settings. So while the simulations may be very beautiful to the eye, it is not at all clear that the results obtained are robust, in the sense that small changes in parameters or initial conditions may result in significant changes in the outcome.

An additional complication arises due to the fact that many of the components in the system may not be fully specified. The lack of specification arises in part because of the limits imposed by the computational resources, and the fact that including more detail in a model does not always lead to more information about the system—the goal is to accurately capture the essential features of a system with the minimal input ingredients. Of greater concern are the impacts of the intrinsic uncertainties associated with individual components.

There is perhaps no greater uncertainty in modeling than that associated with trying to adequately capture the behavior of the humans in the loop. In some cases, the choices

of an individual are represented by a set of predetermined rational rules. For example, in traffic simulations a driver may have a preset path, say, from home to the office, and the mechanism which guides his driving decisions— speed, lane changes, turns— are a set of deterministic rules related to the path, as well as optimal following distances.³ The assumption is that drivers are all behaving perfectly rationally within the rules that have been defined, with no spectrum of skills or strategies, and that the drivers do not have the flexibility to innovate, learn, or adapt to their environment. Similar simplifications are made in models of combat. The crucial question is whether a more accurate representation of the individual will alter in a substantive way the outcome of large scale simulations.

One indication that these features are likely to qualitatively alter the outcome of simulations comes from model studies of drug smuggling.⁴ In simulations performed at the Rand Corporation, it was found that when the smugglers' rules were preset, the projected success of specific interdiction efforts were much greater than in more realistic models which allowed smugglers to change their strategies. Numerous other examples exist in which qualitative changes in the macroscopic behavior of a system result from including the possibility of adaptation, evolution, and learning of the individual agents.

Of course, there are intrinsic complications associated with trying to model the behavior of individuals in a systematic way. This is to large extent why such features have been left out of the large scale simulations previously discussed. However, in the context of much more simplified models there is beginning to be some progress along these lines developed primarily within the Artificial Intelligence community. The basic idea is to define a simple model composed of interacting agents, e.g. members of an economy who try to maximize their profits or biological organisms or nations which are competing for resources. The key concept of adaptation is one in which the strategies used by the agents to determine their actions are allowed to evolve with time as a consequence of innovation, learning, or adaptation. A collection of simple models governed by these kinds of rules have been shown to exhibit qualitative behaviors reminiscent of the real evolutionary systems. For that reason, the basic concepts and philosophy associated with complex adaptive systems are making their way into a broad class of socio-political arenas. However, it remains an enormous challenge to systematically evaluate the behavior of these models. Progress along these lines will be a crucial step towards developing more realistic and predictive models in the future.

In the remainder of this paper, we summarize some of the recent progress which has been made in modeling complex adaptive systems. We will focus our attention on the Iterated Prisoner's Dilemma model, which is among the systems which has been studied most extensively. In Section II we define the model. In Section III we summarize some of the results which have been obtained, and discuss their implications in the context of military applications. Finally, in Section IV we conclude with a discussion of some future studies which we plan to perform which will more quantitatively assess the implications of these models in areas such as predictability and reaction to specific policy decisions.

We emphasize that our conclusion is *not* that modeling the behavior of interacting adaptive individuals is a solved problem. Rather, on the basis of this study we have found that recent efforts on modeling of individual agents have introduced a new class of interesting dynamical systems, with a wide range of socio-political applications. While much work remains to be done to make these systems sufficiently realistic to be useful for predictive purposes, it is crucial that we realize now that the results of such studies are likely to have a significant qualitative and quantitative impact in a wide array of military applications of computer simulations.

II. Models of Complex Adaptive Systems: The Iterated Prisoner's Dilemma

In this section we discuss some of the kinds of features which are contained in models of complex adaptive systems, and for illustrative purposes we define a particular model referred to as the Iterated Prisoner's Dilemma (IPD). While more realistic models of combat exist, few, if any, have been carefully considered as adaptive systems. In contrast, the IPD has the advantage of being among the most carefully studied adaptive models (of course it is far from being a complete model of evolutionary systems). Similar qualitative features have been observed in a broad class of adaptive systems. A major component of future research in this area will involve development and study of additional model systems for which more systematic analysis is possible.

The Prisoner's Dilemma (PD) is a simple game which is played by two players, and has been used in both experimental and theoretical studies of cooperative behavior. The game was introduced in the 1950's by Flood and Dresher at the Rand Corporation,⁶ and is defined as follows. One imagines that the two players have committed a crime together, and have subsequently been captured by the police. There is not enough evidence to convict them unless one of them confesses. Therefore, if both stay quiet (cooperate: C) both will be released. If only one of them confesses (defects: D), then that player will be released and rewarded, while the other player will be convicted and receive a stiff penalty. If both players confess, then both will be imprisoned, but for a shorter period of time. The outcome of the game can be defined in terms of a payoff matrix (see Figure 2), in which R is the payoff each receives if they both cooperate, if one defects and the other cooperates, the defector receives a payoff of T , and the cooperator receives a payoff of S , and finally, P is the payoff each receives if they both defect. The game is defined by the inequalities

$$T > R > P > S, \quad (2)$$

and

$$2R > T + S. \quad (3)$$

The first inequality is simply an ordering of the payoffs, from best to worst, for a particular individual playing the game: the best case is to defect and have your partner cooperate (you are rewarded and released), next best is to both cooperate (you are released), third best is to both defect (you receive a light sentence), and the worst case is to cooperate when your partner defects (you receive a stiff penalty). The second inequality insures that the players cannot get out of their dilemma by taking turns exploiting one another. In other words, and even chance of exploiting and being exploited is not as good an outcome for a player as always cooperating. A common choice in the literature is to take $R = 3$, $T = 5$, $S = 0$, and $P = 1$, which are the choices we will take here.

If the game is played only once (or a predetermined finite number of times) then it is always in a player's best interest to defect—the expected payoff to the individual averaged over all possible strategies of the opponent is maximized by defecting. However, less trivial behavior is obtained if the number of iterations of the game is unknown (i.e. determined stochastically) or infinite. Such scenarios were investigated first by Axelrod,⁷ and are referred to as the Iterated Prisoner's Dilemma (IPD). As a consequence of conducting extensive computer tournaments, the best strategy he found was a very simple rule based on reciprocity, referred to as "Tit-for-Tat." That is, a player begins by initially cooperating, and then on each subsequent game, repeats the previous move of his opponent. However, this particular strategy is not very robust to variations of the game where, for example, players can make mistakes. More generally, the implication is that cooperation can emerge from a society of individuals each optimizing his own situation. In contrast to the case

where the game is played only once, in the iterated game cooperation results from the fact that players can be expected to continue to meet and thus reap the benefits of reciprocity.

In this paper we will focus on a version of the IPD in which there is a large collection of players. In the simplest case, each player plays every other player in an infinite game, the strategies are updated, and then the game is repeated. This case has been studied by Lindgren, whose results we will summarize here (parameter settings given below correspond to the ones chosen by Lindgren in Ref. [8]). Numerous generalizations of this case have been considered. For example, Lindgren and Nordahl⁹ considered a generalization to a spatial game in which each player is associated with a site on a lattice, and only plays his nearest neighbors.

In the IPD each player has a strategy, which is based on a finite memory of the past consisting of the moves made in previous games played with the opponent. A length m history h_m consists of the opponent's last move a_0 , the player's last move a_1 , the opponent's next to last move a_2 , and so on (where $a_k = 0$ when the player defects, and $a_k = 1$ if the player cooperates):

$$h_m = (a_{m-1}, \dots, a_1, a_0). \quad (4)$$

A history of length m corresponds to the both the player's and opponent's moves in the last $m/2$ games if m is even, and the opponents moves in the last $(m+1)/2$ games and the players moves in the last $(m-1)/2$ games if m is odd. Each of the N players (Lindgren takes $N = 1000$) has a strategy which determines his decision to either cooperate or defect on the next move. A strategy based on a memory of length m is encoded in a *genome* of length 2^m :

$$S = [A_0, A_1, \dots, A_{2^m-1}]. \quad (5)$$

That is, the player chooses the strategy A_k when history h_k is observed. Here the histories are interpreted as binary numbers, and are assumed to appear in numerical order, e.g. h_0 is a vector of length m consisting of all 0's, corresponding to the case in which both players defect every time in the remembered games, and $A_0 = 1$ means that the player will choose to cooperate if he sees that history, while $A_0 = 0$ means that the player will choose to defect if he sees that history. There is also a small probability $p = 0.01$ that a player will make a mistake. This small noise makes it possible to solve the game analytically for the expected payoffs for each pair of players playing the game.⁸

What makes this model a complex adaptive system is the fact that the genomes S describing the players' strategies evolve with time according to standard techniques referred to as genetic algorithms. Three kinds of mutations are allowed, corresponding to point mutations, gene duplications, and split mutations. A point mutation changes a single bit of the genome, e.g., $[00] \rightarrow [01]$. The gene duplication attaches a copy of the genome to itself, e.g., $[01] \rightarrow [0101]$, which increases the memory length of the player by one, but the additional information is not used in the strategy of the player until a point mutation takes place somewhere in the genome. The split mutation randomly removes the first or second half of the genome, e.g., $[0011] \rightarrow [00]$, which reduces the memory length of the player by one. The rates of mutation are all taken to be very small (2×10^{-5} per symbol per genome for point mutations, and 10^{-5} per genome for the other mutations), which is the same order of magnitude as estimated mutation rates in living systems.¹⁰

The game begins with all players having a memory one strategy (i.e. remembering what the opponent's move was in the previous game), chosen randomly with equal probability from the four possibilities: $[00]$ (always defect), $[11]$ (always cooperate), $[01]$ (Tit-for-Tat), and $[10]$ (Anti-Tit-for-Tat, which is the opposite of Tit-for-Tat). All players simultaneously play each other in pairwise infinite games, and (assuming there is some small stochasticity)

the payoff of each player is calculated analytically. The set of strategies is then updated. First, the fraction of players holding each of the current strategies is reset according to a weighting which depends on their payoff in the previous round. In particular, the fitness of an individual w_i is the difference between his payoff s_i and the average s :

$$w_i = s_i - s. \quad (6)$$

The fraction x_i of individuals with a given genotype changes from one round of the game to the next according to

$$x_i(t+1) - x_i(t) = d w_i x_i(t), \quad (7)$$

where the growth rate is set to $d = 0.1$. (Note that in the spatial version of the game all nearest neighbor pairs of sites play the infinite game. At the end of the round, the site adopts the highest scoring strategy among itself and its neighbors.) Second, the strategy of each player is subject to mutation at the rates mentioned above. Natural selection arises as a consequence of the increasing fraction of players assigned to the most successful strategies. It is therefore these strategies which are most likely to mutate, so that out of the potentially infinite state space of strategies the model is most likely to explore the strategies which are close to ones that have previously been successful.

III. Interpretation of the Results in Military Context

While the IPD is clearly an abstract representation of social interactions, even the earliest studies of Axelrod included interpretations of the results in military contexts. Axelrod draws analogies between the IPD and trench warfare between British and German soldiers along the Western Front in World War I, which was often the scene of horrible battles, though between these battles the soldiers (often much to the consternation of the commanders) were known to exercise considerable restraint, adopting a live-and-let-live strategy, analogous to the Tit-for-Tat strategy of reciprocity. The analogy goes as follows:

"In a given locality, the two players can be taken to be the two small units facing each other. At any time, the choices are to shoot to kill [the analog of defection] or deliberately to shoot to avoid causing damage [the analog of cooperation]. For both sides, weakening the enemy is an important value because it will promote survival if a major battle is ordered in the sector. Therefore, in the short run it is better to do damage now whether the enemy is shooting back or not. This establishes that mutual defection is preferred to unilateral restraint ($P > S$ in Figure 2), and that unilateral restraint by the other side is even better than mutual cooperation ($T > R$). In addition, the reward for mutual restraint is preferred by the local units to the outcome of mutual punishment ($R > P$), since mutual punishment would imply that both units would suffer for little or no relative gain. Taken together, this establishes the essential set of inequalities: $T > R > P > S$ (Eq. (2)). Moreover, both sides would prefer mutual restraint to the random alternation of serious hostilities, making $2R > T + S$ (Eq. (3)). Thus the situation meets the conditions [defined in the previous section] for a Prisoner's Dilemma between small units facing each other."

Given this analogy, it is easy to generalize the game to many players. For example, strategies based on reciprocity which were producing positive outcome in one small sector of the battle might spread to another nearby (an analog of the spatial IPD). Analogies between the IPD with the genetic evolution of strategies and Richardson-type models have also been made¹¹ in which the players correspond to countries competing in an arms race.

With such analogies in mind, we next summarize some of the features which are observed when the IPD, as outlined in the previous section, is simulated on the computer. A much more detailed discussion of these results can be found in Ref.[8]. For computational simplicity, the length of the genetic code is limited to be less than or equal to 32

(i.e. at most memory five strategies will be considered). In all of the simulations in the course of its evolution the system passes through long periods of metastability (periods of stasis), which are usually interrupted quite abruptly, and subsequently replaced by either dynamical instability or new periods of stasis.

A typical simulation is illustrated in Figure 3, which shows the fraction of the system which has a given strategy as a function of time for 30000 generations.⁸ Within the typical periods of stasis, examples of coexistence between strategies, exploitation, spontaneously emerging mutualism, and unexploitable cooperation can all be found. For example, in Figure 3 during the first period of stasis, Tit-for-Tat [01] coexists with Anti-Tit-for-Tat [10]. This occurs because [01] suppresses invasion attempts by the always defect strategy [00] which alone would outcompete [10], while [10] suppresses invasions of [11] which alone would outcompete [01].

The average scores s , and the number of genotypes n present in the system are illustrated as a function of time in Figure 4. It is interesting to note that s varies enormously through the course of the simulation, corresponding to wide variations in, say, the overall success of defense efforts. In addition, while the number of strategies tends to increase, the behavior is not monotonic. It appears that more mutations survive in periods just before the transitions from metastability, and during the periods of unstable behavior, suggesting that evolution (the development, and success rate of new strategies) primarily occurs during these periods.

Note that in a pattern of evolution reminiscent of both biological organisms and the arms race, the system tends to be dominated by strategies of increasing complexity. Furthermore, while a given strategy may dominate the system for some period of time it will typically perform poorly if inserted artificially in some different time period, or a different simulation of the system, when the other strategies it would be competing against are different. This is another general feature of complex adaptive systems—the fact that there are many possible metastable states, and the exact one a particular system evolves to will depend on the details of the history. On the other hand, general features (e.g. and evolution towards increasing complexity, composed of periods of stasis, interrupted abruptly by rapid change) will be common to essentially all of the realizations of the system.

IV. Conclusions

Concepts associated with complex adaptive systems provide an inspiring paradigm for a wide array of socio-political systems. Features such as periods of stasis, rapid transitions to dynamic instability, large extinctions, and evolution of cooperative strategies have been observed. However, so far the analogies which have been made are mostly qualitative, and few, if any, of the models which have been developed are sufficiently realistic to be used for predictive purposes. Of course, this is a relatively new field, and developing realistic models takes time. At the most basic level it is certainly useful to know that these effects can be important, and that simple rules for interacting individuals can lead to nontrivial collective effects.

The question remains, however, of how we might best make the transition from philosophical concepts to predictive models. The answer clearly lies in a coupled effort to push the simple models harder, and to develop more realistic models which better represent specific situations and which can be compared directly to existing data. To some extent development of more realistic models has already begun—for example, in the work mentioned previously related to the smuggling of illicit drugs,⁴ as well as efforts to model the interaction of traders in an artificial stock market.¹² The danger in all of this is similar to the danger associated with simulations in general. That is, one must be careful not to put too much emphasis on the particulars of individual realizations of the simulation—e.g.,

what cause led to what effect in one particular run— and instead to focus on statistical information accumulated on the basis of many trials, and to test this statistical information against data from real systems. In addition, because no model is entirely realistic it is equally important to compare the behaviors of different models. This is where the first half of the effort comes in— i.e. the systematic analysis of simple models. Only by studying the detailed behavior of simple models can we begin to determine how different input ingredients of a model influence the outcome. The development of new, analytically tractable, models is likely to play an important role in these developments.

It is easy to overlook these crucial stages in model development. With today's technology, simulations can be made to look very realistic, and the focus is often on getting the code up and running on a tight budget and with an impending deadline. As a result, it seems likely that detailed tests for instabilities and chaos in the models, or sensitivity to unknown parameters such as behavior of individuals are typically overlooked. However, as simulations begin to play more and more of a role in military decision making it is crucial that these kinds of questions receive adequate study.

Some statistical tests have been performed for the stock market model,¹² in which a price index is generated when traders buy and sell according to strategies which evolve with time. For example, it was found that the model produces correlations between fluctuations in the stock price, and the price itself, a known phenomenon in real stocks which is not described by traditional economic theories. It is a more difficult task to identify the kinds of tests that might be performed on models of combat. However, in addition to the more basic problems of model development discussed above, some more detailed questions we would like to investigate further in the context of the IPD and other systems are discussed briefly below. Here we have focused on features which have some military implications. In some instances we expect there will be data to compare with.

(1) Distribution of the mean rate of defection. If the intention to kill corresponds to defection in the combat analog of the IPD discussed in the previous section, then the rate of defection should be the quantity most closely related to the number of fatalities. It would be of interest to determine if the IPD generates a scaling law analogous to that illustrated in Figure 1. It would also be of interest to compare results for the non-spatial IPD, with the analogous results for the spatial model.

(2) Predictability of catastrophic events. In the analogy between the IPD and arms races, the periods of dynamics instability which follow periods of stasis might correspond to the collapse of one or more superpowers. We plan to determine whether or not these periods of turmoil can be predicted. On the coarsest scale, it would be of interest to look at the distributions of the time durations of stasis and of dynamic instability to determine the breadth of the distributions, and whether features of the distributions (e.g. the mean and the width) vary with time as the system evolves. Of more immediate predictive use is the identification of precursors to the catastrophes. Lindgren has indicated that the number of genotypes n in Figure 4 increases as the period of stasis draws to a close. If this is true, then an increase in n is one possible precursor. It remains to be seen how well (quantitatively) this can be used for prediction, and whether other such precursors exist. For example, Lindgren also indicates that destabilization is usually the result of a slowly growing group of mutants (rebels) reaching a critical level.

(3) Effects of externally imposed policy decisions. All of the structure which is present in the IPD is the result of the individual agents making decisions based on optimizing their individual payoffs. Of course, in realistic military applications many policy decisions are imposed externally by, e.g., Command and Control. (Of course, another alternative is to think of a hierarchically composed model, reminiscent of military structure, which ultimately we plan to do. However, the first step is to see how decisions at a higher

level effects a population of individuals.) Of particular interest are interventions which might more rapidly stabilize periods of dynamic instability. Society is much more skilled at forming alliances than it is at dissolving alliances. It would be interesting to consider the effects of allowing alliances to restructure more freely (e.g. to increase the growth rate d in Eq. (7)) during times of dynamic instability.

Finally, it is noteworthy that the concept of complex adaptive systems is featured in "Command and Control," a Marine Corps Concept Paper currently under development.¹³ The paper points out that there are two conflicting points of view regarding command and control. The first is to pursue certainty (that is, to be in "complete control"). The second is to cope with uncertainty. Our knowledge of complex, chaotic systems suggests that the first option is hopeless, and it is in fact the second strategy that the paper advocates. By centrally outlining broad objectives, but then subsequently decentralizing decision making, the proposed doctrine shares the basic philosophy of complex adaptive systems, in which systems naturally evolve to efficient strategies based on the adaptation of the individual agents rather than by detailed mandates imposed by an external source. In the case of the Marines, the basic objective is to reach effective military decisions faster than an enemy, in conflicts which take place in different settings and over a broad range of scales. We close with a quote from this paper, outlining certain elements of this philosophy:

"The fundamental point [presented in this paper] is that any military action, by its very nature as a complex adaptive system, will exhibit messy, unpredictable and often chaotic behavior that defies orderly, efficient, and precise control. Our approach to command and control must find a way to cope with this inherent complexity... This view recognizes that effective command and control must be sensitive to changes in the situation. This view sees the military organization as an "open" system, interacting with its surroundings (especially the enemy), rather than as a "closed" system focused on internal efficiency. An effective command and control system provides the means to adapt to changing conditions. We should liken the military organization to a predatory organism—information seeking, learning, and coping in its quest for survival and success—rather than to some "big green machine." Like a living organism, a military organization is never in a state of stable equilibrium, but is instead in a constant state of flux—continuously adjusting to its surroundings... Command and control is not so much a matter of one part of an organization "getting control over" another, as something which connects all the elements together in an cooperative effort... [and] is thus fundamentally an activity of reciprocal influence."

REFERENCES

- [1] L.F. Richardson was a pioneer in the application of mathematical models and statistical concepts to combat and arms races. See, L. F. Richardson, *Arms and Insecurity*, The Boxwood Press, 1960, and Ref.[2].
- [2] L. F. Richardson, *The Statistics of Deadly Quarrels*, The Boxwood Press, 1960.
- [3] Darrell Morgenson, (private communication).
- [4] J.P. Caulkins, G. Crawford, and P. Reuter, "Simulation of Adaptive Response: A Model of Drug Interdiction," *Mathl. Comput. Modelling* **17**, 37-52 (1993).
- [5] G. Mayer-Kress, "nonlinear Dynamics and Chaos in Arms Race Models," *Proc. Third World Conference: Modelling Complex Systems*, (Lui Lam, Ed.), San Jose, 4/12-13/91.
- [6] M. M. Flood, "Some Experimental Games," (Report RM-789-1). Santa Monica, CA: The Rand Corporation, 1952.
- [7] R. Axelrod, *The Evolution of Cooperation*, New York: Basic Books, 1984.
- [8] K. Lindgren, "Evolutional Phenomena in Simple Dynamics," *Artificial Life II*, SFI Studies in the Sciences of Complexity, Vol. X. Ed. C.G. Langton, C. Taylor, J.D. Farmer, and S. Rasmussen, Addison-Wesley, 1991, pp. 295-312.
- [9] K. Lindgren and M. G. Nordahl, "Cooperation and Community Structure in Artificial Ecosystems," *Artificial Life* **1**, 15-37 (1994).
- [10] D. J. Futuyma, *Evolutionary Biology*, (2nd. ed.). Sunderland: Sinauer Associated, 1986.
- [11] S. Forrest and G. Mayer-Kress, "Genetic Algorithms, Nonlinear Dynamical Systems, and Models of international Security," in *Handbook of Genetic Algorithms*, Ed. by L. Davis, 1991, pp 166-185.
- [12] R.G. Palmer, W. B. Arthur, J. H. Holland, B. LeBarron, and P. Tayler, *Physica D* **75**, 264 (1994).
- [13] "Command and Control," A U.S. Marine Corps Concept Paper, C41 Division, headquarters U.S. Marine Corps (Draft).

FIGURES

- (1). Scaling laws for combat.
- (2). Payoff matrix for the Prisoner's Dilemma.
- (3). A typical simulation of the evolutionary IPD.
- (4). The average score and the average number of genotypes for the simulation shown in Figure 3.

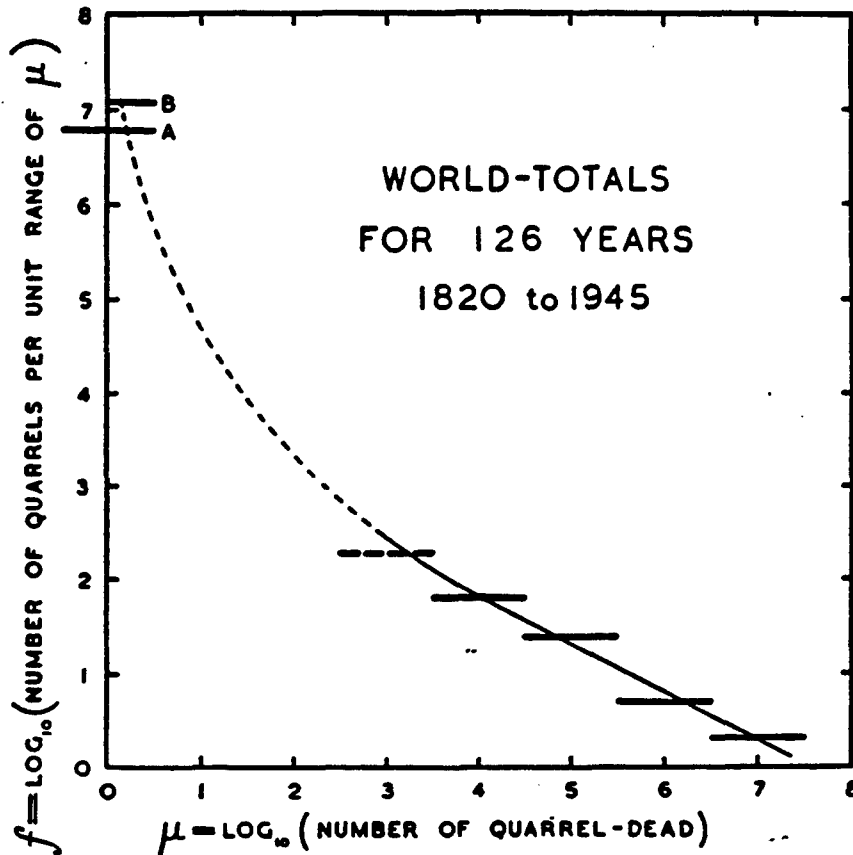
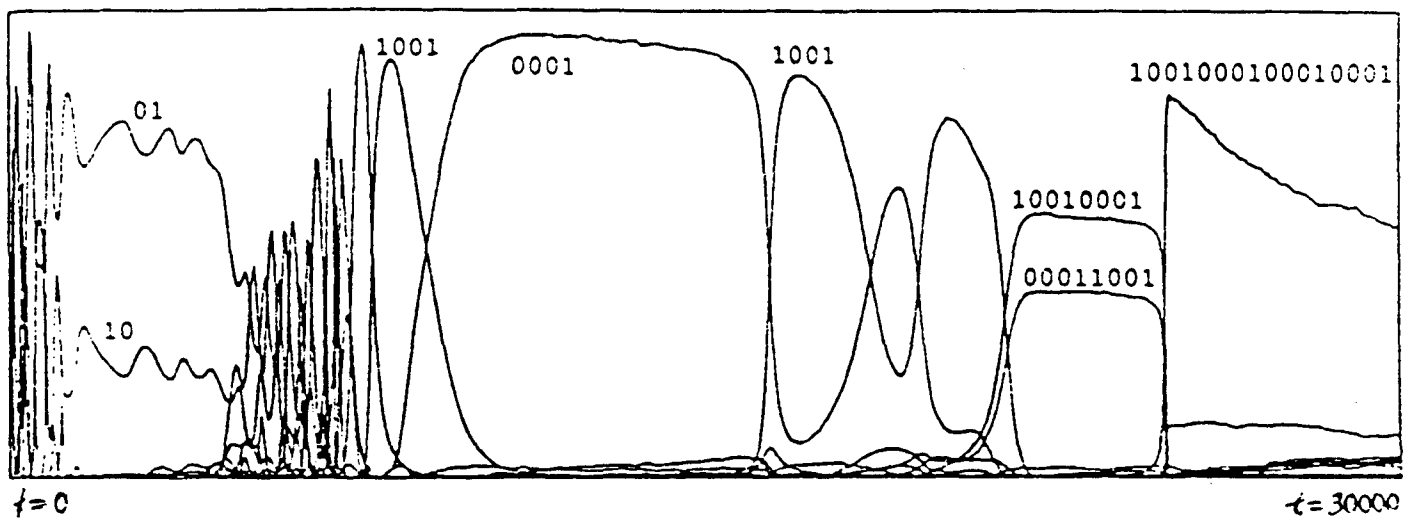


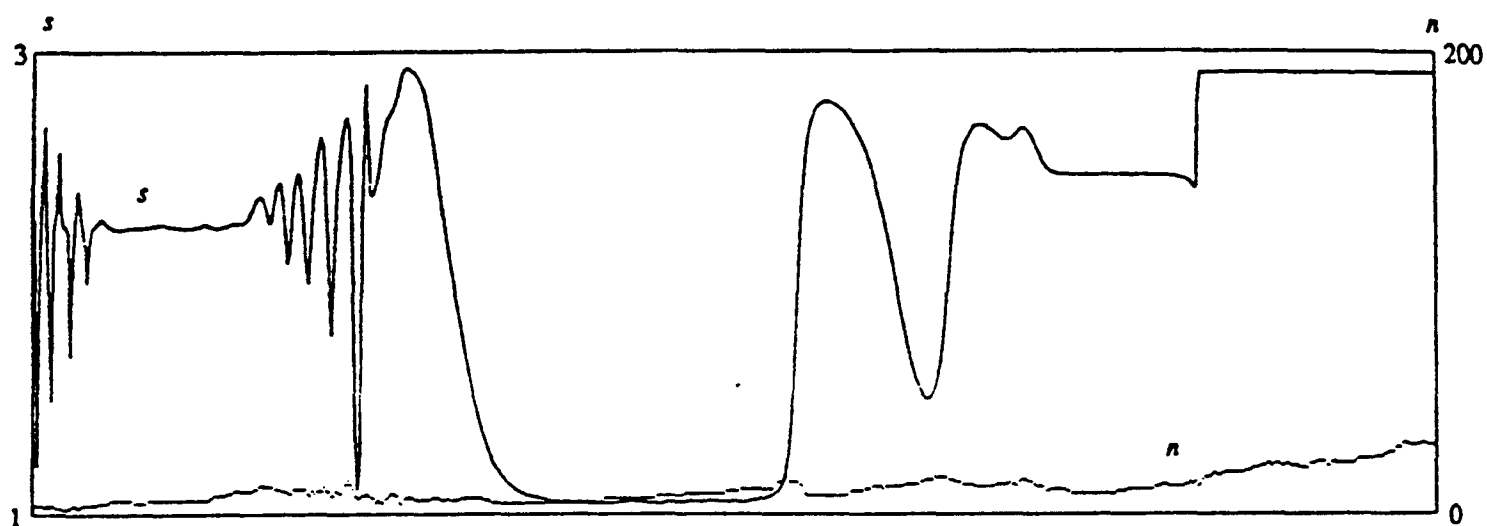
Figure 1. Scaling Laws for Combat

The payoff matrix we use in the Prisoner's Dilemma is the same as the one used by Axelrod. The pair (s_1, s_2) denotes the scores to players 1 and 2, respectively

		Player 2	
		Cooperate	Defect
Player 1	Cooperate	(3, 3)	(0, 5)
	Defect	(5, 0)	(1, 1)



The simulation of Figure 1 is continued for 30000 generations, showing that four periods of stasis appear in the evolution. The oscillations observed in Figure 1 are damped and the system reaches a period of stasis with coexistence between [01] (TFT) and [10] (ATFT). This stasis is punctuated by a number of memory 2 strategies, and after a period of unstable behavior the system slowly stabilizes when the strategy [1001] increases in the population. This strategy cooperates if both players performed the same action last time. For two individuals using this strategy, an accidental defection by one of the players leads to both players defecting the next time, but in the round after that they return to cooperative behavior. Thus, the strategy [1001] is cooperative and stable against mistakes, but it can be exploited by uncooperative strategies. Actually, one of its mutants [0001] exploits the kindness of [1001], which results in a slow increase of [0001] in the population. This leads to a long-lived stasis dominated by the uncooperative behavior of [0001]. A slowly growing group of memory 3 strategies is then formed by mutations, and the presence of these species causes the fractions of the strategies [0001] and [1001] to oscillate. Two of the memory 3 strategies, $M_1=[10010001]$ and $M_2=[00011001]$, manage to take over the population, leading to a new period of stasis. Neither M_1 nor M_2 can handle mistakes when playing against individuals of their own kind, but if M_1 meets M_2 they are able to return to cooperative behavior after an accidental defection. This polymorphism is an example of mutualism which spontaneously emerges in this model. The stasis is destabilized by a group of mutants, and we get a fast transition to a population of memory 4 strategies which are both cooperative and unexploitable.



The average score s (continuous line) and the number of genotypes n (broken line) are shown for the simulation of Figure 2. When the exploiting memory 2 strategy dominates the scene, the average score drops close to 1. The last stasis, populated by the evolutionary stable memory 4 strategies, reaches a score of 2.91, close to the score of 3 achieved by the best strategies in a noise-free environment. Before the transitions and in the periods of unstable behavior, it appears that there are more mutants that survive and the number of genotypes increases, suggesting that most of the evolution takes place in these intervals.

K. EXPLOITING AND CONTROLLING CAVITATION

Michael J. Shelley

Courant Institute of Mathematical Sciences

New York University

New York City, NY

EXPLOITING AND CONTROLLING CAVITATION

1. INTRODUCTION

The phenomena of cavitation is usually treated in technological design as a feature of highspeed hydrodynamics that should be avoided. The growth and collapse of cavitation bubbles near walls (such as in rotors, propellers, pumps, valves, ...) is a central factor in structural fatigue and surface roughening (see the many relevant articles in Furuya (1968)). Cavitation bubbles can also induce resonant response in structures – the near loss of the Glen Canyon dam several years ago is thought to be such an example (see also Li *et al* (1989)). In certain settings, such as in the flow at propeller or hydrofoil tips, the occurrence of cavitation is especially unwelcome due to the strong acoustic signature. One also suspects that the often random space/time distribution of cavitation, say about a body moving rapidly through water, can make control very difficult.

While the control of cavitation has traditionally focused on the modification of structural design so as to avoid its occurrence, other approaches have emerged. For example, Fruman & Aflalo (1989) have demonstrated that the addition of drag-reducing polymers to the flow tends to inhibit the inception of cavitation. In work perhaps more relevant to my considerations here, various researchers (for example, Wade & Acosta (1966), Kato *et al* (1979)) have shown that the creation of a coherent cavitation region that covers a substantial part of the cavitating body – i.e., is *supercavitating* – can lead to substantial reductions in the body's drag coefficient. As the authors also point out, the creation of a large cavitation region is by itself not sufficient for drag reduction, but rather follows from careful experiment and design. On a related point, if a body is moving sufficiently fast, it will naturally create cavitation bubbles, and the creation of a large cavitation bubble about such a body must result in special issues concerning the stability and control of both the bubble, and of the body.

The structure of this piece is as follows. In Section 2, the assumptions on the relevant physics is discussed. In Section 3 is given a brief and rudimentary description of the phenomena of cavitation. In Section 4 I discuss what I believe might be some of the relevant physics regarding high-speed, supercavitating torpedoes – such as have been reported recently in the open press (see, for example, *Popular Science*, October 1995, p.

34). My ruminations are made in the absence of any direct intelligence on the matter, but arise rather from knowledge planned from conversations with several government scientists who had some indirect knowledge on these matters. Some concluding remarks are given in Section 5.

2. ASSUMPTIONS

I shall assume that the fluid – water – can be treated as incompressible, which is entirely reasonable given the very small compressibility of water. One condition that must be satisfied in order for compressibility effects to be ignored is that

$$M^2 = \left(\frac{U}{c}\right)^2 \ll 1 ,$$

where $M = U/c$ is some characteristic speed of the flow, and c is the sound speed of the fluid medium. For water, we have $c \approx 1500\text{m/sec}$. Let's say that the body is moving at 150m/sec (approximately 300 knots, this is the number given in the *Popular Science* report), then $M^2 = 0.01$.

The equations of motion for an incompressible, constant density, Newtonian fluid are the Navier-Stokes equations (in non-dimensional form):

$$\frac{\partial \mathbf{u}}{\partial t} + (\mathbf{u} \cdot \nabla) \mathbf{u} = -\nabla p + \frac{1}{Re} \Delta \mathbf{u} , \quad (1)$$

$$\nabla \cdot \mathbf{u} = 0 \quad (2)$$

where $\mathbf{u}$ is the fluid velocity, and p is the fluid pressure. These two equations are respectively momentum and mass conservation statement. An important quantity in incompressible flow is the *vorticity*, $\boldsymbol{\omega} = \nabla \times \mathbf{u}$. The vorticity gives the local rate and direction of fluid rotation.

For a high-speed flow in water, the effect of viscosity can often be ignored. The measure of viscous effects lies in the size of the Reynolds number

$$Re = \frac{UL}{\nu} ,$$

where ν is the kinematic viscosity of water, and U and L are characteristic length and velocity scales. For large Re , viscous effects will be weak, at least outside of (very thin) boundary layers and wakes formed by shedding. In water, we have $\nu = 10^{-6}\text{m}^2/\text{sec}$, and so by taking $L = 10\text{m}$ and the above value for U , we have $Re > 10^9$. By these considerations, the last term of the Navier-Stokes equations can be neglected, giving the so-called Euler equations.

The flow within the cavitation bubble – which is gaseous – is an entirely different matter, and one which I shall mostly ignore. In an atmospheric flow at speeds as given above, compressibility effects begin to be significant. There may be other things going on within the bubble – mixing of gas bubbles, heating, etc. – that would have to be looked at very carefully.

3. ELEMENTS OF CAVITATION

Within an incompressible fluid, it is entirely allowable and possible for the fluid pressure to become very small, zero, and even negative and so put the fluid in that region under a negative tension. If the pressure falls below the vapor pressure of water (this is typically a very small fraction of an Atm), voids or bubbles can form and expand within the fluid. These voids are called *cavities* or *cavitation bubbles*, and occur typically where there are already “seeds,” or small bubbles, from which the cavities can grow.

Where do cavitation events usually occur? Obviously one needs to look for regions of low pressure. There are several typical situations in high Reynolds number flow:

(1.) In high-speed flow around bodies, the fluid can be separated into three regions: (a) The thin boundary layer near the body, necessary to satisfy the no-slip boundary condition. (b) The wake behind the body, produced by the flux of vorticity from the body and its transport downstream. (c) The remaining fluid, which is mostly irrotational compared to the other two vorticity-bearing regions, and that makes up the majority of the fluid volume.

Considering steady flow in the frame moving with the body, then by using the “rotation” form of the Euler equations it is easy to show that in the irrotational region, the pressure p and the velocity $\mathbf{u}$ satisfy

$$\Delta p = -\frac{\rho}{2}\Delta|\mathbf{u}|^2 = -\rho|\nabla\mathbf{u}|^2 < 0. \quad (3)$$

This inequality implies that the minimum of p must occur upon the boundaries of the region – neglecting the thin boundary layer, this will be upon the body surface. Fortunately for the phenomena of cavitation, it is typical that gas pockets are often trapped within small fissures or holes of a surface. It is from such trapped pockets that cavitation bubbles are typically observed to grow. If the pressure around the bubble increases, say by transport of the bubble or by natural pressure fluctuations, the bubble can violently collapse and yield a strong compression wave.

(2.) For rotational flows, pressure lows can occur at the center of extended vortices. In the simplest case of a steady, axisymmetric flow without swirl, incompressibility gives that the velocity has the form,

$$\mathbf{u} = v(r)\hat{\theta}, \text{ with } v(r) \approx Cr^2 \text{ for } r \ll 1. \quad (4)$$

The Euler equations then imply that

$$\frac{\partial p}{\partial r} = \frac{\rho v^2(r)}{r} > 0. \quad (5)$$

Thus, the pressure increases monotonically away from $r = 0$, and so has a minimum at $r = 0$. Why might this be important for cavitation? Extended vortices form dynamically in many flow situations – boundary layer detachment, shedding from sharp edges, vortex rings formed by fluid jets. If such a vortex is subjected to an axial strain, it will begin to stretch and further decrease the pressure low at its core. The formation of cavitation bubbles within such stretched vortices has been reported by Appel (1961). Appel's work is interesting in that he makes clear the relation between intense vorticity in the form of vortex tubes, and the observed cavitation.

As an aside, a very interesting question is what happens to the vorticity that had made up the core of the now cavitating vortex, and how its disposition affects the propagation of a cavitating vortex. A nice experiment would be to shoot intense vortex rings at a wall. As the vortex ring approaches the wall, it sees its image vortex ring approaching it from the other side of the wall, and it begins to stretch. The stretching could induce a spontaneous, but controlled, cavitation within the core.

(3.) Blunt bodies, or bodies with sharp trailing edges, can form permanent, attached cavitation bubbles (see the photographic plate 18 in Batchelor). In this case, the boundary of the cavitation bubble is sometimes called a *cavitation sheet*. If the sheet covers most of the body, it is sometimes referred to as *supercavitating*. If the cavity is "pressurized," that is, a source of gas is maintained by external means, then the bubble boundary can remain very smooth and laminar. If the bubble is not so maintained, then the pressure of the bubble is likely near the vapor pressure of dissolved gases, and the bubble boundary is rough and highly fluctuating, presumably due to the continuous boiling of gases out of the fluid. The steady-state cavity formed behind a body of given shape is often characterized (in its shape, length, drag, etc.) by its *cavitation number*

$$K = \frac{p_0 - p_c}{\frac{1}{2}\rho U^2}.$$

Here p_c is the pressure within the cavity, and p_0 would be the uniform pressure in the fluid in the absence of the cavity.

4. SUPERCAVITATING BODIES

The case of very high-speed torpedoes is an instance where cavitation simply cannot be avoided, and so must be designed for. It is reported in the open press that rocket torpedoes can achieve speeds of 300 knots, compared to the 60 knots of a

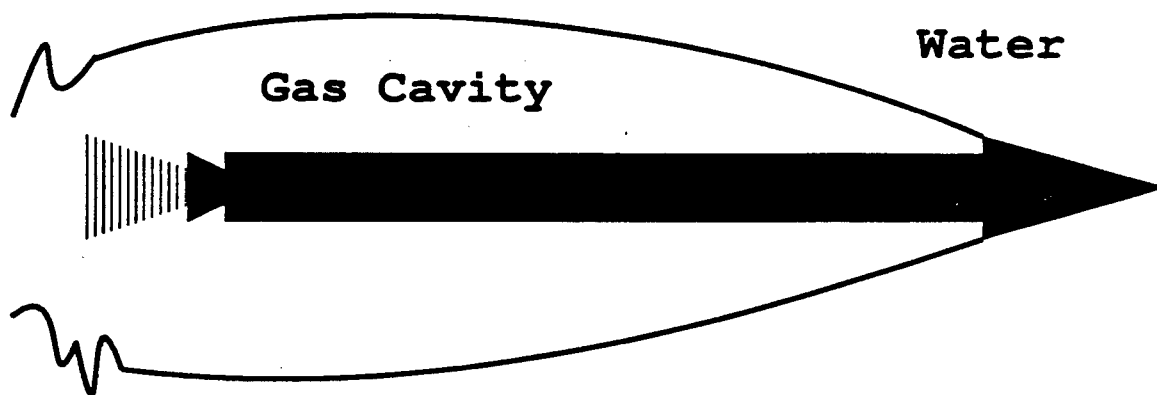


Figure 1: A Supercavitating Torpedo?

conventional torpedo. Because of their high speed, the cavitation number K is likely to be very small. And as K decreases, this leads typically to longer cavitation bubbles, perhaps enveloping the object. Such cavities eventually collapse somewhere behind the splitter, usually forming two hollow cored vortices (see Fig. 19 of Cox & Maccoli (1956)). Figure 1 is a (quite likely inaccurate) schematic of a rocket powered torpedo for which a leading splitter whose sharp trailing edge is used to create a permanent cavity in which the bulk of the body rides. I make several comments:

- The *Popular Science* article reports that the torpedos are conically shaped. I do not know whether this is accurate, but I have chosen the splitter to have a conical shape because such shapes lead to smaller drag coefficients than do other shapes behind which permanent cavities form (see Fig. 6.13.6, Batchelor, p.504). And as the cone becomes more slender, this tends to lead to smaller drag coefficients.
- It would be interesting to calculate what the thrust requirements would be to achieve these reported speeds. A useful formula characterizes the drag on a body with an attached cavity in terms of its cavitation number:

$$C_D = C_{D_0}(1 + K)$$

where C_{D_0} is the drag coefficient for $K = 0$. As I mentioned above, K is likely very small for these flows. Knowing the drag coefficient would then yield an estimate of the thrust required to maintain a constant speed. A back of the

envelope calculation, which mainly involved eye-balling an engineering table, and making several questionable assumptions about its validity, gives that maintaining 300 knots would require the (very high) thrust of $2.5 \cdot 10^5$ lbf. This is (I am told) about the thrust of a single engine on the Titan 2 lift vehicle.

- I am told that rockets work much better when thrusting against gas than against water (why is this so?). Given this, having the nozzle within the cavitation bubble may be an advantage. Also, having venting some exhaust gases into the cavity would lead to pressurization, which *might* have the effect to making the cavitation sheet a more coherent and controllable object.

4.1 IS THERE A STABILITY PROBLEM?

The dynamics of the torpedo within the bubble, as it pushes against the gas/fluid envelope, appears very similar to that of the inverted pendulum – a rod balanced upright upon a table under the force of gravity. This is of course a very unstable situation, and the rod wants to fall away from this position (see Figure 2). Given this analogy, how might stability (i.e., attitude) of the torpedo within the envelope be maintained? One possible solution (which has apparently been attempted or discussed) would be to deploy skis out from the torpedo, allowing it to “surf” on the inside of the bubble.

Are there also dynamical solutions to the stability problem? A much different approach to stabilization might be found through harmonic forcing. It is well known that the inverted pendulum can be stabilized by the rapid vertical oscillation of its support (see, for example, Kevorkian & Cole (1981)). Most presentations of this involve studying the Mathieu equation:

$$\ddot{y}(t) - (a^2 + \epsilon \cos \omega t) \sin y(t) = 0 . \quad (6)$$

However, an elegant and simpler example of such stabilization is given by Arnold (1978). He considers the model equation

$$\ddot{y}(t) = (\omega^2 \pm d^2) y(t) , \quad (7)$$

where the sign changes every r time units. It is clear that if $d = 0$, then there is a solution of the form $y(t) = ae^{\omega t}$, and the fixed point $y = 0$ is unstable. However, Arnold uses Floquet theory to show that for sufficiently rapid oscillations, the stability of the solution near $y = 0$ is restored.

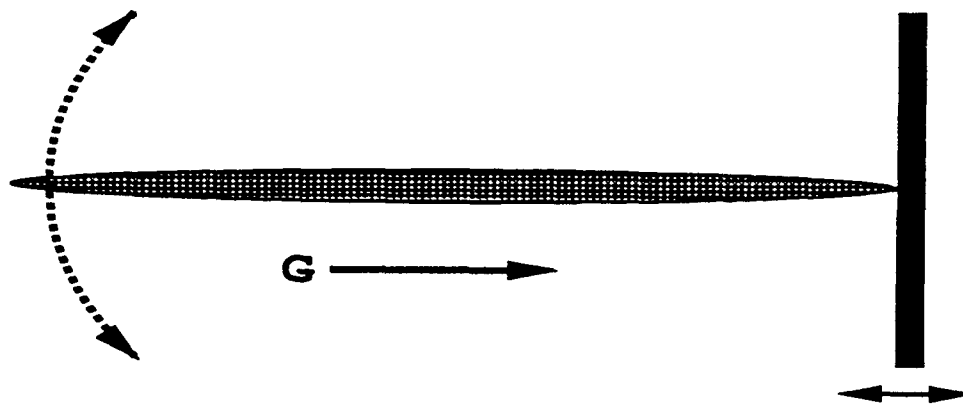


Figure 2: An Inverted Pendulum

Such an approach might be useful in helping to stabilize a torpedo riding in a cavitation bubble, say by oscillating the splitter up and down. Obviously there are many other things going on – that might have to be included in such a stabilization scheme. Considering a dynamic feedback scheme might also be useful.

5. CONCLUSIONS

There are several other topics that would be interesting to pursue:

- Guidance. Would such a torpedo have to be guided by directed thrust? Also, use of sonar for homing would be very difficult due to the noisiness of its environment.
- Could adding drag-reducing polymers be useful in reducing propulsion requirements? What are the real propulsion requirements?
- What is a defense against such a device. What is done in conventional torpedo defense? Is there something specific in their design that would make them vulnerable to some specific defense? For example, could something be added to seawater to destabilize the cavitation sheet? This brings to mind surface tension, but surface tension length-scales are likely too small to be relevant.
- A related question is the stability of the cavitation sheet. Obviously stable cavitation sheets exist. Are they convectively stabilized? That is, the flow along the sheet is so fast, that perturbations are simply swept downstream before they can grow.
- Are there other instances of intentional cavitation that might prove useful, such as a cavitating vortex ring? How would the presence of a cavity, and its

displacement of the vorticity of the ring, affect velocity, drag, and stability. Vortex rings are known to be unstable to the Widnall instability.

I would like to thank Spiro Ledeudis of the Office of Naval Research for telling me about this problem. I also thank Jerry Brackbill, Frank Harlow, and Otis Petersen for some very useful and stimulating conversations.

REFERENCES

1. D. W. Appel, *Cavitation Along Surfaces of Separation*, ASME 1961.
2. V. Arnold, *Mathematical Methods of Classical Mechanics* Springer-Verlag, New York 1974.
3. G. Batchelor, **An Introduction to Fluid Dynamics**, Cambridge University Press, Cambridge 1967/
4. C. Bender & S. Orszag, **Advanced Mathematical Methods for Scientists and Engineers**, McGraw-Hill, New York, 1978.
5. C. Brennen, **Cavitation and Bubble Dynamics**, Oxford University Press, New York, 1995.
6. R. Cooper, **Hydrodynamics Noise, Cavity Flow**, Office of Naval Research, Washington, DC, 1960.
7. R. Cox & J. Maccoli, *Recent Contributions to Basic Hydroballistics*, appeared in **Symposium on Naval Hydrodynamics**, Publication 515, NAS-NRC, Washington, DC, 1957.
8. D. Fruman & S. Aflalo, *Tip Vortex Cavitation Inhibition by Drag-Reducing Polymer Solutions*, J. Fluids Engineering **III**, 211 (1989).
9. O. Furuya (ed), **Cavitation and Multiphase Flow Forum – 1989**, ASME, 1989.
10. J. Hoyt & O. Furuya (ed), **Cavitation and Multiphase Flow Forum – 1985**, ASME, 1985.
11. H. Kato, M. Miura, H. Yamaguchi, & M. Miyanaga, *Drag Reduction by International Cavitation*, in Furuya, 1989.
12. J. Kevorkian & J. Cole, **Perturbation Methods in Applied Mathematics**, Springer-verlag, New York 1981.
13. S. Li, A. Boldy, Z. Liu, & H. Zhou, *Cavitation Resonance in a Venturi Loop*, in Furuya, 1989.
14. R. Wade & A. Acosta, *Experimental Observations on the Flow Past a Plano-Convex Hydrofoil*, in Trans. ASME, Ser. D., No. 88, 273 (1966).

**L. MATHEMATICS AND THEORY IN VIRTUAL
ENGINEERING**

**John C. Doyle
Electrical Engineering
California Institute of Technology
Pasadena, CA**

Mathematics and Theory in Virtual Engineering

John Doyle¹

California Institute of Technology²

Abstract

We will tentatively use the rubric Virtual Engineering (VE) to describe the general multidisciplinary engineering domain that involves using Virtual Reality (VR) interfaces, simulations, and integrated databases for taking complex systems from concept to design, including manufacture, operation, maintenance, and training. It also includes simulation-based approaches to decision support in policy making and real-time C³I. We will discuss the technical issues in VE and argue that conventional methods of modeling, simulation, and analysis will be inadequate in several ways in dealing with large complex systems.

We need a new discipline that focuses on developing mathematical and computational methods to integrate the currently diverse approaches for modeling and simulating heterogeneous systems with interacting fluid, structural, material, chemical, electromagnetic, and electronic subsystems. This theory of VE, or VE mathematics, must deal with heterogeneous, hierarchical, multiresolution and variable granularity models with explicit uncertainty representations. We will discuss how control and dynamical systems are the fields which have the most relevant foundations to contribute to a theory of VE, but they must expand their vision if they are to be successful.

The discussion will be kept as informal as possible, and a minimum amount of mathematics will be introduced as needed. The details of the research described here currently require substantial mathematics, particularly in functional analysis, operator theory, and computational complexity, in addition to standard machinery from control and dynamical systems theory. While it seems that significant simplifications of the theory is unlikely, it is hopefully possible to discuss the implications without all the details.

1 Introduction

Virtual reality (VR), paperless and simulation-based design, virtual prototyping, distributed interactive simulation, immersive VR, synthetic environments, and simultaneous process/product design have already gone from buzzwords to cliches. Although the enormous potential of VE is widely recognized and is being broadly pursued, a sound theoretical framework is needed to avoid compromising the potential power of VE. VE is driven by the need to improve the design and analysis of large, complex systems, the expected availability of sophisticated VR human/computer interfaces, and parallel, distributed, high-performance computers, networks, and databases.

In a VE environment, an engineer (or decision maker, analyst, etc.) with a sophisticated VR interface will be connected to networks of other engineers in various disciplines (such as materials, structures, fluids, electromagnetics, controls, electronics, computer science, manufacturing, maintenance, training, tactics) sharing a common database. Engineers will work at different levels in a multiscale, multiresolution, variable granularity model. Design changes would automatically propagate throughout the model and other engineers whose designs are affected will respond by

¹Many people contributed to the ideas in this paper. I want to particularly thank my graduate students, whose names appear prominently in the references, and Jan Willems.

²Caltech 116-81, Pasadena, CA 91125, doyle@hot.caltech.edu

evaluating the consequences of that change, including manufacturability and maintenance. Sophisticated immersive visualization methods would take data from experiments or tests on physical prototypes and facilitate the comparison of data with theory and simulation.

The entertainment industry will push VR with increasingly realistic look and feel and an emphasis on fooling human senses. While there are many challenges in further developing VR, they are already heavily funded and widely appreciated. What is important for VE is that the VR paradigm of "realistic=looks good" should not dominate VE as well. Otherwise, engineers and programmers would ultimately design systems fine-tuned for VR which do not work in reality. Without a correct and fundamental mathematical structure, VE could fail spectacularly and the potential for abusing it could be tremendous. For example, there are and will be frequent instances where competing simulations of the same system visually "prove" opposite conclusions. We must not be seduced into equating good looks with high fidelity, but must have reliable and systematic methods to evaluate the accuracy and fidelity of competing simulations. In another scenario, military planners using simulation-based C³I tools could find that actual battlefield outcomes differ radically from simulations if the planner is given inadequate indication of sensitivities to assumptions and initial conditions.

Thus we believe that in addition to refining and advancing the technology behind VE, we also need a new discipline to address the theoretical issues that will be essential to VE's success. We will tentatively call this discipline VE mathematics (VEM), although perhaps VE theory would be equally appropriate. What VEM must provide is a structure for VE that imposes a common discipline and language regarding simulation and analysis across different disciplines, streamlining model development and integration, providing common and systematic simulation and analysis tools, and helping make reliable predictions of real system performance and evaluate the impact of various assumptions. We must create as transparent a relationship as possible between assumptions and results, and a means to communicate both.

Although there will not be a single solution to the VEM problem, well-conceived research should produce software and methodologies backed up by theory. The VE engineer will always have to do case-specific, and sometimes ad hoc fixes that go beyond the theory, just like other engineers. But just as VE will improve with advances in visualization, computation, and networking, VE will also improve with advances in VEM. The goal of VEM is to push back the boundaries where the ad hoc approaches must take over, so that the ad hoc parts of an engineering solution are smaller and more manageable.

2 Modeling complex systems

We will now discuss some generic issues that arise in developing a modeling methodology aimed at modeling large complex interconnected systems. We will focus first on problems that are inherent in this setting and later we'll discuss specific approaches that we believe are most promising. By a modeling methodology we mean both the theoretical framework and the software implementation. Hopefully, the initial remarks will seem obvious and perhaps even banal to anyone who has thought about these issues.

2.1 Fidelity vs. complexity

There is usually an inherent tradeoff in modeling a physical system between fidelity and complexity, in that higher fidelity models require greater complexity. By model complexity, we mean the difficulty it takes to answer questions about the model. This is obviously not a function of just

the model, which we view as a purely declarative description of a system's behavior, but also the analysis and simulation questions that will be asked about the model. Informally, we may suppose that we have certain problems in mind, and then we can view complexity as the computational cost associated with solving these problems for a given model. Fidelity of the model is a measure of the predictability of a physical system based on the model. So if complexity measures how hard it is to say something about the model, fidelity measures how easy it is to conclude something about the true system from results about the model. Of course, a model's fidelity is also a function of the problems posed using it.

A well-known example of the tradeoff between fidelity and complexity occurs in weather prediction. For standard computer simulations used in weather prediction, the error between the prediction and the actual weather grows with time. A measure of fidelity would be how many days before the prediction departs significantly from the real weather. Complexity would be the computer time that the computation of the prediction takes. Obviously, to be useful, this computation must allow the prediction to be computed faster than real time.

We can formalize complexity in mathematical terms using computing machine models and measures of complexity in terms of time and space usage. The notion of fidelity of a model, in contrast, cannot be as readily formalized since it involves the relationship between a mathematical model and a physical system. No matter how high fidelity a model is, there is always some uncertainty about how it compares with the real system being modeled. There is no way to mathematically formalize this beyond a certain point, because the real system is not a mathematical object. Thus an engineer using a model must always make a leap of faith in evaluating its fidelity. The role of theory and a modeling methodology is to make this leap as small as possible.

While we will remain vague for the time being about what is meant by these terms, it is clear that the tradeoff between them is at the heart of understanding the modeling process. Suppose for the moment that we think of the true system as being a set of possible measured behavior and the model as predicting what those measurements can be. If we have some norm to quantify the difference between prediction and observation, then we will call the modeling error the norm of this difference maximized over some experimental conditions. We will expand on this vague description in the next section.

We will call a high fidelity model one which has small modeling error, and concentrate on the tradeoff between modeling error and complexity. If we have a good modeling methodology, then we should expect a relationship like the following:

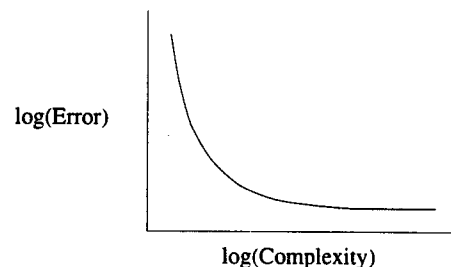


Figure 1: Tradeoff between model error and model complexity

Ideally, we want low error, low complexity models so that we can cheaply and accurately predict the behavior of our modeled system. In some sense, a good modeling methodology can be thought of as one that minimizes the error for any level of complexity. We will discuss the elements that

must go into a robust modeling methodology, focusing particularly on generic issues that arise in modeling complex systems.

We must keep in mind that there are limits to the fidelity that is possible in any modeling problem, such that eventually increases in complexity don't lead to smaller errors. For the weather prediction problem, long term predictability is viewed by experts as being impossible. More generally, any example where the physical system does not have perfect repeatability will have limits on how much modeling error can be reduced. This is reflected in the figure as the incremental benefit of higher complexity modeling going to zero before the error goes to zero.

2.2 Hierarchies

High fidelity models are often developed by "tearing and zooming", where a complex system is modeled by "tearing" it into components and then "zooming" into the component and modeling it in some detail. This can continue at several levels to produce a hierarchical model, and the full system must then be built up from the component models. This is essentially the only way that complex system models can be developed, but it leads to a number of difficult problems, including the need to connect heterogeneous component models, and the need for multiresolution or variable granularity models.

Heterogeneity arises from the presence in complex systems of component models dealing with material properties, structural dynamics, fluids, chemical reactions, electromagnetics, electronics, embedded computer systems, etc. Typically, modeling methods in one domain may be incompatible with those in another. Even when we can break the system into multiple levels with the lowest level containing only one modeling domain, the component models must be combined. Thus a critical issue that must be dealt with in every aspect of modeling is how to combine heterogeneous component models.

2.3 Multiresolution and variable granularity models

There is also a need for multiresolution or variable granularity models, because the process of tearing and zooming can continue indefinitely to create an infinitely complex model. Unfortunately, it often cannot be determined a priori what modeling error, and thus complexity, is appropriate for a given component, because the required fidelity will depend on the full system model and the problems that will be solved with it. Thus the error/complexity tradeoff is made even more difficult in a hierarchical model.

The standard approach to developing variable error/complexity component models is to allow multiresolution or variable granularity models. Simple examples of this include using adaptive meshes in finite element approximations of continuum phenomena, or multiresolution wavelet representations for geometric objects in computer graphics. In hierarchical models, the problem of developing analogous variable resolution component models is quite subtle and will surely be an important research area for some time. Using these same examples, there are difficult problems associated with modeling continuum phenomena involving fluid and flexible structure interaction, as well as phase changes. Similarly, building aggregate multiresolution models of interconnected geometric objects from multiresolution component models is a current topic of research in computer graphics.

2.4 Uncertainty

There is always uncertainty in modeling any physical system, but there are several strategies for improving the predictability of models in the presence of this uncertainty. One method that has proven particularly useful in control engineering is to explicitly model uncertainty using noise, parametric uncertainty, and bounds on unmodeled dynamics. This allows for analysis and simulations to be done using a set of uncertain models rather than a single idealized model, yielding a much more systematic study of model predictability. Modeling uncertainty will be the primary focus of the remainder of this paper and be discussed in more detail.

2.5 Computational complexity

While many of the results we will describe below rely heavily on the theory of computational complexity, it is beyond the scope of this paper to give a tutorial account of these techniques. Readers unfamiliar with machine models and the classes P, NP, and NP hard should still be able to follow the main points. We will consider a problem computationally easy if it can be solved on a standard computer in a time that grows polynomially in the size of the problem. Such problems are said to be in class P, or simply in P. Problems not in P are considered to be hard. NP is a class of problems that includes P, and the central problem in theoretical computer science is the conjecture that $P \neq NP$, which experts believe. The hardest problems in NP are called NP hard, and if $P \neq NP$, then these problems are not in P, and therefore really are hard.

We will assume that $P \neq NP$ for the purposes of this paper, so there are basically 2 types of complexity results that are relevant: 1) a problem is proven to be in P by exhibiting an algorithm that solves it in polynomial time, 2) a problem is proven to be NP hard by showing it is equally hard as a known NP hard problem. We will use the reminder ($P \neq NP$) for results that involve proofs of NP hardness, because the "hardness" of such problems relies on this conjecture.

A final important feature of the computational complexity proofs we will refer to is that NP hardness implies that some subset of the worst instances of any problem are hard. It is possible that an algorithm can be developed for an NP hard problem that runs in polynomial time on all problems that ever arise from engineering applications, because the "hard" subset on which it takes exponential time don't ever arise. Unfortunately, it can be extremely difficult to prove such an assertion, perhaps harder than proving $P \neq NP$, so we must rely on experimental evidence.

These remarks may make it sound like NP hardness is a useless abstraction, but our experience suggests otherwise. NP hardness does not preclude the existence of practical, polynomial time algorithms, but it does preclude the practicality of certain styles of algorithms. For example, any algorithm that does the same amount of computation on each instance of a problem must be extremely slow. Roughly speaking, if we can say a priori how much time an algorithm will run on an NP hard problem, then that time must grow exponentially ($P \neq NP$). This is an example of how the theory can direct research in practical algorithm development.

3 The modeling process

In this section we will discuss in somewhat more detail the modeling process, including comparisons with experiments and the role of uncertainty. The differences between our models and the systems they represent is a much more critical problem in large, complex interconnected systems than in conventional settings. In order to understand this issue and the technical approaches to dealing with it, we'll start by considering the simplest possible experimental setting that illustrates the essentials.

3.1 An abstract experiment

Suppose we have an experiment with the following assumptions:

- 1) Specified vector of measurements with time trajectory $y(t)$. $t \in [0, T]$
- 2) Repeat the experiment to get a set Y_m of measured trajectories.
- 3) Normalized so $y(0) = 0, \forall y \in Y_m$.
- 4) The set of all possible trajectories is Y , (so $Y_m \subset Y$).

The assumption that $y(0) = 0$ is simply a choice of units. We are assuming that as the experiment is repeated, the initial conditions are identical ($y(0) = 0$) as far as the measurements y are concerned. Alternatively, we could imagine that we build multiple copies of the same system and then run experiments whose initial conditions are indistinguishable from our measurements. The set Y of all possible trajectories is what we think we would get if we did all possible experiments. This is unknowable in principle and requires a leap of faith on the part of the modeler, but it is a useful abstraction. In a real experiment we would also have finite precision, sampled data which we would store on a computer. Since the finite and sampled nature will not be an issue initially, we will worry about it later and just assume that Y is a subset of an abstract function space (such as $L_2[0, T]$) with a norm $|\cdot|$.

Of course, we are interested in varying initial conditions, too, but we want to focus on the fact that even if we make all the measurable initial conditions identical, the trajectories are not identical. That is, Y is not a single trajectory and the set may be large, because even if we make our experiments have identical measured initial conditions, there will be other differences that we cannot measure. Even if we view our physical experiment as a purely deterministic system, we can never measure the whole state of the system, and these variations can lead to growing deviations between trajectories.

$E_{max} = \frac{1}{2} \max\{|y_1 - y_2|, y_i \in Y\}$ is the radius of the set Y in the norm $|\cdot|$ and may be thought of as a measure of fundamental unpredictability. If we have a good experiment, the set Y will have some structure that we will want to capture in our model. That structure may not be at all obvious, of course.

3.2 Conventional modeling

As a starting point, suppose we want a model that produces a "measurement" $x(t)$ from $x(0) = 0$ such that $x(t)$ is close to Y . We can think of our modeling error as either $M_{max} = \max\{|y - x|, y \in Y\}$ or $M_{min} = \min\{|y - x|, y \in Y\}$. Note that no matter what our model x is, $M_{max} > E_{max}$. We want a model that minimizes M_{max} and M_{min} (ideally, $M_{max} = E_{max}$ and $M_{min} = 0$), but also minimizes complexity. Basically, we view "understand" to be having a simple model with small error.

Of course, there is a tradeoff between complexity and error, and complexity goes up as error goes down. Fortunately, for most modeling techniques using differential or integral equations or variational principles, complexity of simulating the model to produce $x(t)$ typically grows in polynomial time with system dimension. The rough intuition is that higher fidelity modeling usually involves using finer temporal and spatial resolution, or expanding the spatial domain. Thus complexity of generating $x(t)$ typically grows as some polynomial in n^m where n is the grid size ($1/n$ is the resolution) and may be very large, and $m \leq 4$ is the space-time dimension. This may lead to large n^m but the growth in adding additional grid points increases n , not m , and

growth remains a polynomial. So adding a bit of complexity to the model won't create a sudden overwhelming increase in the complexity of simulation.

This is essentially the modeling paradigm that dominates science and engineering, and dynamical systems has recently proven useful in understanding this paradigm. The most profound discovery is that behavior that appears to be very complex or even random can actually be the chaotic trajectory of a low dimensional nonlinear system (this is close to an informal definition of chaos). Chaotic systems may have some short term predictability, but they have severely limited long term predictability.

The main application of this modeling paradigm is that if we can get a good model (small $M_{\min}$), we can study that model and conclude that if it has some problem (sensitivity, instability, chaos, etc.) we can expect the true system to have potential for the same problem. For example, current models of weather have reasonably good short term predictability (days) but very poor long term predictability (weeks or more), because these models exhibit extreme sensitivity to initial conditions. In a sense, the model does not predict itself very well. If we run a simulation with certain initial atmospheric measurements and then change one measurement by an amount well below the resolution of the measurement devices and rerun the simulation, the resulting trajectory departs significantly in a week or so (of simulated time). This says only that the model is not predictable, and does not directly imply that weather is not as well. However, it is reasonable to conclude that without some truly profound discovery about the structure of atmospheric dynamics, weather will never be predictable on a long time scale. This is the type of conclusion that scientists and engineers routinely make from models.

Put simply, this modeling paradigm answers some question about an experiment (e.g. predicting the weather, or the performance of a proposed design for an airplane or a microprocessor) by giving a necessary condition in terms of a model. That is, a necessary condition for something to be true about the real system is that it be true about a high fidelity model. This is usually most convincing when 1) the result is negative and 2) we can check with experiments. Again, the weather is a good example of this, where our experience with models confirms what is suspected from doing real weather prediction. We are generally less convinced when the results are positive ("this new design is great") and there is limited experimental experience ("so let's build it"). Unfortunately, one of the main motivations of VE is to avoid the expense of physical prototypes and produce complex systems that work well in simulation.

It is perhaps only a small oversimplification to say that this standard modeling paradigm is poorly suited to VE. In VE we want to build complex systems, but not "complex behavior" in the sense of dynamical systems. We want reliable, predictable, even boring, system behavior, because the environments that our systems must operate in are not reliable and predictable. If our model is a component in a larger hierarchical model, and we want high levels of predictability for our hierarchical model, this paradigm will not work. Small errors at the component level can lead to large errors at the system level, and this modeling approach offers little help other than necessary conditions. In other words, this modeling approach is not robust.

The typical "consumer" of a conventional model is a human expert in the problem domain being modeled. What constitutes a good model and how experiments and models compare is part of the common educational background in the problem domain. Models serve as a means to communicate our understanding of physical systems. Model uncertainty and the interpretation of assumptions is typically implicit and part of the domain-specific expertise. In contrast, in VE the consumer of a model will be a larger model in a hierarchy. We cannot rely on anything that is domain-specific, implicit, and requires human interpretation at every level.

In summary, the problem with this standard modeling paradigm is that it 1) lacks robustness, 2) relies on human interpretation, 3) may generate overly domain-specific models, 4) has only implicit

uncertainty descriptions. Next we outline an alternative approach which has been successful in robust control and has great promise for VE modeling.

3.3 Robust modeling

The key idea in robust modeling is to use model *sets* X that cover the true set Y ($Y \subset X$). One measure of model error is then something like

$$R_{error} = \max_y \min_x \{|y - x| \mid y \in Y, x \in X\}$$

and $R_{error} = 0$ if and only if $Y \subset X$. In this type of modeling, the real goal is to cover Y with as little extra behavior as possible, that is the content of $X - Y$ should be minimized, and the appropriate measure of error for this is something like

$$R_{extra} = \max_x \min_y \{|y - x| \mid y \in Y, x \in X\}$$

and the best model is one that minimizes R_{extra} and is low complexity.

In contrast to the necessary conditions of conventional modeling, this provides sufficient conditions for performance: If it works on the uncertain model, it will work on the true system. This is essential for robust simulation of hierarchical interconnected systems. The problem with this approach is analyzing/simulating the uncertain system. Complexity growth can be exponential in the dimension of the uncertainty description, depending on the uncertainty mechanism used. This has been studied extensively in robust control and recently there has been substantial progress in characterizing the complexity tradeoffs. This will be discussed later.

Note there is a leap of faith in assuming that $Y \subset X$. Often however, the leap of faith in assuming $Y \subset X$ for robust modeling is less than that in assuming $x \in Y$ in conventional modeling. The advantage is that once this is done at the component level, the consequences then can be evaluated at the system level.

Traditionally, when uncertainty was introduced explicitly, the theory relied heavily on stochastic models. In contrast, robust control has focused on deterministic set descriptions. These have recently been reconciled with a useful set-based description of white noise that can be used in the robust control paradigm. Generalizations of robust control methods have also led to new techniques for model identification and validation from data. This will be discussed further in later sections.

3.4 Necessary and sufficient conditions

If we have a single model producing x and a set model X such that $x \in Y \subset X$, then x provides a necessary condition for analysis and X provides a sufficient condition. These can be combined and refined to yield methods that converge and are thus necessary and sufficient. This basic idea is behind the most promising existing analysis and robust simulation tools for both linear and nonlinear systems. The challenge again is the computational complexity associated with this procedure, since most problems of interest can be shown to be NP hard. We have great success in developing algorithms that overcome this theoretical intractability on problems of engineering interest, and expect that generalizations of these algorithms will be essential parts of our VE robust modeling and simulation software infrastructure.

4 Uncertainty models

In this section we will expand on our discussion of robust modeling by considering explicit mechanisms that have proven useful in modeling uncertainty in physical systems. We'll use aircraft models as an example, but these types of uncertainties are quite generic and will naturally appear in any robust modeling effort.

In the previous sections we talked abstractly about experiments with measurement sets Y_m and all possible trajectories Y , and models x and model sets X that cover Y in the sense that $Y \subset X$. We will now think in terms of models that are differential equations so that x and X would be generated by numerically integrating the differential equations. It would be possible to think about introducing uncertainty just in terms of the abstract behavior x and X , but moving immediately to the more concrete approach taken here is more familiar and natural to most engineers and scientists, if perhaps less theoretically elegant.

4.1 Noise

Consider the standard wind gust noise models that are used to describe the forces generated on an airplane due to unsteady, time-varying atmospheric winds. Typically, the atmosphere is not modeled in detail, but the forces are described as signals which are constrained in some way, either in terms of some probability distribution, or via statistical tests. Traditionally, noise is modeled as a stochastic process even when it is more appropriate to think of it as chaotic. That is easy to see in the case of wind gusts which are generated by fluid motions in the atmosphere. Just because the wind appears random and we don't want to model the atmosphere in detail doesn't mean that wind is naturally viewed as a random process. Modeling wind as a stochastic process must be motivated by convenience.

Other examples of phenomena that is often modeled as noise include coin-flipping, arrival times for queues, thermal noise in circuits, etc. In each case, the mechanism generating the noise is typically not modeled in detail, but some aggregate probabilistic or statistical model is used. Robust control has emphasized set descriptions for noise, in terms of statistics on the signals such as energy, autocorrelations, and sample spectra, without assuming an underlying probability distribution. Recently the stochastic and robust viewpoints have been reconciled somewhat by the development of set descriptions that recover much of the characteristics and convenience of stochastic models.

4.2 Parametric uncertainty

Standard rigid body airplane models typically consist of a generic form of the model as an ODE with vehicle-specific parameters for mass distributions, atmospheric conditions (dynamic pressure), and aerodynamic coefficients. While these parameters vary considerably between vehicles, there is also uncertainty in their values for a specific vehicle, due to manufacturing tolerances, measurement limitations, or because a constant is used in place of more complicated phenomena that aren't modeled in detail.

Physical mechanisms that might lead to parametric uncertainty include "constants" that change with time (e.g., dynamic pressure, center of mass), but whose variations are very slow compared with the other dynamics of the system and the detailed mechanisms responsible for this variation are not modeled. Physical constants that change as a system ages would be an example of this. As another example, surface deflections don't directly generate forces, but rather act through complicated changes in the fluid. We may crudely capture this with a quasi-steady model

that assumes surface deflections do directly generate forces, but with parametric uncertainty in the aerodynamic coefficient that quantifies the magnitude of the force generated.

Examples in other domains where parametric uncertainty would arise would be values of resistance, capacitance, and inductance in lumped circuit models, mass, spring, and damping constants in lumped mechanical models, thermal resistance and capacity, and so on. Of course, parametric uncertainty doesn't only arise in lumped models but also occurs, for example, in coefficients in PDE models as well.

4.3 Unmodeled dynamics

The use of constant parameters to model aerodynamic forces generated by, say, surface deflections neglects the complex, unsteady fluid flows that more accurately describe the physics. Such quasi-steady approximations may be accurate for small, slowly changing surface deflections but become increasingly inaccurate when the movements are large and fast. We may choose to model this type of uncertainty with a mechanism similar to noise above, but involving relationships between variables, such as surface deflections and forces. Like noise modeling, we use bounds and constraints, but now on the ratio of signals.

Similarly, rigid body models assume that forces directly generate rigid body motion, while models that included flexible effects allow for bending as well. We might choose not to model the flexibility in detail but bound the error between forces and rigid body motions as constraints on signals.

4.4 Mixed uncertainty descriptions

As an example, suppose we model the force F generated on the wing of an aircraft by a deflection d_{ail} of the aileron as the static model $F = c_0 d_{ail}$, where c_0 is some constant. We might add some parametric uncertainty δ_1 and some unmodeled dynamics δ_2 to give $F = (c_0 + c_1 \delta_1 + c_2 \delta_2) d_{ail}$ where c_1 and c_2 are constants. The simplest view of this would be with δ_1 an unknown real number bounded by $|\delta_1| \leq 1$ and δ_2 is an unknown operator bounded in some induced norm as $\|\delta_2\| \leq 1$. With this normalization, the constants c_1 and c_2 give the size of the allowable perturbations. The perturbation δ_1 would model our presumed uncertainty due to unmodelled mechanisms such as variations in dynamic pressure or uncertainty in the effective area of the aileron when deflected. If only δ_1 were present then the model would be uncertain but have no dynamics.

A more refined model might include more details about the way the forces are generated as a consequence of changes in the local flow which are in turn generated by the surface deflection. If we choose to ignore such details, it would not be possible to account for all such uncertainty with just δ_1 since the fluid mechanisms have memory and dynamics. We can, however, cover the fluid mechanisms with the uncertain operator δ_2 . The resulting set of systems includes dynamics, but without a detailed description. Of course, this set includes behavior that could not possibly occur in any physical system, since the set of operators with $\|\delta_2\| \leq 1$ includes mathematical objects with no physical meaning. We could add additional constraints to try to capture more physics, but at the expense of a more complex model.

Additional uncertainty could be added in the form of noise to reflect the fact that forces are generated on the aileron due to wind gusts which in turn generate forces on the wing. And so on. One inadequacy with this example is that it is artificial to think of the forces as being a function of the deflections, rather than thinking of the aileron/wing system as imposing certain physical constraints on the relationship between forces and deflections. This more "correct" modeling could

include uncertainty mechanisms exactly as was done here, so it doesn't illustrate anything different and would be more complicated to explain.

4.5 Multiscale/multiresolution/variable granularity

The tradeoff between model fidelity and complexity is reflected in uncertainty modeling as well. As we add more detail to our models in an attempt to capture more phenomena and get higher fidelity models, we add finer structure to our uncertainty descriptions. Each mechanism that is modeled introduces the potential for new parametric uncertainty and new sources of noise and unmodeled dynamics.

As we add more detail to a model, for example, adding flexibility to our rigid body model, or some details of the flow around the vehicle, or the interaction of flexible and fluid effects (aeroelasticity), more constants would be introduced, which would in turn naturally lead to more uncertain parameters. At the same time, the size of the unmodeled dynamics may be reduced as more phenomena is modeled in detail. This potentially reduces the total uncertainty at the expense of greater complexity, but doesn't ever eliminate uncertainty due to unmodeled dynamics and noise. There are always phenomena, typically at small and large spatial and temporal scales, that are not modeled in detail, or cannot be modeled accurately. A complete uncertainty description must account for these effects.

We may also move the other direction, looking for aggregate models that cover details with unmodeled dynamics. Thus for a given component we may have a nested family of models that vary from simple descriptions with a crude and large uncertainty description involving a few parameters and large unmodeled dynamics to a more refined description with more parameters and smaller unmodeled dynamics.

5 Analysis and simulation of uncertain models

The engineer's main tools for the analysis of uncertain nonlinear systems are still simulation coupled with robustness analysis of linearizations at equilibria. Exploring regions of uncertain parameter values for nonlinear systems by repeated simulation involves prohibitive computation growth rates. It is quite easy to see intuitively why this might be so, even for linear systems with parametric uncertainty. Suppose that for fixed parameters the cost of running a simulation is C , and that we have p uncertain parameters, and we want to evaluate r values for each parameter. The total number of possible combinations of parameter values is r^p and the cost of all the evaluations is Cr^p . Thus the growth rate in p is exponential, which means that the addition of even a few parameters to a model can cause severe increases in computation. This intuition is supported by theoretical results that show that even for linear models with parametric uncertainty, evaluating virtually any performance measure for a simulation is NP hard in the number of parameters.

Given the reliance on such simulation in VE, this is a severe technological problem and one of the major challenges to a robust theory of VE. There are several approaches to overcoming the apparent intractability. An indirect approach which has been the industry standard for decades is so-called Monte Carlo simulation. The essence of this idea is that if we assume some probability distribution on the parameters, then any performance measure has an induced probability distribution. To experimentally estimate the performance distribution we can select random values of the parameters and run simulations to get a sample distribution of performances. This can then be subjected to standard statistical tests like any experimental data.

5.1 Monte Carlo methods

The beauty of the Monte Carlo approach is that the accuracy of the estimates trivially does not depend on the dimension of the parameter space, so there is no growth whatsoever in the computation cost with the number of parameters. The only cost is that of the simulation itself, and the number of times it must be repeated to get a statistically significant sample size. This approach is particularly useful when the probability distributions have natural interpretations. Examples of this would be estimating the yield of chips as the result of some manufacturing process, or estimating the mean time between failures for components of some system.

There are several weaknesses to the Monte Carlo approach. One is that it often may be difficult to interpret the probabilities that describe both the assumptions and the results. We may simply want to know if anything bad can happen for some set of parameters, and there is no natural way to interpret the probability distribution of the parameters or the resulting probability distribution of the performance. Also, we can't actually compute the probability distribution of the performance, but only indirectly assess it. That is, we don't get results like "the plane will not crash 99% of the time," but instead like "we can be 95% confident that the plane will not crash 99% of the time." The need to use confidence levels can be particularly annoying when they are not naturally motivated by the model.

Fortunately, there are alternatives to the Monte Carlo technique that provide different and complementary answers, and the development of such alternatives has been a driving force behind the research in robust control theory for the last 20 years. The difficulty is that if we want to avoid purely probabilistic estimates we must overcome the worst-case intractability implied by the NP hardness of our problems. The approach that has proven to be most successful involves computing hard bounds, as opposed to the soft bounds available by Monte Carlo, and refining the bounds by using branch and bound. The meaning of soft bounds, hard bounds, and branching will be discussed next.

5.2 Hard bounds

To discuss the computation of hard bounds, we need to be a bit more specific about our performance objective. Let's suppose that we have some desired behavior for our system, and the deviation by our model from that ideal behavior is our error e , which can then be viewed as just another measurement. Suppose our units are chosen so that the performance objective is $E < 1$ for $E = |e|$ in some norm, and $E_{max} = \max |e|$, where the max is taken over all uncertainties in our set. We will assume that we have three kinds of norm-bounded uncertainty: noise, parameters, and dynamics, and an allowable uncertainty is one that satisfies our norm constraint. We will further assume that our model is linear in that for any fixed values of the parameters, the system is linear, although the parameter dependence may be nonlinear. Our analysis question can then be taken as:

Question: Q : Do there exist allowable values for noise, parameters, and unmodeled dynamics that make $E \geq 1$?

We are interested in how hard it is to answer this question, for various assumptions on the model and the uncertainty. It turns out that noise and unmodeled dynamics are equivalent, in that they both can be described entirely in terms of norm inequalities. If we have a perturbation δ_2 and signals $y = \delta_2 u$, then if δ_2 is an otherwise unrestricted operator, $\|\delta_2\| \leq 1$ is equivalent to the condition $|y| \leq |u|$. So we can represent unmodeled dynamics either in terms of a norm bounded operator or as norm inequalities on signals. Further, the existence of a noise $|n| \leq 1$ such that $|e| \geq 1$ is equivalent to existence of $|e| \geq |n|$. Further, the existence of $|e| \geq |n|$, $y = \delta_2 u$,

$\|\delta_2\| \leq 1$, is equivalent to existence of $|n| \leq |e|$ and $|y| \leq |u|$. Thus noise and performance and unmodeled dynamics are all mathematically equivalent, provided they are described in terms of the same norm and $\|\delta_2\|$ is in terms of the induced operator norm for $|\cdot|$.

Some of the recent major results in robust control give a fairly complete classification of which problems are in P and which are NP hard. The most significant P result is roughly that if the uncertainty and performance can be described entirely with norm inequalities as discussed above, then $\mathcal{Q}$ can be solved by a finite dimensional convex optimization problem which in turn can be solved in polynomial time in the dimension of the uncertain signals and system state. Thus for this type of uncertainty, $\mathcal{Q}$ is in P. It is not easy to describe the intuition behind this result without developing substantially more mathematical machinery.

Unfortunately, parameters cannot be described as norm inequalities on signals, and it can be shown that $\mathcal{Q}$ is NP hard in the number of real parameters. As noted before, this does not preclude our developing practical polynomial time algorithms, just that they will be difficult to analyze and will probably run in exponential time on some subset of problem instances. We then try to find algorithms for which this subset is exceedingly rare. Of course, we can only verify this experimentally, but extensive numerical experience suggests that we can be successful.

We will now describe algorithms that have been successful in overcoming the apparent intractability of NP hard problems involving real parameters. Suppose that we denote our set $\mathcal{D}$ of uncertain perturbations Δ with $\Delta \in \mathcal{D}$, which includes real parameters and norm bounded operators, and our performance as $\mathcal{Q}$: Is $\max_{\Delta \in \mathcal{D}} E(\Delta) \geq 1$? It is difficult to compute $\max E(\Delta)$ if Δ includes many real parameters, but we can still get bounds. Note that we can trivially get a lower bound, because for any $\Delta_0 \in \mathcal{D}$, $E(\Delta_0) \leq \max E(\Delta)$. Local search over Δ can give a better lower bound $lb < E_{max}$, and there are various algorithms generalizing power iterations for eigenvalue computation that are even better. All of these algorithms have polynomial growth rates, so a lower bound can easily be computed and with the best algorithms this lower bound is usually close to the global maximum. Of course, if we could find the global maximum, then we'd have $\max E(\Delta)$, but this problem typically doesn't have any nice properties that make it possible to find the global maximum in polynomial time ($P \neq NP$).

We can get an upper bound by covering the set $\Delta \in \mathcal{D}$ by $\bar{\Delta} \in \bar{\mathcal{D}}$ where $\mathcal{D} \subset \bar{\mathcal{D}}$. Then if the uncertainty $\bar{\Delta} \in \bar{\mathcal{D}}$ is described entirely in terms of norm inequalities, we can easily compute $\max E(\bar{\Delta})$ and get an upper bound. Of course, there are many ways to do this covering, some giving better bounds than others. Fortunately, there are nice ways to parametrize all coverings in such a way that the covering can be optimized simultaneously with the bound while maintaining convexity and polynomial growth.

Thus we have ways to easily compute a lower bound lb and upper bound ub . These bounds can not be guaranteed to be close, because even the approximation problem is NP hard, and there are examples where the upper bounds can be arbitrarily bad. Fortunately, the bounds are typically quite good, usually $lb/ub > .8$ for typical ensembles of problems, and this ratio seems not to get worse with problem size. A ratio of .8 is adequate for most engineering problems because it provides a perturbation that is within 20% of the worst possible. In other words, the true answer is within about 10% of $.9ub$.

Thus computing the bounds gives an answer good enough for engineering purposes most of the time, but we are still left with the problem of what to do when the bounds are poor, when $lb/ub < .8$. One successful approach has been to use branch and bound to refine the bounds, and experience suggests that careful use of branch and bound can effectively eliminate the bad cases. It is less successful at making the bounds much better than the average without excessive computational cost. Next, we'll describe how branch and bound works and try to explain the intuition behind these observations.

5.3 Branch and bound

If the ratio lb/ub is too small, then we can refine the bounds by splitting into a disjoint union $\mathcal{D} = \mathcal{D}_1 \cup \mathcal{D}_2$, referred to as “branching.” The simplest way to do this is to split one of the real parameters into 2 parts. We may now compute new bounds lb_i and ub_i for $\max(E(\Delta_i))$, and the new global bounds are then $lb = \max(lb_i)$ and $ub = \max(ub_i)$. Branching naturally makes the bounds tighter. The upper bound is lower since the set being covered is smaller and the bounds can be tighter. The lower bound is similarly larger because there are now 2 searches over smaller regions. Even if this doesn’t yield a new worse perturbation, the previous worst known perturbation can still be used and the lower bound won’t change.

We can continue to subdivide the $\mathcal{D}_i$ further and compute bounds on the resulting tree of subdomains. It is standard in branch and bound schemes to prune the tree by noticing that if $ub_i \leq lb$ then $\mathcal{D}_i$ can be eliminated, because the worst perturbation can’t possibly be in $\mathcal{D}_i$. It is almost trivial to prove that under very mild conditions, this scheme converges, and convergence doesn’t depend on the quality of the bounds. Intuitively, the continued subdivision of $\mathcal{D}$ eventually produces a very large number of systems, each with small uncertainty, so eventually the uncertainty almost doesn’t matter, and the bounds converge. This also makes clear why this convergence is useless, since we are simply gridding the space of real parameters, which is what we wanted to avoid in the first place.

For branch and bound to be effective, it is essential that the bounds get reasonably close without extensive subdivision, because just splitting each real parameter once would still yield 2^p subdomains. Thus pruning must eliminate most branches, essentially preventing the tree from getting too broad. As we said earlier, our bounds are not always tight ($P \neq NP$), but our best bounds are usually tight. When we get a situation where $lb \ll ub$ the branching must produce prunable branches. So branching must take us from domains which have poor bounds to domains that have good bounds. Of course, this can not always happen ($P \neq NP$), but the question is whether it happens on all but a very small set of problems. Proving this appears hopeless, so we must rely on experimental evidence to guide us. The evidence so far, gathered on many thousands of examples chosen to be representative suggests that branch and bound can be used to get the worst-case bounds ratio as good as the average bound ratio with modest additional computational cost.

For this scheme to work, it is essential that the bounds have a good average lb/ub . Extensive experimentation with several alternative bounds and with attempting to get lb/ub much better than the bound average has shown that in all cases that have been tried, the average lb/ub has been obtained on all problems in polynomial time, and improving much beyond that has produced exponential time. While it is possible to construct examples which violate this, they seem to be very special in that they do not arise either when problems are selected pseudorandomly or when selected with engineering motivation.

5.4 Soft vs. hard bound summary

We can distinguish the results from Monte Carlo, which we will call soft bounds, from the hard bounds approach above. By soft bounds, we mean that we compute E for a large number of random Δ , the worst case in that set is always a lower bound, and an upper bound can be estimated with some confidence level. The number of random Δ required to achieve a specified confidence level does not depend on any problem dimension, so the computation growth with problem size depends only on the cost of computing E , which usually has polynomial growth. Recall that the upper bound is not guaranteed, and it is possible that it is very misleading. These

soft bounds are particularly useful when the probabilistic interpretation is natural, and are less so when the worst-case estimates are desired.

In contrast, the hard bounds are guaranteed and can be refined by branch and bound. The potential difficulty is that there exist problems for which this refinement may take prohibitively long. Fortunately, it seems that such problems are extremely rare to the point that it is unlikely that anyone would ever encounter one without specifically constructing it. It is interesting to note that this is true of course with essentially all numerical algorithms, even those we think of as having polynomial time algorithms, such as eigenvalue and singular value computation. For these problems, however, there is a much clearer picture of the nature of "hard" problems, whereas for the problems discussed here, the evidence is entirely numerical.

One of the important differences that we have discovered in initial numerical comparisons of soft vs. hard bounds is that if one seeks anything like comparable information, the hard bounds approach is much faster. Typically, only a few branches are needed to refine the bounds to a reasonable level, while many thousands of points must be evaluated for soft bounds to have any reliability. Since the cost of each branch is essentially the same as the cost of each Monte Carlo evaluation, the difference in the number of evaluations is the critical factor in favor of the hard bounds.

There remain major advantages to the soft bounds. Obviously, when a probabilistic interpretation is natural, then soft bounds are appropriate. But another consideration is that Monte Carlo methods apply immediately to any problem for which we can do simulations efficiently, whereas good bounds are available for only a special class of problems. Fortunately, we don't have to choose, since we will certainly want to use both methods and they give complementary information. At the same time, the success of the hard bounds approach so far strongly motivates us to expand the class of problems that we can compute good bounds for.

6 Extensions and further reading

We have focused primarily on one area of VE mathematics having to do with robust modeling, but the issues discussed more superficially in the introduction regarding software architectures and all the problems of modeling large hierarchies of heterogeneous components, which represent challenges even for conventional modeling, are also important. Of course, these issues are all ultimately inseparable, but the starting points for research will come from current disciplines such as computer graphics, numerical methods, computational mechanics, applied math, etc.

We have discussed research growing out of control and dynamical systems, as they have the right foundation to treat the various issues of predictability of complex models. But it would be very misleading to suggest that these fields are currently solving problems directly relevant to VE, except in special cases. To be relevant to VE, both areas must greatly expand their view of their problem domains. Control theorist must get beyond the plant/controller paradigm and consider more general interconnections of components, must view modeling, model validation with data, and analysis of models in a more unified manner, and more effectively apply robustness analysis methods to nonlinear models.

It appears that mainstream nonlinear control theory is unfortunately for the most part irrelevant to VEM, and a more promising direction appears to be the integration of robust control with dynamical systems, whose perspective on nonlinear systems and limits to predictability is clearly relevant to VEM. On the other hand, mainstream dynamical systems must also expand its vision to contribute to VEM, going beyond chaos and the celebration of "complexity" to the more difficult questions of how to design and verify systems that behave predictably. For example,

the use of normal forms and symmetry has been very effective in both solid and fluid mechanical systems. It has also provided a framework in which computational tools, such as dstool, have been developed.

In the remainder of this section, we will outline some research areas in control theory that appear to have relevance to VEM, and suggest some additional reading. The viewpoint taken here is extremely parochial, covering only research at Caltech, and is obviously not intended in any sense to be an overview of control theory. It is primarily for experts who are interested in finding further details.

6.1 Integrated modeling, identification, analysis, and design.

The use of mathematical models to design large systems involves various instances and possibly iterations of modeling, identification, system design and controller synthesis, and various forms of analysis and simulation. These activities are currently performed using a variety of mathematical machinery, but we have recently developed an extremely promising unified framework for these various aspects of system design [2, 8, 14]. In addition to providing an "interface" between system ID and control, this framework allows advances in different directions to combine readily. For example, it is clear how to combine progress in linear system ID with progress in nonlinear robustness analysis to produce nonlinear system ID methods, a major need in VE. Our framework also clarifies many of the computational issues and what basic algorithms must be developed.

6.2 Robustness analysis computation.

As discussed earlier, robustness analysis with real parametric uncertainty is NP hard, generally viewed as implying worst-case intractability. Except for special cases, the more general methods of identification and implicit analysis are also NP hard. Conventional numerical analysis notions of guaranteed algorithm convergence are irrelevant in these problems, because global convergence is computationally prohibitive and local convergence is of little value. Thus the only reasonable strategy is to aim for algorithms which exhibit experimentally good performance on problems of engineering interest, and here our success in robustness analysis is very encouraging. A unified framework allows for experience gained in algorithms for one particular problem to be transferred to the more general class.

Thus, to obtain acceptable computation, we do not attempt to solve the various hard problems exactly but rather to obtain good bounds, and aim for acceptable growth rates on problems of engineering interest, rather than for all problems that are mathematically possible. Upper bounds are usually convex feasibility problems, which have tractable computation. We have demonstrated that branch-and-bound can be successfully used to overcome the intractability of mixed μ , and developed power algorithms for the lower bounds which are much faster and produce better bounds than conventional local optimization [1, 9, 21, 22, 23, 8].

6.3 Analysis of implicit systems.

We have extended robustness analysis techniques to systems described in implicit form [16, 15], developing new tools for analysis of systems under a combination of time-invariant/time-varying perturbations [12] and exact conditions for robust H_2 performance analysis [13]. There is strong engineering motivation for this extension, particularly for VE. In fact, the standard control theory I/O formulation is only adequate for systems which are deliberately built to match the "signal flow" conception, and it appears awkward when modeling physical systems from first principles,

where physical laws such as mass, momentum, or energy balances or physical laws such as Newton's second law, Ohm's law, and so on are more naturally thought of as relations between variables than as I/O maps. The important uncertainty modeling machinery from robust control can not only be generalized appropriately to implicit systems, but also greatly extended to treat entirely new problems [14].

The implicit form analysis allows for over-constrained problems, such as those involving an uncertain system and a finite number of integral quadratic constraints (IQCs), which may be used to obtain richer signal characterizations. It also provides a framework for the formulation of model validation/system identification questions. For systems with structured uncertainty involving a combination of linear time-invariant and linear time-varying perturbations, an exact test for analysis was obtained, based on a finite augmentation of the original problem. Conditions based on scaled small-gain are also available for this case, and the class of perturbations for which these conditions become necessary has been characterized. A necessary and sufficient condition was obtained for worst-case H_2 -performance analysis under structured uncertainty. This test is a convex feasibility condition across frequency, of the same nature and computational complexity as the corresponding conditions for H_∞ performance. The proof is based on a deterministic characterization of white noise signals, and the necessity proofs involve an extended "S-procedure losslessness" result on quadratic functions on L_2 [14].

6.4 Model validation/ID.

The extension of our analysis algorithms to the implicit formulation is currently being pursued, as are issues arising from the incorporation of data, both of which are necessary for solving ID/model validation problems. Although results with a preliminary algorithm are very encouraging, development is needed to further improve the convergence properties as was done for the standard case. An LMI upper-bound has recently been coded in matlab using the LMI-Lab toolbox. Application of these bounds to computation for experimental control problems at Caltech has recently begun. This work should help lead to an identification methodology appropriate for robust control [10, 7].

Geir Dullerud has also been collaborating with Roy Smith at UCSB on developing an approach for time-domain model validation of an uncertain, noisy continuous-time model with a discrete and finite data record; which is the most directly relevant type of model validation for control as well as VE. They have derived validation conditions in the presence of both LTI and LTV uncertainty which are convex and can be evaluated by LMI methods, have coded optimization software to implement it, and have successfully applied the method to several laboratory experiments [3, 4, 5, 17].

6.5 Robust control and nonlinear extensions.

So far, the most successful applications of robust control techniques have occurred in problem domains (flexible structures, flight control, distillation) where there may be substantial uncertainty in the available models, and the degrees of freedom and the dimension of the input, output, and state may be high, but the basic structure of the system is understood, the uncertainty can be quantified. Nonlinearities are bounded and treated as perturbations on a nominal linear model, or handled by gain-scheduling linear point designs. Robustness for nonlinear systems was proved to be equivalent to the existence of solution to Hamilton-Jacobi equations or nonlinear matrix inequalities [6]. However, computational methods to establish the existence of these solutions have not been developed to a level comparable to their linear counterpart (i.e., existence of solutions of Riccati equations and linear matrix inequalities), and are theoretically intractable, even for the

cases which are easy for linear systems.

The state of the art in industry, as discussed above, still consists in obtaining lower bounds to the performance indices through extensive simulation or local optimization techniques. However these methods require large amounts of computation; standard optimization techniques fail even for small problems, and a search over parameter space exhibits exponential growth with the number of parameters. The methods actually used in industry share two main characteristics: performance specifications are made over a finite time horizon, and the interface between the analysis method and the system is a simulation.

In recent work at Caltech, we have begun to extend the robustness analysis techniques of linear systems, and in particular the associated computational methods, to nonlinear systems. Given the diversity of nonlinear behavior, it is clear that this cannot be done in complete generality and still maintain the efficiency and usability of the methods. We have focused on analysis methods for a specific nonlinear robust performance problem: tracking a trajectory in the presence of noise and uncertainty. Many nonlinear analysis problems of engineering interest can be reduced to such a problem. A common example is an airplane performing an automatic change of altitude and heading. The pilot enters the new heading and altitude and the flight computer determines nominal commands to perform it. A second control loop maintains the airplane around the planned trajectory. The designer has in mind an appropriate path to be completed in a finite predetermined time, and designs the control system accordingly. Since the real system is not exactly the one used for the design, and since it is also subject to noise, the system will not follow the intended trajectory. The question of interest becomes: will the real trajectory, under the worst conditions possible, remain close enough to the nominal one in an appropriate norm?

In [19] we presented a power algorithm to compute a lower bound on the performance index associated with the robust trajectory tracking problem (i.e., the distance from the actual to the nominal trajectory). This algorithm is similar in nature to the one developed for the structured singular value [11, 22], and has similar behavior. Since, as was the case for linear systems, the algorithm is not guaranteed to converge in general, its analysis is done empirically. We test this algorithm by applying it to simulations of real systems. We carry out several different performance tests on two different platforms: the Caltech ducted fan experiment and a simplified model of an F-16 jet fighter. The results of these tests are reported in [19, 18]. These results indicate that without significant additional computation, and avoiding computationally expensive parameter searches, a lower bound on the given performance index can be computed that gives more information on the worst case behavior of the system than the standard Monte Carlo procedures.

We have also begun to develop computable upper bounds for the trajectory generation problem by using rational approximations to nonlinear systems and restricting the types of uncertainties which can enter into the dynamics [20, 18]. At present, these results are still far from being practical, but they are a starting point in developing computational machinery for performing robust modeling of nonlinear systems.

REFERENCES

- [1] R. Braatz, P. Young, J. Doyle, and M. Morari. Computational complexity of μ calculation. *IEEE Transactions on Automatic Control*, 1993.
- [2] J. Doyle, M. Newlin, F. Paganini, and J. Tierno. Unifying robustness analysis and system id. In *Proc. IEEE Control and Decision Conference*, 1994.
- [3] G. E. Dullerud and R. S. Smith. A continuous-time extension condition. *IEEE Transactions on Automatic Control*, 1995. (to appear).
- [4] G. E. Dullerud and R. S. Smith. Experimental application of time domain model validation: Algorithms and analysis. *Int. J. of Rob. and Non. Con.*, 1995. (submitted).
- [5] G. E. Dullerud and R. S. Smith. The experimental validation of robust control models for a heat experiment: A linear matrix inequality approach. In *Proc. IEEE Control and Decision Conference*, 1995.
- [6] W.-M. Lu and J. C. Doyle. Robustness analysis and synthesis for uncertain nonlinear systems. In *Proc. IEEE Control and Decision Conference*, pages 787–792, 1994.
- [7] J. C. Morris and M. P. Newlin. Model validation in the frequency domain. In *Proc. IEEE Control and Decision Conference*, 1995.
- [8] M. Newlin. *Model Validation, Control, and Computation*. PhD thesis, California Institute of Technology, 1995.
- [9] M. P. Newlin and S. T. Glavaski. Advances in the computation of the μ lower bound. In *Proc. American Control Conference*, 1995.
- [10] M. P. Newlin and R. S. Smith. A generalization of the structured singular value and its application to model validation. *IEEE Transactions on Automatic Control*, 1995. (to appear).
- [11] A. Packard and J. C. Doyle. The complex structured singular value. *Automatica*, 29(1):71–109, 1993.
- [12] F. Paganini. Analysis of systems with combined time invariant/time varying structured uncertainty. In *Proc. American Control Conference*, 1995.
- [13] F. Paganini. Necessary and sufficient conditions for robust h_2 performance. In *Proc. IEEE Control and Decision Conference*, 1995.
- [14] F. Paganini. *Sets and Constraints in the Analysis of Uncertain Systems*. PhD thesis, California Institute of Technology, 1995.
- [15] F. Paganini and J. Doyle. Analysis of implicit uncertain systems, parts i and ii. CDS Technical Report CIT/CDS 94-018, California Institute of Technology, 1994.
- [16] F. Paganini and J. Doyle. Analysis of implicitly defined systems. In *Proc. IEEE Control and Decision Conference*, pages 3673–3678, 1994.
- [17] R. S. Smith and G. E. Dullerud. Validation of continuous-time control models by finite experimental data. In *Proc. American Control Conference*, 1995. (submitted, IEEE T. Automatic Control).

- [18] J. Tierno. *Numerically Efficient Robustness Analysis of Trajectory Tracking for Nonlinear Systems*. PhD thesis, California Institute of Technology, 1995.
- [19] J. Tierno, R. M. Murray, and J. C. Doyle. An efficient algorithm for performance analysis of nonlinear control systems. In *Proc. American Control Conference*, 1995.
- [20] J. E. Tierno and R. M. Murray. Robust performance analysis for a class of uncertain nonlinear system. In *Proc. IEEE Control and Decision Conference*, 1995.
- [21] P. Young, M. Newlin, and J. Doyle. Computing bounds for the mixed μ problem. *International Journal of Robust and Nonlinear Control*, 5:573–590, 1995.
- [22] P. Young, M. Newlin, and J. Doyle. Let's get real. In *Robust Control Theory*. Springer Verlag, 1995. pp. 143–174.
- [23] P. Young, M. Newlin, J. Doyle, and A. Packard. Theoretical and computational aspects of the structured singular value. *Inst Sys Contr Inf Eng*, 1995. (to appear).

M. MILITARY APPLICATIONS OF DIAMOND

Richard B. Kaner

Department of Chemistry

University of California, Los Angeles

Los Angeles, CA

MILITARY APPLICATIONS OF DIAMOND

A. INTRODUCTION

Diamond has fascinated people throughout history due to its unique combination of physical properties ranging from extreme hardness and high thermal conductivity to great transparency in the visible, UV and X-ray parts of the spectrum. Diamond has always been considered valuable due to its appealing appearance and relative scarcity. However, modern techniques of synthesis for both single crystals and thin films have now surpassed the amount of diamond extracted from the earth each year and promise to lower diamond's cost and availability. Because of this, applications where diamond may be considered the ideal material, but have until now been precluded due to cost, need to be reconsidered. Novel applications with military significance should also be considered. This paper will first look at the special physical properties of diamond, next discuss its availability from mining and laboratory syntheses, then present current applications ranging from cutting tools to protective coatings, and finally, explore important future applications of diamond including domes to protect infrared guided missiles, substrates for thermal management, electron emitters for flat panel displays, and as semiconductors for electronics.

B. PHYSICAL PROPERTIES OF DIAMONDS

Diamond is by far the hardest material known, as shown in Figure 1 [1]. On the conventional Mohs scale (which is logarithmic) it is assigned a 10, being considerably harder than the mineral corundum (Al_2O_3), which is assigned a 9. In fact, on the Vickers hardness scale (which is linear) diamond exhibits a hardness of $1 \times 10^{11} \text{ N/m}^2$, making it about five times harder than corundum, with a Vicker's hardness of $2.2 \times 10^{10} \text{ N/m}^2$. The second hardest material after diamond is cubic boron nitride, BN ($4.8 \times 10^{10} \text{ N/m}^2$), followed by zirconium carbide, ZrC ($3 \times 10^{10} \text{ N/m}^2$) and silicon carbide, SiC ($2.5 \times 10^{10} \text{ N/m}^2$). Diamond is the most abrasive resistant mineral known, as well as possessing the highest modulus of elasticity [2].

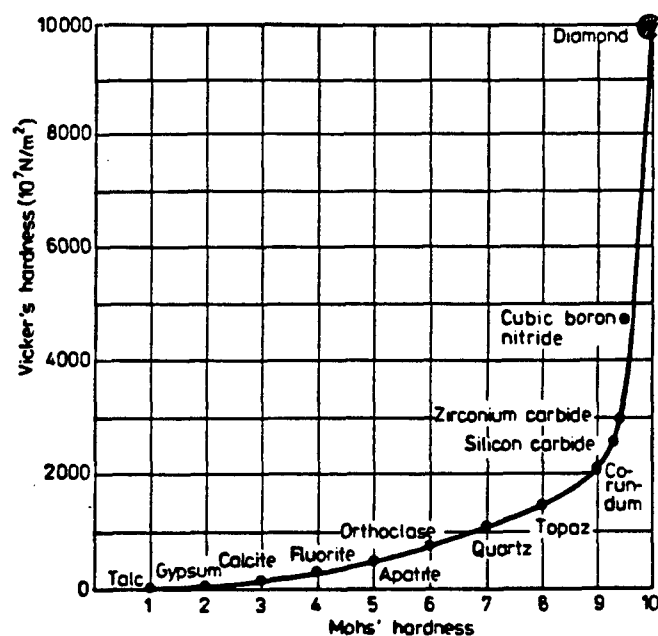


Figure 1. Hardness of Diamond on Logarithmic Moh's Scale vs. Vicker's Hardness (which is Linear) [1]

The hardness of diamond stems from its light weight and short three-dimensional covalent bonds. Diamond is composed of pure carbon (98.9 percent ^{12}C and 1.1 percent ^{13}C , giving an average molecular weight of 12.011 amu) arranged in a cubic zinc blend type structure (Figure 2). An equivalent description of diamond is a face-centered cubic carbon lattice with half the tetrahedral sites filled with other carbons. This leads to the highest number density of atoms of any material (1.76×10^{23} atoms/cm³). Each carbon atom is sp^3 hybridized forming four single bonds to other carbons situated at the ideal tetrahedral angles of $109^\circ 28'$. The C-C bond lengths are $1.54\text{\AA}$ leading to a density of 3.515 g/cm^3 . A summary of the physical properties of diamond with a comparison to zinc sulfide, sapphire (corundum) and silicon is given in Table 1.

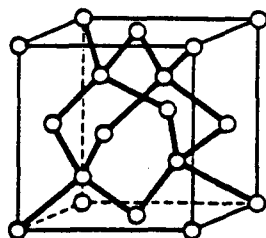


Figure 2. Face Centered Cubic Unit Cell of Diamond

Table 1. Physical Properties of Diamond and Other Materials

Material	Trans- mission (μm)	Density (g/cm^3)	Hardness (kg/mm^2)	Strength (MPa)	Young's Modulus (GPa)	Thermal Expansion (K^{-1})	Thermal Conduc- tivity (W/cmK)	Dielectric Constant	Loss Tangent ($\tan \delta$)	Electrical Resis- tivity (Ωcm)	Band Gap (eV)
Diamond (C) [3]	0.25- 100 ^a	3.52	10,000	3,000	1,050	1.3×10^{-6}	26	5.7	<0.0004	10^{16b} [4]	5.7
Zinc Sulfide (ZnS)	0.3-13	4.09	160	70	88	7.0×10^{-6}	0.27	8.4	0.0024	10^{10} [6]	3.6 [6]
Sapphire (Al_2O_3)	0.14 ^b	3.99 ^c [5]	2,200	400	344	5.3×10^{-6}	0.34	9.4	0.00005	$\sim 10^{11}$ [1]	8.9 ^c [5]
Silicon (Si)	1.2-100	2.33 [6]	1,150	120	131	2.6×10^{-6}	1.63	11	~ 0.009	3×10^5 [6]	1.11 [6]

^a Diamond has a weak absorption in the 3-5 μm region.

^b The electrical resistivity of diamond is sensitive to impurities.

^c At 1000 °C.

The phase diagram for diamond is shown in Figure 3 [4, 8]. Note that diamond is only metastable under ambient conditions. However, its conversion to graphite, the thermodynamically stable form of carbon at room temperature (see Figure 4), has an extremely high activation barrier and is exceedingly slow under ambient conditions. However, above 1500 °C graphitization does become a problem. For example, at 2100 °C 0.1 carat (1 carat = 0.2 g) of single crystal diamond converts to graphite in under three minutes [9, 10]. The thermal stability is decreased by impurities and inclusions. In fact, diamond decomposition is catalyzed by elements that react chemically or dissolve it, including oxygen, iron, cobalt, and nickel. Diamond is not stable in air much above 600 °C due to its rapid reaction with oxygen to form gaseous carbon dioxide. Iron, cobalt, and nickel are used as catalysts for diamond synthesis under high pressures and will catalyze reconversion to graphite at low pressures. Diamond is highly resistant to chemical attack, such as by acids, even at elevated temperatures [11]. Only molten hydroxides and the salts of oxy-acids have corrosive effects on diamond. Hydrogen can be used to increase the stability of diamond likely due to chemisorption on its surface [12].

Diamond is an electrical insulator with a band gap of about 5.7 eV and a calculated resistivity of 10^{70} Ωcm for a perfect crystal [12]. However, real diamond has many defects and impurities that lower its resistivity to a typical range of 10^{14} - 10^{16} Ωcm. Boron can be used to p-dope diamond to further lower its resistivity to 10^3 Ωcm [13].

Diamond's brilliance, which makes cut gems especially attractive, stems from its high index of refraction and high dispersion, due to a large difference in refractive indices for red and violet light [1]. Total internal reflection of light occurs at an angle of incidence as small as 24.5°. Only minerals such as rutile (TiO₂), garnet (YAl₂O₄) and cubic zirconia (ZrO₂, now commonly used as an inexpensive diamond substitute in jewelry), exceed the dispersion of diamond. Zircon (ZrSiO₄) shows a comparable dispersion and is also used as a diamond substitute in jewelry. Diamond has excellent long wavelength transmission properties. Its cutoff edge in which a 2-mm thick window has <10 percent transmission extends well beyond 100 μm, giving it a wider window than any ceramic or alkali halide salt [3] (see Figure 5).

Diamond is the best heat conductor known at room temperature with a value of 26 W/mK [1] (see Figure 6). This value is four times higher than copper. The actual value of thermal conductivity depends on the perfection of the crystal structure, grain size, and impurities. The more perfect the diamond, the higher is its thermal conductivity. Synthetic diamond films often have thermal conductivities as low as 2 W/mK (see Figure 7). Work at General Electric using CVD (chemical vapor deposited) diamond from pure ¹³C precursors

followed by high pressure diamond growth, resulted in a single crystal diamond with a thermal conductivity of 33 W/mK due to low phonon scattering [14].

Because of its excellent thermal conductivity diamond feels cold when touched. Thus, a simple method to distinguish diamond from "paste," i.e., glass, which is a relatively poor heat conductor, is by touching a specimen with one's tongue. Alternatively, when breathed on, diamond mists and clears much more readily than glass. The specific heat of diamond at room temperature is about 550 J/kgK [1].

Diamond possesses an extremely small coefficient of thermal expansion. Based on X-ray diffraction measurements, diamond's coefficient of thermal expansion is $1.3 \times 10^{-6}/\text{K}$ at room temperature and $7 \times 10^{-6}/\text{K}$ at 1400 °C [1].

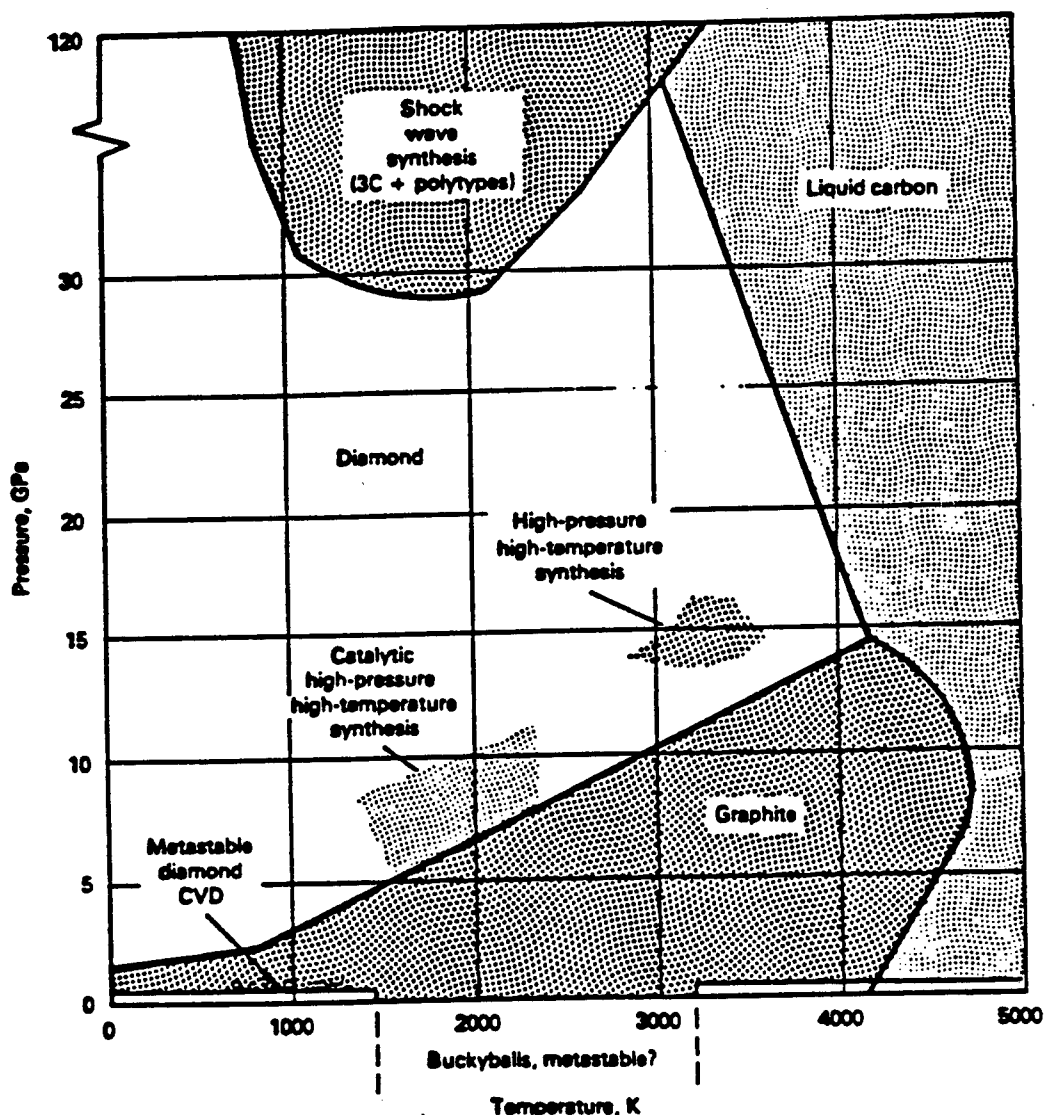


Figure 3. Carbon Phase Diagram [4]

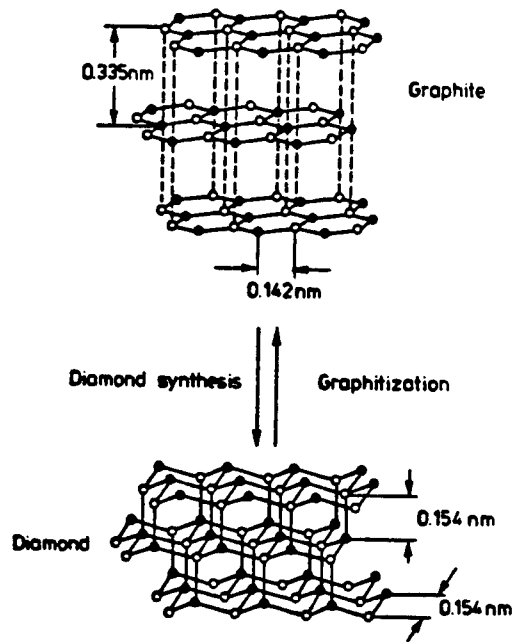


Figure 4. Graphite, the Thermodynamically Stable Form of Carbon, Under Ambient Conditions Can Be Converted to Metastable Diamond at Elevated Temperatures and Pressures (Generally Greater Than $1500\text{ }^{\circ}\text{C}$ and 6 GPa) [1]

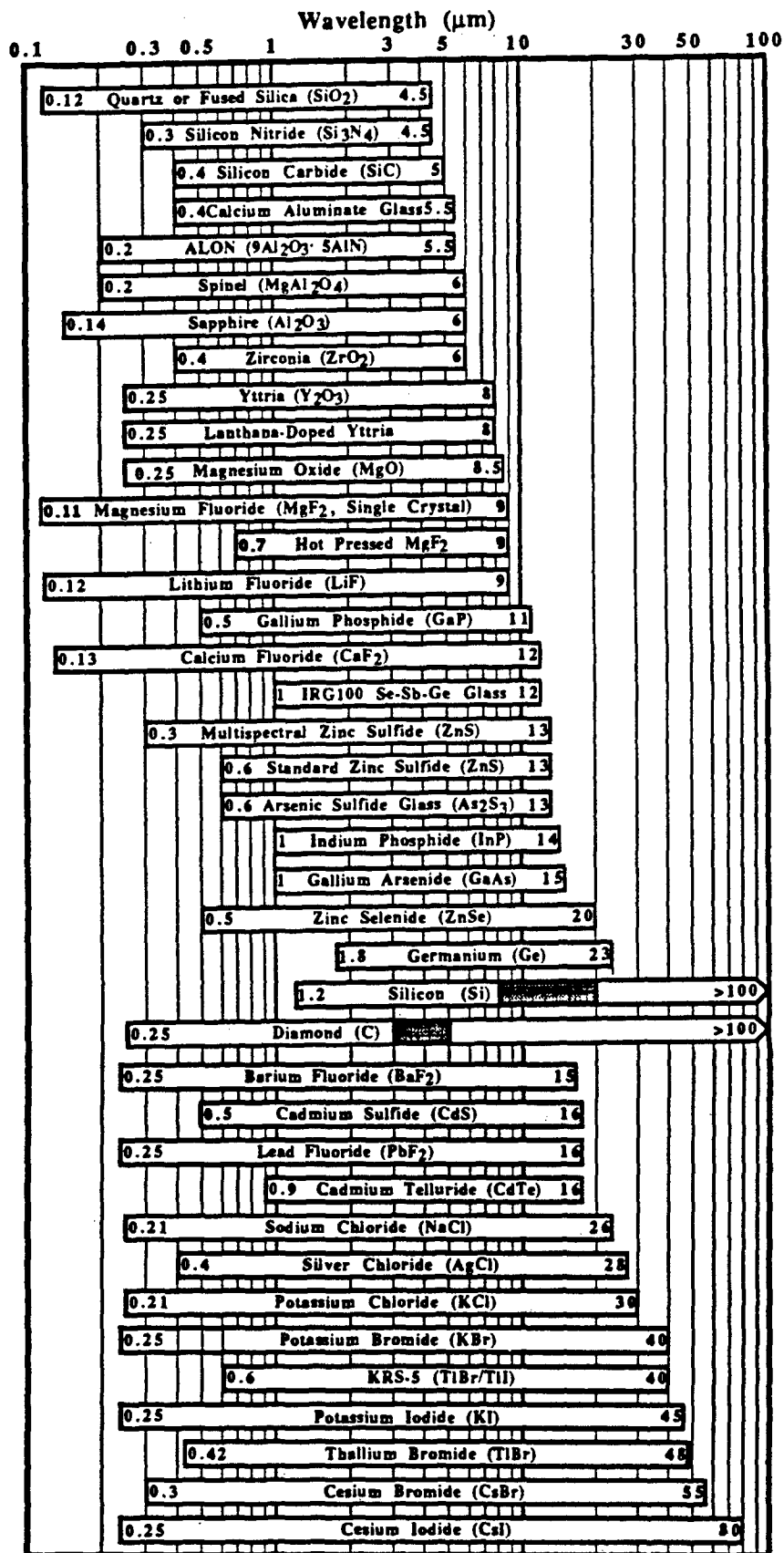


Figure 5. Transmission Windows for Selected Materials [3]

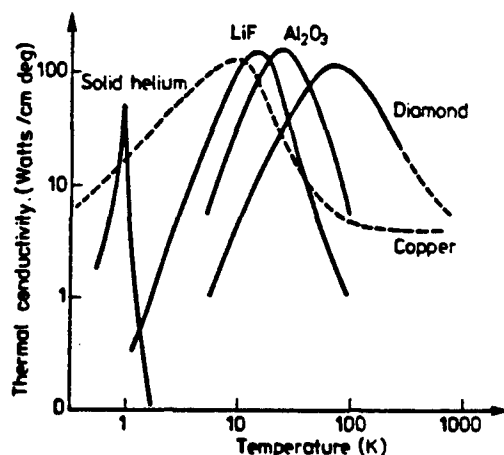


Figure 6. Thermal Conductivity of Diamond as a Function of Temperature Compared to Copper, Alumina, Lithium Fluoride, and Solid Helium

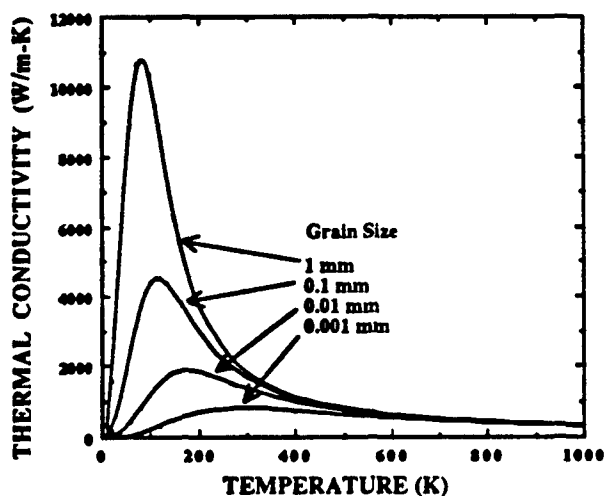


Figure 7. Calculated Grain Size Dependence of Thermal Conductivity for Diamond

C. THE AVAILABILITY AND SYNTHESIS OF DIAMOND

Diamond has been known for thousands of years with relatively large deposits occurring in Zaire, South Africa, Russia, and Australia [1]. It generally takes excavation of about 10 tons of rock to produce one carat (0.2 g) of diamond. Over 50 megacarats of diamond are mined from the earth each year. The price ranges from as low as \$2 per carat for diamond grit used for abrasives up to \$10,000 or more per carat for gem quality material. Larger stones cost even more per carat due to their scarcity.

Synthetic diamond has been known since the 1950s when scientists at General Electric subjected graphite to pressures of ≥ 55 kbar at ~ 2000 °C using transition metal catalysts [15].

This process has been refined over the years and is carried out on a megacarats scale each year. Although gem quality stones can be made in the laboratory, their production is slow, energy intensive and, therefore, currently uneconomical. Essentially all high pressure synthetic diamonds currently produced are used to coat cutting tools. A few gem-quality synthetic diamonds, up to two carats, are produced in Japan by the high pressure/high temperature process for use as heat sinks.

Chemical vapor deposition (CVD) is rapidly becoming another important synthetic route to diamonds. Currently, over 1 megacarats is produced annually as coatings for wear resistance. The mechanism behind CVD diamond growth is still not fully understood [16]. Generally, about 1 percent of a hydrocarbon, such as methane, is used in an atmosphere containing 99 percent hydrogen. The hydrogen suppresses the growth of graphite, likely due to its reacting with graphite much faster than it reacts with diamond. Hydrogen also prevents reconstruction of the diamond surface, enabling continuous diamond growth.

Many methods have been developed for growing CVD diamonds. These can be classified into two major categories: diffusion-controlled processes and convection-driven processes. Diffusion-controlled processes, such as hot filament or microwave plasmas, enable very high quality diamond films to be grown, but suffer from slow growth on the order of a few microns per hour. Much higher growth rates are possible using convection-driven processes, such as an arcjet, but film quality suffers.

Hot filament methods use a tungsten, tantalum, or rhenium wire resistivity heated to ~ 2100 °C in an evacuated chamber with a flow feed gas mixture of typically 1 percent methane and 99 percent hydrogen [17, 18]. Diamond is deposited on a heated substrate located 5-15 μm away. More than 5 percent of the hydrogen decomposes into atoms, which then activates the methane for deposition and removes sp^2 graphitic carbon from the surface. Active species for diamond growth rate are short-lived and can only diffuse over small distances. This leads to low deposition rates of 0.5-20 $\mu\text{m/hr}$, with 5 $\mu\text{m/hr}$ being typical. Although carbon monoxide, water, or oxygen additions can increase the growth rate up to 40 $\mu\text{m/hr}$, high-quality diamond has so far only been grown at ~ 1 $\mu\text{m/hr}$. Although plasma methods are better for ultra-clean films, hot filament methods may be useful for wear resistance applications [17].

Low-pressure microwave plasma methods appear most promising for deposition of high-quality diamond films. A mixture of 0.3-1 percent methane, up to 1 percent oxygen and hydrogen at 0.1 atmospheres or less is excited to a plasma using high frequency (e.g., 2.45 GHz) microwave radiation. Diamond is deposited on a substrate, such as (100) Si heated to 450-1150 °C. Deposition areas can be large, up to 75 mm² so far, while deposition rates for high quality diamond are small, ~1 µm/hr [19]. Since the plasma forms at the convergence of the microwave energy and the feed gas, the substrate can be located in the middle of the plasma away from the container walls. This avoids contamination from, for example, glass walls. A prototype reactor system using six tubular reactors running in parallel has been developed in Japan.

Other diffusion controlled methods for diamond synthesis and their limitations include: low-pressure glow discharge, which uses a high bias voltage between the substrate and a filament and appears limited to growth rates of < 0.1 µm/hr [19], and RF thermal plasmas, which are less stable than microwave plasmas and generally produce poor quality diamond films [20].

Among the convection-driven processes, arcjets are well advanced. In the DC arcjet method, a methane/hydrogen feed gas flows at high velocity through the arc generated by passing a high current across an annular electrode (see Figure 8). The activated gas mixture is sprayed onto heated substrates resulting in growth rates up to 100 µm /hr for 1 cm² areas [8, 21]. The rate decreases as larger substrates are used. Norton has deposited 4-inch diameter films up to 1-mm thick for thermal management and tribological applications [22]; 800 µm thick optical quality films have also been produced.

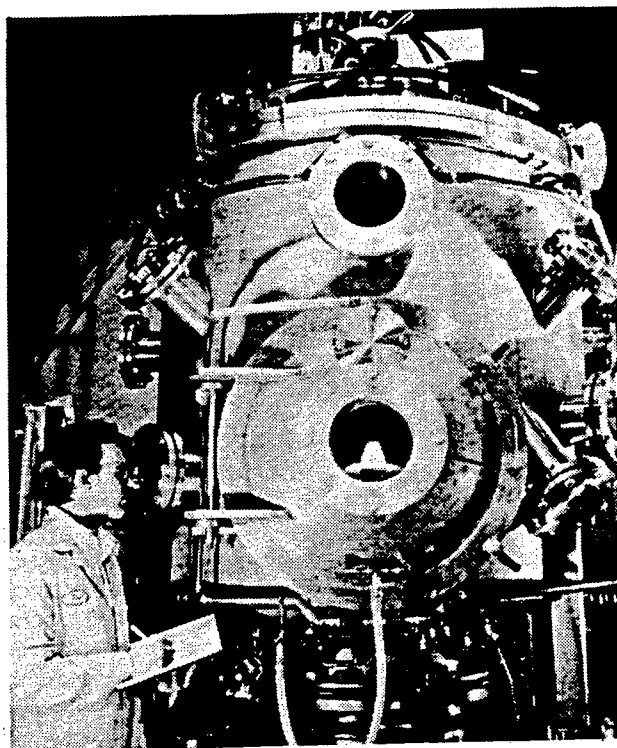


Figure 8. Industrial DC Arcjet Diamond Deposition Chamber Showing Jet Impinging on Substrate

A commercial oxygen/acetylene welding torch can also be used to deposit diamond [19]. A slight excess of acetylene leads to uncombusted carbon, which can form diamond in the reducing part of the flame. Although combustion flame synthesis has the advantage of very simple equipment, which can be used in air rather than in an enclosed chamber, it suffers from the disadvantages of inhomogeneous diamond deposits, a need to cool the substrate and low conversion rates of acetylene (<1 molecule out of 10,000). This leads to a cost estimate for diamond produced from welding gases of \$2-\$20 per carat. This can be compared to the cost of diamond produced by the conventional high temperature/high pressure method of about 20 cents per carat (for diamond grit).

D. CURRENT APPLICATIONS OF DIAMOND

The jewelry market accounts for less than 5 percent of diamond produced worldwide. By far the largest uses of diamond is as abrasives, drillbits, and toolstones. Diamond tools include dressers, shapers, dies, indentors, drills, and scapels, which can be used to cut and shape materials ranging from rock to glass to human tissue [1]. In fact, diamond-coated scapels are known for making cleaner incisions for surgery (e.g., eye surgery) than even lasers. Surprisingly, the only common material diamond cannot cut is steel. This is due to a

reaction between carbon and iron at the elevated temperatures produced during cutting operations. For this reason, cubic boron nitride, which is both isoelectronic and isostructural with diamond (a BN unit replaces every two carbons), is used. Cubic boron nitride does not occur naturally and must be produced in megacarats quantities each year from high temperature/high pressure reactions [23]. The current price of BN ranges from about \$2-\$5 per carat.

Other applications for diamond include low friction coatings for ball bearings and other tribological uses, wear resistant coatings, and diamond anvil cells for high pressure research.

E. FUTURE APPLICATIONS OF DIAMONDS

1. Missile Domes

Infrared detection has important military applications ranging from heat-seeking missiles to night-vision goggles. All objects emit infrared radiation due to thermally induced vibrations. The wavelength of maximum emissions, $\lambda_{\max}$, can be calculated from the Wein displacement law which is valid above 100 K [3]:

$$\lambda_{\max} \approx \frac{2.878 \times 10^{-3} \text{ mK}}{T} \quad (1)$$

At room temperature, 298 K, $\lambda_{\max}$ is about 9.6 μm , which falls in the long wavelength atmospheric window (see Figure 9). At 500 °C (773 K), the temperature of jet aircraft exhaust, $\lambda_{\max}$ is about 3.6 μm , which falls in the midwave infrared atmospheric window. To protect the infrared detector in a guided missile, an infrared transparent dome is needed as shown in Figure 10. Currently, the most common dome material is magnesium fluoride, due to its wide transparency in the infrared region from 0.11-9 μm . Unfortunately, there is generally a tradeoff between long wavelength transmission and desirable mechanical properties, since materials containing heavier atoms have longer wavelength cutoffs (see Figure 5), but longer, weaker bonds. Thus, magnesium fluoride missile domes must be protected during aircraft flights, which reduces their effectiveness and forces pilots to "arm" their missiles by blowing off the protective covers during an engagement. This is costly both in time, which is critical during combat, and in expense, since once exposed, the magnesium fluoride domes need replacing due to erosion from sand and water droplets.

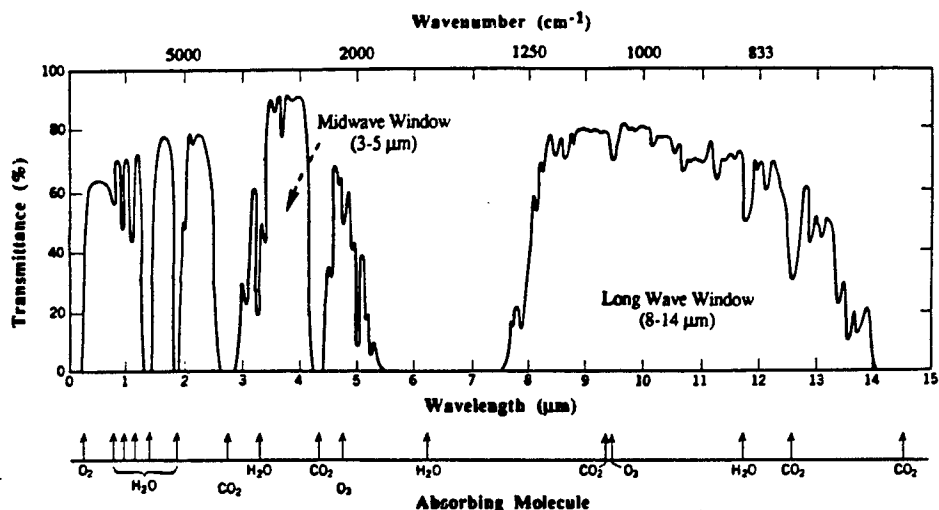
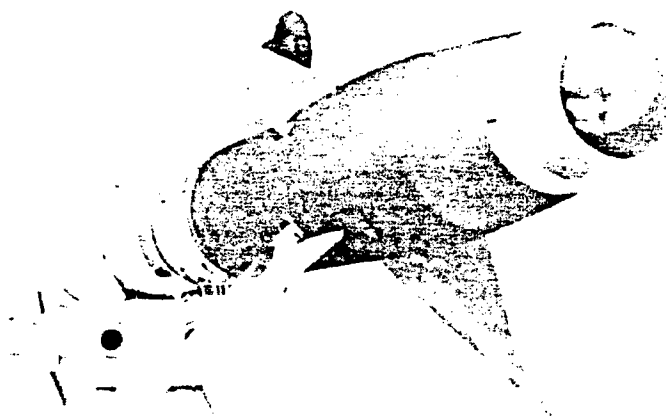


Figure 9. Infrared Transmission Throughout the Atmosphere for a 1.8-km Pathlength at Sea Level with 40 Percent Humidity [3]



Source: Naval Air Warfare Center, China Lake, CA

Figure 10. A Sidewinder Missile Showing the Protective Dome Covering the Infrared Detector [3]

The only exception to the tradeoff between long wavelength transmission and good mechanical properties is diamond. Diamond has the best mechanical performance of any material along with good long wavelength transmission due to its high symmetry, which does not allow a change in dipole moment (a necessary condition for infrared absorption). Diamond is especially resistant to erosion by particulate matter. In a sand erosion test at 26

m/sec impact speed, the mass lost by diamond was 20,000 times less than that of silicon nitride Si_3N_4 (at 47 m/sec) and 3,000,000 times less than alumina, Al_2O_3 (at 34 m/sec) [24]. Unfortunately, pure diamond domes produced by CVD processes take several days to produce, cool, and polish, and are estimated to cost around \$100,000 each. Since no obvious civilian use for these domes are known, the price is unlikely to be reduced significantly due to a lack of a sufficiently large demand, and hence it is unlikely that this application will become economically feasible.

An alternative dome design is to use a thin protective coating of diamond on an infrared transparent window. Unfortunately, most infrared window materials are soft, low melting compounds such as zinc sulfide or zinc selenide, which degrade under diamond reactor conditions, such as a hydrogen plasma at 950 °C. The inverted structure of a thin layer of zinc selenide grown on diamond at 950 °C, will likely delaminate due to a large difference in thermal expansion coefficients leading to high sheer stresses estimated at up to 130 MPa [25]. Interestingly, a successful method for attaching diamond to zinc selenide windows, called optical brazing, is known [26, 27], as shown in Figure 11. In this novel process, a diamond film is grown on a smooth piece of silicon. The rough side of the diamond film is coated with an arsenic/selenium/sulfur glass, which when softened, fills in all crevices and can be used to glue the diamond to the zinc selenide window. The silicon substrate is then etched away with hydrofluoric acid leaving a smooth diamond face. As the remaining diamond arsenic/selenium/ sulfur and zinc selenide have nearly the same index of refraction, the scatter due to roughness of the diamond layer is reduced to a negligible amount. This appears to be a promising process for coating optical materials with diamond. However, it is unclear if such sandwiches can withstand the impact and temperature conditions needed for use in missile domes.

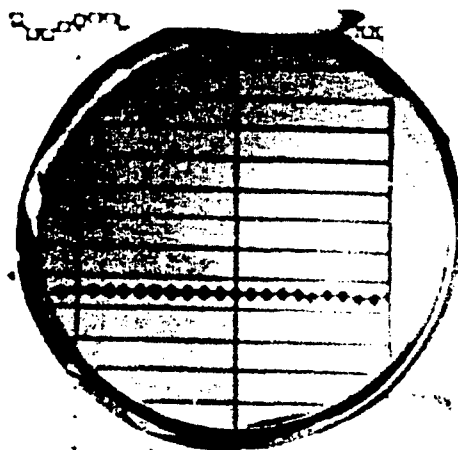


Figure 11. A 7-mm-Thick Diamond Coating Attached to a 2-Inch Selenide Window by Optical Brazing at Westinghouse Corp. Is Optically Clear [26]

2. Thermal Management Substrates

Semiconductor chips, the heart of the electronics revolution, have continued to become smaller and smaller by increasing the packing density of transistors. However, semiconductor chips have just about reached the limit where the problems of electrical interconnections dominate over size, weight, and operating speed. The interconnect problems can be overcome by three-dimensional multichip modules (MCMs), as shown in Figure 12. The shorter lengths of interconnects enable faster electronic systems packed into smaller volumes. However, a new problem of higher power dissipation arises due to the higher packing densities and higher operating frequencies. Thus, efficient thermal management substrates become a necessity. Diamond, with the highest room temperature thermal conductivity, can be considered the ideal substrate as long as the cost is reasonable and the quality high.

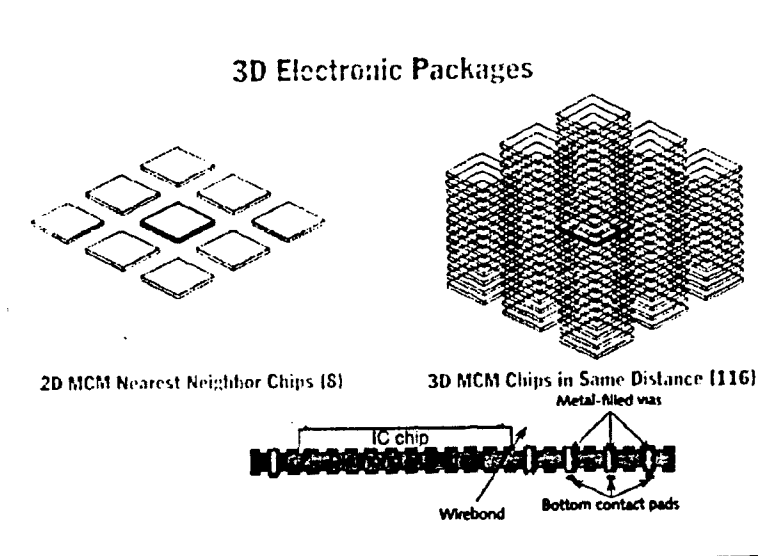


Figure 12. A Schematic Diagram of Three-Dimensional Electron Multichip Module (MCM) Packaging for Semiconductors

Current CVD diamond films suffer from rough surface morphology with 2-5 μm variations (see Figure 13), which limits the heat diffusion efficiency. Several novel approaches, including laser/ion beam treatments, hot-metal methods, and/or chemical-assisted mechanical polishing will likely lead to a solution to this problem [28, 29]. The cost of polishing CVD diamond can be currently estimated at \$1-\$2 per cm^2 with reductions likely with future scale-up.

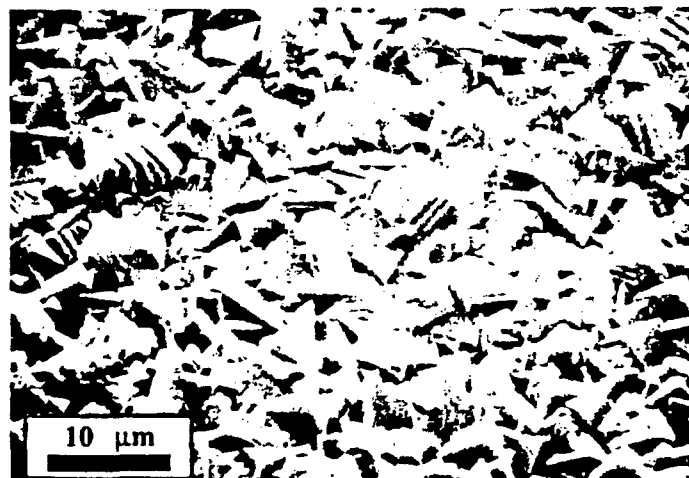


Figure 13. Course Microstructure of Microwave Plasma-Grown Diamond as Observed by Scanning Electron Microscopy

Another difficulty with diamond is making interconnections directly through the material. A novel approach using laser drilling in a liquid medium has recently been developed [30]. Uniform, chemically clean holes can be drilled in a diamond substrate submerged in water using a Q-switched Nd-YAG laser (see Figure 14).

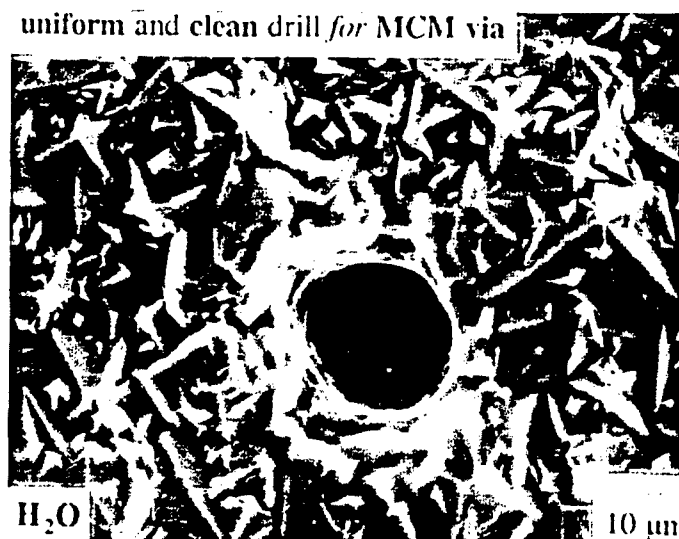


Figure 14. Laser Drilling Under Water Through a Diamond MCM Substrate [31]

For thermal management substrates to spread heat efficiently they will need to be 0.5-1 mm thick [31]. Growth rates of 50-100 $\mu\text{m/hr}$ over 3-6 inch diameter areas will be needed. In addition, the cost must be less than \$10 per m^2 , or about \$5 per carat for CVD diamond to replace other potential substrates such as aluminum nitride (AlN) or sapphire (Al_2O_3). A company in Norburn, MA, named Astex, Inc., is already producing synthetic diamond

substrates for about \$10 per carat using a 5-kW microwave reactor. Increasing the reactor power to 250 kW could, in theory, produce diamond for about \$2 per carat. Even adding in the cost of polishing, this looks like a future winning technology.

3. Electron Emitters for Flat Panel Displays

Diamond cathodes hold exciting promise for use as field emitters for display devices. Field emission displays could potentially overtake liquid crystals in the flat panel market. The flat panel display market currently exceeds \$9.6 billion per year and is expected to generate more than \$14 billion before the turn of the century [32].

The market driver is pocket size computers and other electronic devices. A major impact for the military lies not only in portable devices for individual soldiers but also in heads up displays for airplane pilots, tank drivers, and radar operators. Flat panel displays are the enabling technology for many applications from laptop computers to global positioning systems. In the future, they could make possible such things as personal digital assistants and virtual reality driven robots for handling nuclear materials or hazardous waste. Military applications will likely follow rather than lead the market, especially if current liquid crystal displays continue to dominate the field. This is because Japanese manufacturers currently have a virtual lock on the world market with the U.S. share at a mere 2-3 percent. Sharp, which leads the world in production with 95 percent of the market in active matrix liquid crystal displays (LCDs), has said that is not interested in low-volume type custom displays such as those needed by the US military [32]. Hence it may be in DOD's best interest to invest in other emerging technologies.

The drawbacks of active matrix LCD's include high manufacturing costs, large power requirements, and narrow viewing fields, whose images can be tough to see in bright sunlight. The high costs of manufacturing stem from a rather complex design (see Figure 15), which utilizes a light source (generally continuous full power backlighting), a polarizer plate, a circuit plate followed by a solution of nematic phase liquid crystals, a set of color filters, and a second polarizer. Each pixel of an active matrix LCD has a capacitor attached. Light is free to pass through the unoriented liquid crystal. However, when a current is passed across the dielectric nematic liquid crystal, the molecules twist, blocking polarized light and turning a given pixel black. Gray scales can be obtained by changing the voltage across each pixel, which effects the degree of twist and hence the intensity of the transmitted light. Color displays are obtained by adding red, green, and blue filters to each pixel.

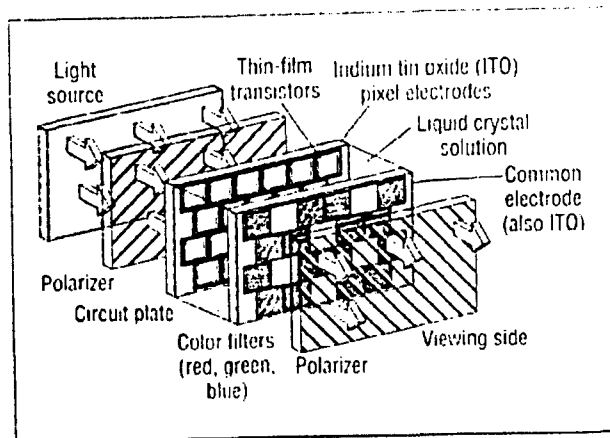


Figure 15. Schematic Diagram of an Active Matrix Liquid Crystal Display [32]

Field emission displays could potentially overcome many of these drawbacks [32, 33]. Field emitters work like conventional cathode ray tubes in that electrons are liberated from a cathode and impinge on a phosphor-coated anode to produce an image (see Figure 16). Since cathodoluminescence is a very efficient way of generating light, the power consumption of field emission displays will be less than one watt at 60 Cd/m^2 , while a typical LCD needs 4 watts. To create a display, each pixel requires its own emitter which can now be made on a $2\text{-}\mu\text{m}$ scale (Figure 17). The emission current, which determines the display brightness, is strongly dependent on the work function of the emitting material. Thus clean materials and uniformity of the emitter surfaces are critical.

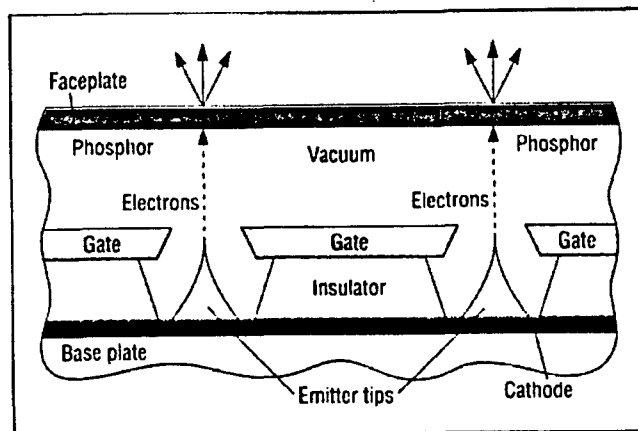
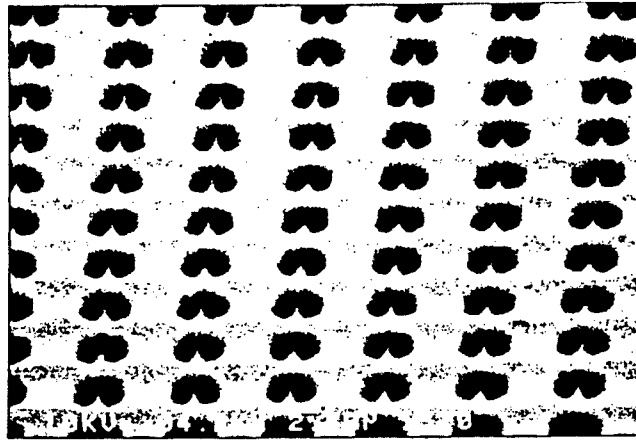


Figure 16. Schematic Diagram of Field Emission Cathode—Voltage Applied to the Gate Causes Emission of Electrons from the Tips [32]



**Figure 17. An Array of Diamond Field Emitters
Created by FED Corp.— Up to 1,600 Emitters
Can Make Up a Single Pixel [32]**

Diamond is chemically quite inert requiring only a 10^{-4} torr vacuum to achieve an adequate mean free path for emitted electrons (molybdenum and silicon microtips require $\sim 10^{-7}$ torr). This simplifies manufacturing by easing requirements for spacer design, substrate cleaning and device sealing. Fabrication can be done using simple printed wire board type lithography. Emission currents up to 100 mA/mm^2 at 350 V have been obtained using laser ablated nanocrystalline diamond films [32]. The challenge is how to mass-produce displays with uniform emission across the entire surface. So far SI Diamond has demonstrated a 125×125 pixel monochrome prototype [32]. This appears to be a promising technology.

4. Semiconductors for Electronics

Wide band gap semiconductors made of diamond would have many advantages over traditional semiconductors such as silicon. They would be faster, they could work with high-power output at microwave or millimeter wave frequencies, could operate at high temperatures, would be radiation hard, and should have low losses [34]. The high dielectric strength of $>5 \times 10^6 \text{ V/cm}$ would allow diamond transistors to operate at much higher voltages than conventional semiconductors. At the same operating voltages used for silicon, the charge carrier path length in diamond would be shorter, leading to less resistance and therefore less loss. Since the atomic mass of carbon is lower than that of silicon, the charge carrier velocity would be over twice that of silicon and at least equal to the peak reached in GaAs. Additionally, in contrast to GaAs, an electric field up to 1,000 times greater would not hurt the carrier velocity, thus potentially enabling high-power output at millimeter or

microwave frequencies. The high chemical binding energies in diamond should make it much less susceptible to radiation damage than conventional semiconductors.

Diamond could be used in high power capacitors due to its low loss tangent [34]. Conventional semiconductors cannot be used as capacitors due to thermally excited carrier leakage and high loss tangents.

Computer memory chips could be improved using diamond based semiconductors [34]. Conventional silicon-based dynamic random access memory (DRAM) suffers from charge loss due to thermally generated carriers. This necessitates that the charge be refreshed every 20 μsec . A wide band gap memory based on a diamond capacitor in series with a diamond transistor would have no measureable leakage, at least up to 200 °C and perhaps to much higher temperatures. Thus the refresh circuitry currently used in computers could be eliminated reducing power requirements by at least 50 percent. Even complete loss of power would not destroy a diamond based memory.

Diamond-based semiconductors await improved methods for growing electronically pure materials economically. In addition, although boron can be used effectively as a p-dopant for diamond since its acceptor levels lie close to the valence band edge in diamond (within 0.35 meV), no good n-dopant has yet been found since the donor levels in nitrogen are quite far from the conduction band edge in diamond (1.7 eV) [35]. Perhaps substitutional phosphorus or interstitial lithium or sodium n-doping will be better.

5. Other Applications

Diamond's unique combination of physical properties can lead one's imagination to envision applications ranging from the very practical to the highly fanciful. On the practical side are precision instruments made from single crystal diamond and diaphragms for loudspeakers and microsensors made from chemically deposited diamond films.

Loudspeaker diaphragms exploit the extremely high speed of sound in diamond [36]. When alumina diaphragms are coated with diamond the upper frequency cut-off is raised from 35 to 50 kHz [37]. Deposition of diamond on silicon followed by etching away the substrate with hydrofluoric acid further increases the frequency cut-off to 80 kHz.

Microsensors for rugged electronics is a potential application for doped diamonds. For example, when strain is applied to a boron-doped diamond substrate, the resistance changes. This piezoresistance can enable diamond to be used as a strain gauge for pressure or acceleration sensing [38]. A picture of such a sensor developed at Vanderbilt University is shown in Figure 18.

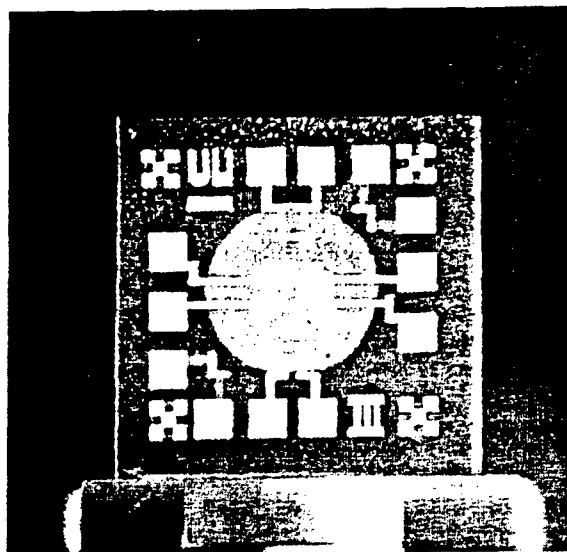


Figure 18. Diamond Pressure Sensor Composed of a Circular Diamond Piezoresistor and Metal Interconnects [31]

A more fanciful application of diamond is as a coating for naval ships. In principle, diamond's exceedingly low coefficient of friction could reduce drag resistance and increase ship speed and fuel economy. By coating propeller blades with diamond, cavitation damage could be reduced due to the excellent mechanical properties of diamond. Unfortunately, it is unlikely that the problems associated with uniformly coating large areas with diamond will be overcome anytime soon. Additionally, it appears barnacles still grow on diamond surfaces [39].

Another application with military interest is windows for radiation detectors. These could operate in space and be useful for detecting all types of ionizing radiation as well as such things as meteorites. Diamond can also be used as the active element in detectors for X-rays and gamma rays [34].

F. CONCLUSIONS

Diamond is the ultimate material of choice for many applications due to its extreme hardness, excellent mechanical properties, high thermal conductivity, transparency in the visible, UV, and X-ray regions, among many other exciting properties. The main drawback to diamond has been its relative scarcity and, therefore, high cost. As modern techniques for the growth of high quality films improve, potential applications, where diamond is the material of choice, become practical. The most important of these as far as the military is concerned may be substrates for thermal management. Diamond substrates will enable

computer chips to take advantage of three-dimensional architecture, leading to higher density chips and smaller computers. Clearly this application, despite the obvious military significance, will be driven by the civilian market. Other applications should be followed closely as, for example, diamond field emitters could be the enabling technology for flat panel displays. Diamond domes for covering infrared guided missiles appear to be a technology that will be too expensive for years to come, since no comparable civilian use exists to drive down the price. Other applications ranging from diamond microsensors to diamond semiconductor based electronics will likely become important in the 21st Century. However, as Jonathan Alford once said "technology always delivers less, arrives later, and costs more than forecast."

REFERENCES

1. A. Bakon and A. Szymanski, *Practical Uses of Diamonds*, New York: Ellis Horwood, 1993.
2. W. Hanusch and G. Bergmann, *Ind. Diam. Rev.*, 34, 126 (1974).
3. D.C. Harris, *Infrared Window and Dome Materials*, SPIE Optical Engineering Press, Volume 10, 1992.
4. P.D. Ownby and R. W. Stewart, *Engineering Properties of Diamond and Graphite*, in *Ceramics and Glasses*, Volume 4, Engineering Materials Handbook, ASM International, 1991.
5. M. Miyayama, K. Koumoto and H. Yanagida, *Engineering Properties of Single Oxides*, in *Ceramics and Glasses*, Volume 4, Engineer Materials Handbook, ASM International, 1991.
6. B.G. Streetman, *Solid-State Electronic Devices*, Prentice-Hall, 1980.
7. R.L. Cadenhead, Materials and Electronic Phenomena, in *Electronic Materials Handbook, Volume I - Packaging*, ASM International, 1989.
8. J.M. Sater, *Diamond Activities in DoD Report*, Institute for Defense Analyses, May, 1993.
9. G. Davies, T. Evans and P.F. James, *Proc. Roy. Soc. London A*, 277, 260 (1964).
10. G. Davies and T. Evans, *Proc. Roy. Soc. London A*, 318, 413 (1972).
11. J.E. Field, *The Properties of Diamonds*, Oxford: Academic Press, 1979.
12. *Physical Chemistry of Condensed Phase Superhard Materials and Their Surfaces*, Kiev: Naukova Dumka, 1975.
13. Y.L. Orlov, *The Mineralogy of Diamond*, London: Wiley.
14. *Diamond Deposition Newsletter*, 1, Superconductivity Publications, Somerset, N.J. (1990).
15. F.P. Bundy, H.T. Hall, H.M. Strong and H. Wentorf, *Nature*, 176, 51 (1955).
16. W.A. Yarbrough and R. Messier, *Science*, 24, 688 (1990).
17. R. Haubner and B. Lux, "Diamond Growth by Hot Filament Chemical Vapor Deposition: State of Art," in *Diamond and Related Materials II*, Amsterdam: Elsevier, 1993, p. 1277.
18. W. Piekarczyk and W.A. Yarbrough, in *New Diamond Science and Technology*, R. Messier, J.T. Glass, J.E. Butler, and R. Roy, eds., Pittsburgh: Materials Research Society, 1991, p. 327.

19. *Status and Applications of Diamond and Diamond-Like Materials: An Emerging Technology*, Report of the Committee on Superhard Materials, National Materials Advisory Board, NMAB-445, National Academy Press, 1990, 115 pages.
20. T.R. Anthony, "Comparative Evaluation of CVD Diamond Technologies," in *Workshop on Diamond and Diamond-Like Carbon Films for the Transportation Industry*, F.A. Nichols and D.K. Moores, eds., Argonne National Laboratory, February 4-5, 1992, p. 117.
21. M.H. Loh and M.A. Capelli, "Supersonic DC-Arcjet Synthesis of Diamond," in *Diamond and Related Materials II*, Amsterdam: Elsevier, 1993, p. 454.
22. K.J. Gray, "Electromagnetic Window Properties of CVD Diamond," in *Diamond Optics*, V.A. Feldman and S. Holly, eds., SPIE Vol. 1759, 1992, p. 203.
23. R.H. Wentorf, *J. Chem. Phys.*, **26**, 956 (1957) and **34**, 809 (1961).
24. I.P. Hayward and J.E. Field, "The Solid Particle Erosion of Diamond," *J. Hard Mater.*, **1**, 53 (1990).
25. R.A. Lucheta, "Thermally Induced Interfacial Stresses in a Thin Film on an Infinite Substrate," *Appl. Opt.*, **30**, 2252 (1991).
26. W.D. Partlow, R.E. Witkowski and J.P. McHugh, "CVD Diamond Coatings for the Infrared by Optical Brazing," in *Applications of Diamond Films and Related Materials*, Y. Tzeng, M. Yoshikawa, M. Murakawa and A. Feldman, eds., Amsterdam: Elsevier, 1991.
27. R.H. Hopkins, W.E. Kramer, G.E. Brandt, J.S. Schruben, R.A. Hoffman, K.B. Steinbruegge and T.L. Peterson, "Fabrication and Evaluation of Erosion-Resistant Multispectral Optical Windows," *J. Appl. Phys.*, **49**, 3133 (1978).
28. J. Wilkes and E. Wilkes, *Properties and Applications of Diamond*, Ch. 9, Oxford: Butterworth-Heinemann Ltd., 1991.
29. D.F. Grogan, T. Zhao, B.C. Bovard, and H.A. MacLeod, *Applied Optics*, **31**, 1483 (1992).
30. A.P. Maishe, G. Salamo, Z. Lu, M. Xiao, H.A. Naseem, W.D. Brown, and L.W. Schaper, *Proceedings of the 5th Annual Diamond Technology Workshop*, May 18-20, 1994.
31. J.L. Davidson, W.D. Brown, and J.P. Dismukes, "Diamond for Electronic Applications, A Unique Dielectric Semiconductor and Thermal Conductor," *The Electrochemical Society, Interface*, Fall, 1995, 22.
32. K. Derbyshire, Flat Panel Displays, "Beyond AMLCDs: Field Emission Displays," *Solid State Technology*, **37**, 55 (1994).
33. N. Kumar, H. Schmidt and C. Xie, "Flat Panel Displays, Diamond-Based Field Emission Flat Panel Displays," *Solid State Technology*, **38**, 71 (1995).
34. M.N. Yoder, *Proc. SPIE* 2151, 72, 1994.

35. R&D Magazine, 42, April, 1995.
36. R.W. Cahn, "High Speed Diamond," *Nature*, **377**, 197 (1995).
37. N. Fugimori, *New Diamond*, **3**, 20, 1987.
38. D. Wur, J.L. Davidson, W.P. King and D.L. Kinser, J. *Microelectronmechanical Systems*, **4**, No. 1, 1995.
39. Private communication with scientist at Courtauld.

Appendix A

GLOSSARY

Appendix A

GLOSSARY

μ	micro
$\mu\text{A/lm}$	microamperes per lumen
$^{\circ}\text{C}$	degrees centigrade
μm	micrometer
$\mu\text{m/hr}$	micrometers per hour
2-D	two dimensional
A	ampere
$\text{\AA}$	angstrom
Al_2O_3	corundum
ACS	attitude control system
AI2ATC	Advanced Image Intensifier Advanced Technology Demonstrator
Al	aluminum
Al_2O_3	alumina
ALMDS	Airborne Laser Mine Detection System
AlN	aluminum nitride
ALT	anode layer thruster
amu	atomic mass unit
APD	avalanche photodiode
APS	active pixel sensor
As	arsenic
BMDO	Ballistic Missile Defense Office
BN	boron nitride
Br	bromine
CCD	charge coupled device
Cd	candela
CIS	Combat Identification System
CLZ	craft landing zone
cm^3	cubic centimeters
CMOS	complementary metal-oxide semiconductor
COBRA	Coastal Battlefield Reconnaissance Analysis
Cr	chromium
CTE	charge transfer efficiency
CVD	chemical vapor deposition

DOE	Department of Energy
DOS	Disk Operating System
DRAM	dynamic random access memory
ECM	electronic countermeasure
EMP	electromagnetic pulse
EOD	explosive ordnance
EP	electric propulsion
eV	electronvolt
FAA	Federal Aviation Administration
FCC	Federal Communication Commission
Fe	iron
FIR	far infrared
FSK	frequency shift keying
FSU	former Soviet Union
g/cm^3	grams per cubic centimeter
Ga	gallium
Ga As	gallium arsenic
Gb/s	gigabits per second
Ge	germanium
GEO	geosynchronous earth orbit
GHz	gigahertz
GMA	gas-metal-arc
GMR	giant magnetoresistance
GN&C	guidance, navigation, and control
GPa	gigapascal
GPOW	Global Precision Optical Weapon
GPS	global positioning system
HEDM	high-energy-density material
Hz	hertz
IFF	identify friend or foe
in	inches
Isp	specific pulse
J/kgK	joules per kilogram kelvin
K	potassium
kb/s	kilobits per second
keV	kibelectronvolt
kHz	kilohertz

km	kilometer
km ²	square kilometer
kW	kilowatt
kWE	kilowatt (electric)
LCACV	landing craft air cushion vessel
LCD	liquid crystal display
LEO	low earth orbit
LIDAR	laser radar
lm	lumen
LO	low observability
lp	line pair
lp/mm	line pairs per millimeter
m	meters
mA/mm ²	milliampere per square millimeter
MAD	magnetic anomaly detection
MBE	molecular beam epitaxy
MCM	mine countermeasures
MCM	multichip module
MCP	microchannel plate
MEMS	Micro Electro-Mechanical Systems
meV	megaelectronvolts
MHz	megahertz
mJ	millijoule
mm	millimeter
mOe	millioersted
MOS	metal-oxide semiconductor
MOUT	Military Operations in Urban Terrain
MPa	megapascal
MPD	magnetoplasma-dynamic
ms	millisecond
MTI	moving target indicator
mW	milliwatt
N/m ²	newton per square meter
NASA	National Aeronautics and Space Administration
Nb	niobium
Nd	Neadymium
Ni	Nickel
nm	nanometer
ns	nanosecond

O	oxygen
ODS	Operation Desert Storm
Oe	oersted
p/m	parts per million
PG	photogate
PPT	pulsed plasma thruster
Pr	proseadymium
psi	pounds per square inch
R&D	research and development
RAM	radar-absorbing material
RAS	radar-absorbing structure
RCS	radar cross-section
RF	radio frequency
RLV	reusable launch vehicle
RTG	radioactive thermal generator
SA	situational awareness
SAM	surface-to-air missile
SAR	synthetic aperture radar
SDIO	Strategic Defense Initiatives Office
Se	selenium
Si	silicon
Si ₃ N ₄	silicon nitride
SiC	silicon carbide
SNR	signal-to-noise ratio
SPT	stationary plasma thruster
STS	Space Transportation System
SW	shallow water
SZ	surf zone
TAV	transatmospheric vehicle
Ti	titanium
TiO ₂	rutile
TOA	Total Obligational Authority
UAV	unmanned aerial vehicle
UEP	underwater electric potential
USAF	United States Air Force
UUV	unmanned underwater vehicle
UV	ultraviolet

V/cm	volts per centimeter
VLSI	very large scale integration
VSW	very shallow water
W	watt
W/mK	watts per meter kelvin
WWII	World War II
YAl ₂ O ₄	garnet
Zn	zinc
ZrC	zirconium carbide
ZrO ₂	cubic zirconia
ZrSiO ₄	zircon

Appendix B

DISTRIBUTION LIST FOR IDA PAPER P-3296

VOLUME I

APPENDIX B
DISTRIBUTION LIST FOR IDA PAPER P-3296

VOLUME I

GOVERNMENT AGENCY	<u>Number of copies</u>
Defense Advanced Research Projects Agency 3701 N. Fairfax Drive Arlington, VA 22203-1714 ATTN: Dr. Ira Skurnick, DSO	1
DSSG ALUMNI	
Professor A. Paul Alivisatos Department of Chemistry University of California Berkeley, CA 94720	1
Professor Gregory L. Baker Department of Chemistry Michigan State University 320 Chemistry Building East Lansing, MI 48824-1322	1
Professor Gaetano Borriello Department of Computer Science and Engineering, FR-35 University of Washington Seattle, WA 98195	1
Professor Kevin F. Brennan School of Electrical Engineering Georgia Institute of Technology Atlanta, GA 30332-0325	1

Professor Jean M. Carlson 1
Department of Physics
University of California
Santa Barbara, CA 93106

Professor S. Lance Cooper 1
Department of Physics
University of Illinois, Urbana-Champaign
1110 W. Green Street
Urbana, IL 61801

Professor John C. Doyle 1
Electrical Engineering 116-81
California Institute of Technology
Pasadena, CA 91125

Professor Brent T. Fultz 1
Mail Stop 138-78
Keck Laboratory of Engineering Materials
California Institute of Technology
Pasadena, CA 91125

Dr. Daniel E. Hastings 1
Chief Scientist of the Air Force
HQUSAF/ST
1075 AF Pentagon
Washington DC 20330-1075

Professor William C. Johnson 1
Department of Materials Science and Engineering
Thornton Hall
University of Virginia
Charlottesville, VA 22903

Professor Richard B. Kaner 1
Chemistry Department
University of California, Los Angeles
Los Angeles, CA 90024-1569

Professor Ann R. Karagozian 1
Department of Mechanical, Aerospace,
and Nuclear Engineering
University of California
Los Angeles, CA 90024-1597

Professor Stephen D. Kevan 1
Physics Department
University of Oregon
Eugene, OR 97403

Professor Christopher S. Kochanek 1
Department of Astronomy, MS-51
Harvard University
60 Garden Street
Cambridge, MA 02138

Professor Clifford R. Pollock 1
Department of Electrical Engineering
Cornell University
Ithaca, NY 14853-2801

Professor Gabriel Robins 1
Department of Computer Science
Thornton Hall
University of Virginia
Charlottesville, VA 22903

Professor Michael J. Shelley 1
Courant Institute of Mathematical Sciences
New York University
251 Mercer Street

OTHERS:

Institute for Defense Analyses 20
1801 N. Beauregard St.
Alexandria, VA 22311-1772

DTIC 2
Defense Technical Information Center
8725 John J. Kingman Rd., STE 0944
Fort Belvoir, VA 22060-6218

REPORT DOCUMENTATION PAGE			Form Approved OMB No. 0704-0188	
Public reporting burden for this collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington, VA 22202-4302, and to the Office of Management and Budget, Paperwork Reduction Project (0704-0188), Washington, DC 20503.				
1. AGENCY USE ONLY (Leave blank)	2. REPORT DATE February 1996	3. REPORT TYPE AND DATES COVERED Final		
4. TITLE AND SUBTITLE Defense Science Study Group IV: Study Reports 1994-1995 Volume I		5. FUNDING NUMBERS DASW01-94-C-0054 Task No. A-103		
6. AUTHOR(S) William J. Hurley, Nancy P. Licato, DSSG IV members				
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Institute for Defense Analyses 1801 N. Beauregard Street Alexandria, VA 22311-1772		8. PERFORMING ORGANIZATION REPORT NUMBER IDA Paper P-3296		
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES) Defense Advanced Research Projects Agency 3701 N. Fairfax Drive Arlington, VA 22203-1714 ATTN: Dr. Ira Skurnick, DSO		10. SPONSORING/MONITORING AGENCY REPORT NUMBER FFRDC Programs 2001 N. Beauregard Street Alexandria, VA 22311-1772		
11. SUPPLEMENTARY NOTES				
12a. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution unlimited.			12b. DISTRIBUTION CODE	
13. ABSTRACT (Maximum 200 words) The Defense Science Study Group (DSSG) is a program sponsored by the Defense Advanced Research Projects Agency (DARPA) and managed by the Institute for Defense Analyses (IDA). The purpose of the program is to expose some of the most talented young scientists and engineers from academia to issues and operations related to national security. Strengthening ties between the scientists and engineers and the national security community will provide the Government with a new source of technical advisors and informed critics. Individuals spend about 20 days per year for 2 years receiving briefings by distinguished speakers; visiting laboratories, intelligence organizations, and military, manufacturing, and industrial facilities; and conducting studies on defense-related topics. This report is a compilation of the studies prepared by the members.				
14. SUBJECT TERMS Defense Science Study Group, defense technology, image intensification, space technology, radio tags, titanium submarines, acoustic arrays, shipboard waste, stealth materials, ship design, IFF, complex systems, cavitation, virtual engineering, diamond			15. NUMBER OF PAGES 297	
			16. PRICE CODE	
17. SECURITY CLASSIFICATION OF REPORT UNCLASSIFIED	18. SECURITY CLASSIFICATION OF THIS PAGE UNCLASSIFIED	19. SECURITY CLASSIFICATION OF ABSTRACT UNCLASSIFIED	20. LIMITATION OF ABSTRACT Unlimited	