

LOAN DOCUMENT

PHOTOGRAPH THIS SHEET

AD-A262 018



DTIC ACCESSION NUMBER

LEVEL

INVENTORY

AFDSR-TR-93-0120

DOCUMENT IDENTIFICATION

Dec 92

DISTRIBUTION STATEMENT

Approved for public release
Distribution Unlimited

DISTRIBUTION STATEMENT

ACCESSION FOR

NTIS GRA&I
DTIC TRAC
UNANNOUNCED
JUSTIFICATION

☒
☒
☒

BY

DISTRIBUTION/

AVAILABILITY CODES

DISTRIBUTION

AVAILABILITY AND/OR SPECIAL

A-1

DISTRIBUTION STAMP

DTIC
ELECTE
MAR 10 1993
S C D

DATE ACCESSIONED

DATE RETURNED

98 3 4 068
98 3 10 015

DATE RECEIVED IN DTIC

93-04696



REGISTERED OR CERTIFIED NUMBER

PHOTOGRAPH THIS SHEET AND RETURN TO DTIC-FDAC

H
A
N
D
L
E

W
I
T
H

C
A
R
E

REPORT DOCUMENTATION PAGE

AGENCY USE ONLY (Leave blank) 28 Dec 92 Annual 1 Sep 91 - 31 Aug 92

REPORT AND SUBTITLE

1992 Summer Faculty Research Program (SFRP)
Volumes 1 - 16

V-16

F49620-90-C-0076

Mr Gary Moore

RESEARCH ORGANIZATION NAME(S) AND ADDRESS(ES)

Research & Development Laboratories (RDL)
5800 Uplander Way
Culver City CA 90230-6600

MONITORING AGENCY NAME(S) AND ADDRESS(ES)

AFOSR/NI
110 Duncan Ave., Suite B115
Bldg 410
Bolling AFB DC 20332-0001
Lt Col Claude Cavender

SUPPLEMENTARY NOTES

ABSTRACT AVAILABILITY STATEMENT

UNLIMITED

ABSTRACT (Maximum 200 words)

The purpose of this program is to develop the basis for continuing research of interest to the Air Force at the institution of the faculty member; to stimulate continuing relations among faculty members and professional peers in the Air Force to enhance the research interests and capabilities of scientific and engineering educators; and to provide follow-on funding for research of particular promise that was started at an Air Force laboratory under the Summer Faculty Research Program.

During the summer of 1992 185 university faculty conducted research at Air Force laboratories for a period of 10 weeks. Each participant provided a report of their research, and these reports are consolidated into this annual report.

1. AUTHOR

2. SECURITY CLASSIFICATION

UNCLASSIFIED

3. SECURITY CLASSIFICATION OF THIS PAGE

UNCLASSIFIED

4. SECURITY CLASSIFICATION OF ABSTRACT

UNCLASSIFIED

UL

5. DISTRIBUTION STATEMENT

UNITED STATES AIR FORCE
SUMMER RESEARCH PROGRAM -- 1992
GRADUATE STUDENT RESEARCH PROGRAM (GSRP) REPORTS

VOLUME 10
WRIGHT LABORATORY

RESEARCH & DEVELOPMENT LABORATORIES
5800 Uplander Way
Culver City, CA 90230-6608

Program Director, RDL
Gary Moore

Program Manager, AFOSR
Lt. Col. Claude Cavender

Program Manager, RDL
Billy Kelley

Program Administrator, RDL
Gwendolyn Smith

Submitted to:

AIR FORCE OFFICE OF SCIENTIFIC RESEARCH
Bolling Air Force Base
Washington, D.C.
December 1992

PREFACE

This volume is part of a 16-volume set that summarizes the research accomplishments of faculty, graduate student, and high school participants in the 1992 Air Force Office of Scientific Research (AFOSR) Summer Research Program. The current volume, Volume 10 of 16, presents the final research reports of graduate student (GSRP) participants at Wright Laboratory.

Reports presented herein are arranged alphabetically by author and are numbered consecutively -- e.g., 1-1, 1-2, 1-3; 2-1, 2-2, 2-3.

Research reports in the 16-volume set are organized as follows:

VOLUME	TITLE
1	Program Management Report
2	Summer Faculty Research Program Reports: Armstrong Laboratory
3	Summer Faculty Research Program Reports: Phillips Laboratory
4	Summer Faculty Research Program Reports: Rome Laboratory
5A	Summer Faculty Research Program Reports: Wright Laboratory (part one)
5B	Summer Faculty Research Program Reports: Wright Laboratory (part two)
6	Summer Faculty Research Program Reports: Arnold Engineering Development Center; Civil Engineering Laboratory; Frank J. Seiler Research Laboratory; Wilford Hall Medical Center
7	Graduate Student Research Program Reports: Armstrong Laboratory
8	Graduate Student Research Program Reports: Phillips Laboratory
9	Graduate Student Research Program Reports: Rome Laboratory
10	Graduate Student Research Program Reports: Wright Laboratory
11	Graduate Student Research Program Reports: Arnold Engineering Development Center; Civil Engineering Laboratory; Frank J. Seiler Research Laboratory; Wilford Hall Medical Center
12	High School Apprenticeship Program Reports: Armstrong Laboratory
13	High School Apprenticeship Program Reports: Phillips Laboratory
14	High School Apprenticeship Program Reports: Rome Laboratory
15	High School Apprenticeship Program Reports: Wright Laboratory
16	High School Apprenticeship Program Reports: Arnold Engineering Development Center; Civil Engineering Laboratory

1992 GRADUATE RESEARCH REPORTS

Wright Laboratory

<u>Report Number</u>	<u>Report Title</u>	<u>Author</u>
1	Point Spread Function Characterization of a Scophony Infrared Scene Projector	Terri L. Alexander
2	The Effect of Nonhomogeneous Interphases and Global/Local Volume Fraction on the Mechanics of a Layered Composite	Vernon T. Bechel
3	Velocity and Temperature Measurements in a High Swirl Dump Combustor	Lance H. Benedict
4	Development of an Enhanced Post Run Data Analysis Program for the Integrated Electromagnetic System Simulator	Benjamin F. Bohren
5	Hard Target Code Assessment and a Qualitative Study of Slide Line Effects in EPIC Hydrocode	Thomas C. Byron
6	Laser Imaging and Ranging (LIMAR) Processing	Ahmet A. Coker
7	Neural On-Line Learning in Missile Guidance	Jeffrey S. Dalton
8	Optimal Detection of Targets in Clutter Using an Ultra-Wideband Fully-Polarimetric SAR	Ronald L. Dilsavor
9	Effects on Intensity Thresholding on the Power Spectrum of Laser Speckle	Alfred D. Ducharme
10	Using X Windows to Display Experimental Data	David E. Frink
11	VLSI Synthesis Guiding Techniques using the SOAR Artificial Intelligence Architecture	Lindy Fung
12	A Chemical Investigation of the Oxidative Behavior of Aviation Fuels	Ann Phillips Gillman
13	Finite Element Analysis of Interlaminar Tensile Test Specimens	Diane Hageman
14	The Design of a NO ₂ Chemiluminescence Test Chamber	Andrew P. Johnston
15	Jacobian Update Strategies for Quadratic and Near-Quadratic Convergence of Newton and Newton-Like Implicit Schemes	Daniel B. Kim
16	Impact of Quasi-Isotropic Composite Plates by 1/2" Steel Spheres	John T. Lair
17	Universal Controller Analysis and Implementation	Shawn H. Mahloch
18	Process Migration Facility for the QUEST Distributed VHDL Simulator	Dallas J. Marks
19	(Report not received)	
20	Preliminary Work on the Design of an Image Algebra Coprocessor	Trevor E. Meyer

Wright Laboratory (cont'd)

<u>Report Number</u>	<u>Report Title</u>	<u>Author</u>
21	(Report not received)	
22	Detection and Adaptive Frequency Estimation for Digital Microwave Receivers	Steven Nunes
23	Efficient Analysis of Passive Microstrip Elements for MMICs	Todd W. Nuteson
24	Built-in Self-Test Design of Pixel Chip	R. Frank O'Bleness
25	A Study of Damage in Graphite Epoxy Panels Subjected to Low and High Velocity Impact	Mohammed A. Samad
26	Low Velocity Impact Damage of Composite Materials	Rob Slater
27	Monitoring of Damage Accumulation for the Prediction of Fatigue Lifetime of Cord-Rubber Composites	Jeffrey A. Smith
28	Effects of Intermolecular Interactions in a Cyclic Siloxane Based Liquid Crystal	Edward Peter Socci
29	(Report not received)	
30	A Literature Review of the Reaction Kinetics of the Pyrolysis of Phenolic/Graphite Composite Material	Kimberly A. Trick
31	A Study of the Chemical Vapor Deposition of a Single Filament	Rose Marie Vecchione
32	A Switched Reluctance Motor Drive using MOSFETS, HCTL-1100, and MC6802 Microprocessor	Lawrence Vo
33	Investigation of the Combustion Characteristics of Swirled Injections in a Confined Coannular System with a Sudden Expansion	David L. Warren

POINT SPREAD FUNCTION CHARACTERIZATION
OF A
SCOPHONY INFRARED SCENE PROJECTOR

Terri L. Alexander
Graduate Research Associate
Center for Research in Electro-Optics & Lasers (CREOL)

University of Central Florida
Orlando, Florida 32816

Final Report for:
AFOSR Summer Research Program
Guided Interceptor Technology Branch
Wright Laboratories Armament Directorate

Sponsored by:
Air Force Office of Scientific Research
Eglin Air Force Base, Florida 32542

August 1992

POINT SPREAD FUNCTION CHARACTERIZATION
OF A
SCOPHONY INFRARED SCENE PROJECTOR

Terri L. Alexander
Graduate Research Associate
Center for Research in Electro-Optics & Lasers (CREOL)
University of Central Florida

Abstract

A Scophony Infrared Scene Projector (IRSP) is being used at Wright Laboratories Armament Directorate, Guided Interceptor Technology Branch, Eglin AFB, for evaluation of thermal-imaging guidance systems. This is a hardware-in-the-loop testing system which reduces the number of necessary field trials and has unlimited potential for in-laboratory simulation where the performance of entire seeker systems can be analyzed. The performance of an optical system in terms of such characteristics as wavefront error, resolution, and transfer factor, can be measured with knowledge of the system's MTF and PSF performance. A slow scan calibration system was used to measure the image plane of the IRSP under three separate configurations of the system. MTFs and PSFs were derived for the IRSP without the use of the scatter screen, with the scatter screen in place, and with the scatter screen rotating.

POINT SPREAD FUNCTION CHARACTERIZATION OF A SCOPHONY INFRARED SCENE PROJECTOR

Terri L. Alexander

I INTRODUCTION

The Kinetic Kill Vehicle Hardware-in-the-Loop Simulation (KHILS) Test Facility is being developed by the Wright Laboratories Armament Directorate, Strategic Defense Division, Guided Interceptor Branch (WL/MNSI), Eglin AFB, FL, to provide non-destructive hardware-in-the-loop performance testing of strategic defense interceptor systems. The main focus of the KHILS system is in performance analysis of seeker systems, signal processing, and guidance, navigation, and control subsystems.¹

The benefits of hardware-in-the-loop testing include saved development time and expense, and the unlimited potential of in-laboratory simulation and testing where the performance of entire seeker systems can be tested and analyzed without the need for as many field trials. A major component of the KHILS system is the Laser Scophony Infrared Scene Projector (IRSP). A critical element of the IRSP is the scatter screen which is designed to eliminate laser coherence effects and to redefine the optical invariant of the projector system to match the unit under test. The infrared scene projector's optical performance and the effects of the scatter screen were investigated for this report.

Fundamental figures of merit for an optical system are its modulation transfer, optical transfer, and point spread functions. This report presents measured data of the IRSP pixel intensity profiles with and without the scatter screen. That data is then used to determine modulation transfer function and point spread function performance.

II THEORY

i) Laser Scophony IRSP

The KHILS IRSP is a scanned laser projection system that employs Scophony techniques with acousto-optic modulation/deflection devices to project high resolution 96x96 pixel imagery in up to four infrared wavebands simultaneously. Figure 1 is a diagram of the IRSP optical layout.

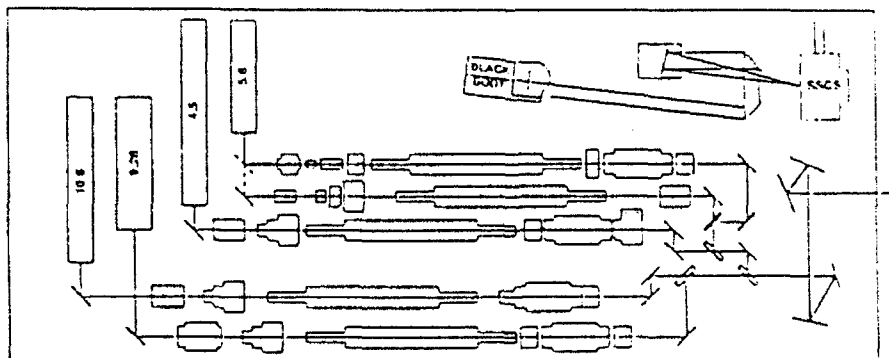


Figure 1. Infrared Scene Projector optical layout.¹

Scophony modulation uses a collimated laser beam to fill a large portion of the acousto-optic modulator cell. Spreading the laser input allows the projection of multiple pixels simultaneously. This method is used to increase the dwell time of the IRSP on respective seeker/focal plane array detectors, which also increases the spatial resolution.^{2,3}

The four laser optical trains can be used in any combination for single or multiple wavelength testing. During multi-wavelength operation, the image scan of all optics trains is synchronized, and the outputs are optically combined. This report focuses on the CO₂-laser-driven 9.28μm optical train.

A functional block diagram of the optical system and the Scophony image scan pattern for a 96x96 pixel format are shown in Figure 2.

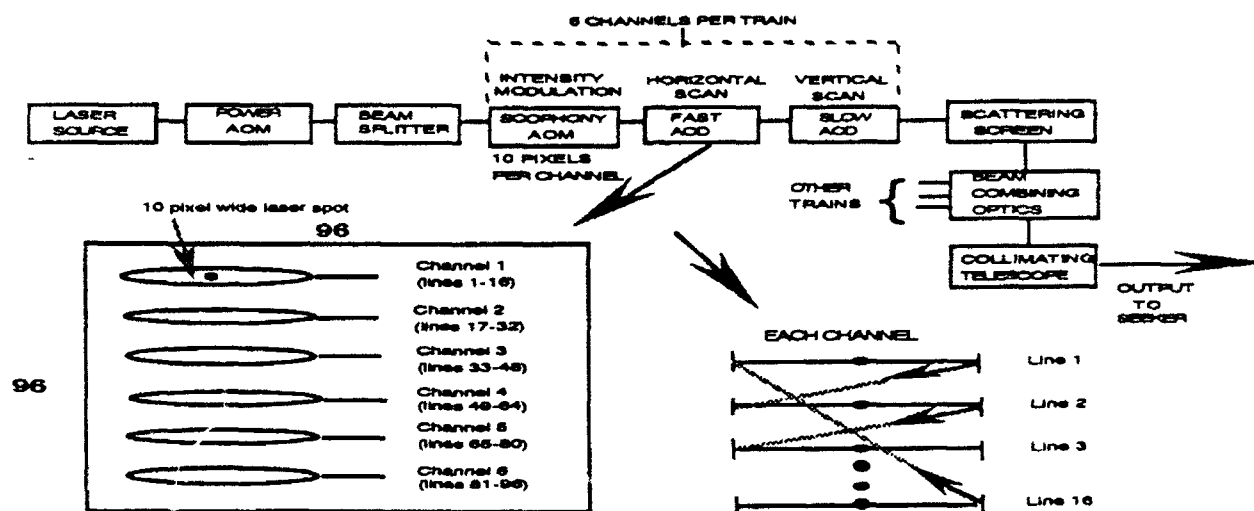


Figure 2. IRSP functional block diagram and Scophony scan pattern.¹

Active elements in the system include the laser source and the acousto-optic modulator and deflectors (AOMs, AODs). The power AOM is used to attenuate the laser output and control the maximum intensity within the overall image. The beam splitter divides the laser beam into six equal intensity segments with a 16-line vertical spacing between each segment. The complete 96x96 pixel image is formed using six channels with each channel consisting of 16 lines of 96 pixels each. Within each channel, the Scophony AOM produces a 10-pixel wide intensity modulated laser spot which is scanned over the 96x16 portion of the image using the fast and slow AODs in the pattern shown in the lower portion of Figure 2. All six channels are scanned simultaneously, so that, at any instant of time 60 pixels in the 96x96 image are being illuminated.¹

ii) Scatter Screens

The laser sources used in the IRSP produce monochromatic light with a high degree of spatial coherence. Although lasers provide an excellent source of high intensity light, their coherence introduces an interference phenomena known as laser speckle which causes a non-uniform intensity pattern. Additionally, as with any optical system, a fundamental characteristic of the IRSP is its Optical Invariant (*Lagrange Invariant*). The optical invariant states that across any surface for a given optical system

$$y_p N u - y N u_p = \text{a constant}$$

where y_p is the chief ray height, u_p is the chief ray angle with respect to the optical axis, N is the index of refraction, y is the axial ray height, and u is the axial ray angle with respect to the optical axis. A result of this theorem is that if the aperture of a system can be varied, then the angular field must change in inverse ratio to the aperture. The product of aperture and field is constant and increasing one must result in reduction of the other.⁵

The scatter screen is a Zinc Selenide (ZnSe) circular plate (3mm thick, 1.5 inch diameter) with an rms roughness of $0.75\mu\text{m}$ and antireflection coating designed uniquely for use with each wavelength of operation. The purpose of the scatter screen is to modify the optical invariant established at the scophony acousto-optic modulator allowing for collimating optics to match the invariant of the seeker. The rms roughness which establishes optimum non-lambertian output without introducing unacceptable attenuation is currently being investigated. In addition to modifying the optical invariant, rotation of the scatter screen causes the beam to encounter different scattering sites and averages out the interference effects. The result is loss of coherence and elimination of laser speckle.⁶

iii) Point Spread Function, Optical Transfer Function and Modulation Transfer Function

Discussion of the point spread function (PSF), optical transfer function (OTF), modulation transfer function (MTF), and their relationships to each other is in order. A graphical representation of their relationships is shown in Figure 3.

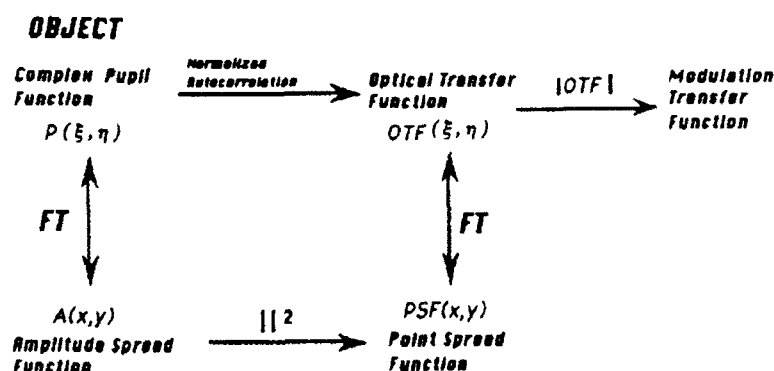


Figure 3. Relationships between different imaging properties of an optical system.⁷

We begin with a complex pupil function $P(\xi, \eta)$, which describes the wavefront shape as it emerges from the IRSP (where ξ and η are the spatial frequencies in the X and Y directions). The Fourier transform of $P(\xi, \eta)$ is the amplitude spread function $A(x, y)$ which is a field amplitude and phase. The squared modulus of $A(x, y)$ is the point spread function $PSF(x, y)$ which is a profile of the resulting irradiance distribution in the image plane. The Fourier transform of the $PSF(x, y)$ is the optical transfer function $OTF(\xi, \eta)$. The OTF is a measure of an optical system's ability to form high contrast images. The modulation transfer function $MTF(\xi, \eta)$ is the modulus of the OTF and is a measure of the reduction in contrast from object to image. If the modulation of a periodic, one-dimensional irradiance (I) distribution is defined as

$$\text{modulation depth} = \{I(\text{max}) - I(\text{min})\} / \{I(\text{max}) + I(\text{min})\}$$

then for a sinusoidal distribution with some spatial frequency ξ , the modulation transfer is the

decrease in modulation depth from the object plane to the image plane

$$\text{modulation transfer} = \text{image modulation} / \text{object modulation}.$$

Plotting the modulation transfer versus spatial frequency is the modulation transfer function $\text{MTF}(\xi)$.⁴

iv) Measurement Methodology

The IRSP output image is projected directly to a seeker under test or folded through a series of mirrors to the Slow Scan Calibration System (SSCS) as shown in Figure 1. The SSCS has a collimated blackbody source for radiometric and spatial reference. The SSCS uses an off-axis-parabola (OAP) to focus the IRSP output to an image plane which is scanned by a single element Mercury-cadmium-telluride (HgCdTe) photovoltaic detector. The resulting irradiance distribution is then used to determine the spatial characteristics of the IRSP image. This arrangement is shown in Figure 4.

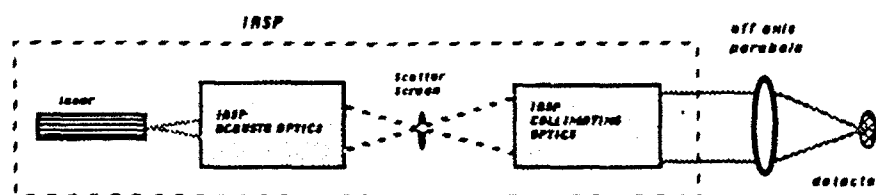


Figure 4. The IRSP output is passed through an off-axis-parabola and focused down the detector.

The irradiance distribution data measured by the SSCS is the resultant point spread function of the IRSP and the SSCS optical systems. This aggregate PSF can be described by the following equation:

$$\text{PSF}_{\text{measured}} = \text{PSF}_{\text{IRSP}} * \text{PSF}_{\text{OAP}} * w(x) \quad (1)$$

where $*$ denotes the convolution operator and the detector has an along-scan width $w(x)$. The detector contribution to $\text{PSF}_{\text{measured}}$ is determined by its dimensions. This leaves two unknowns, PSF_{IRSP} and PSF_{OAP} . The contribution of PSF_{OAP} must now be determined.

Measurement of PSF_{OAP} requires consideration of the consequences of source size, detector size, and collimator quality on the measured PSF. This will be done exclusive of the IRSP. The arrangement for determining PSF_{OAP} is shown in Figure 5.

The scanning detector produces an output voltage $v(x)$ as a function of position. This output voltage is a function of detector width along the scan direction $w(x)$. If the irradiance distribution in the image plane is denoted by $i(x)$ (Watt/cm^2), then the detector output is

$$v(x) = i(x) * w(x). \quad (2)$$

The irradiance distribution in the image plane $i(x)$, is the convolution of the ideal image with the PSF produced by the collimator/OAP system. The ideal image here is $p(x/dM)$, where $p(x)$ is the pinhole function, d is the diameter of the pinhole, and M is the magnification of the collimator/OAP

system. Now we have

$$i(x) = p(x/dM) * PSF_{coll\&OAP}(x). \quad (3)$$

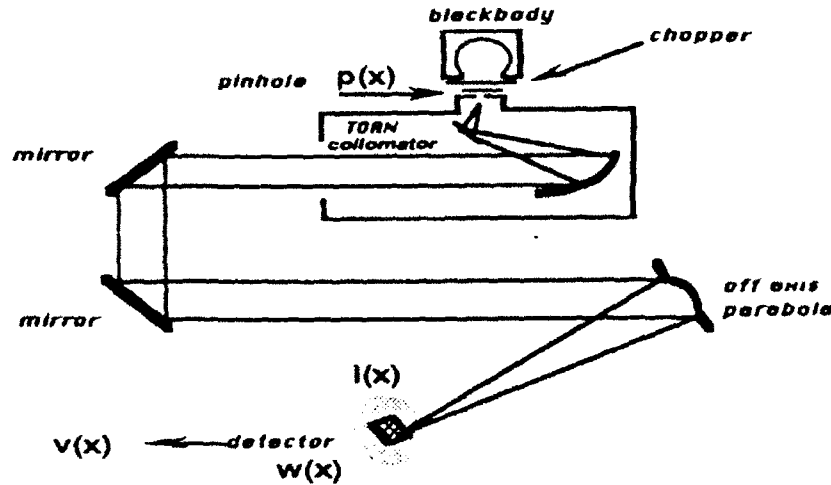


Figure 5. Experimental arrangement for the characterization of the off-axis-parabola.

The point spread function of the collimator and OAP system from equation (3) is the convolution of the PSFs caused by aberrations in the collimator and the OAP, as well as the PSF caused by diffraction in the collimator/OAP system. This is

$$PSF_{coll\&OAP}(x) = PSF_{aberr\ coll}(x) * PSF_{aberr\ OAP}(x) * PSF_{diffraction, coll/OAP}(x).$$

Assuming the collimator is diffraction-limited (this is experimentally verified later), we let $PSF_{aberr\ coll} = \delta(x)$. Using the properties of the convolution with a delta function we now have

$$PSF_{coll\&OAP}(x) = PSF_{aberr\ OAP}(x) * PSF_{diffraction, coll/OAP}(x).$$

Diffraction PSF is calculated once for a whole system and is determined by the limiting aperture for the overall optics train. In this case the $PSF_{diffraction, coll/OAP}(x)$ is determined by the aperture stop of the collimator and OAP system. Since the OAP is the aperture stop of the system it will determine $PSF_{diffraction}$. This also means that the OAP will be operating at a larger F-number ($F/\#$) than the collimator since the OAP is overfilled by the collimator. The $F/\#$ is a way to specify the amount of aperture used in an optical system. However, while this experiment would have the OAP operating at a certain relative aperture determined by the collimator beam overfilling the aperture, once we direct the IRSP into the OAP for characterization of the IRSP, this beam will not overfill the aperture and will cause the OAP to operate at a different $F/\#$. To correct for this, we must know the beam size generated by the IRSP and apply an aperture stop of equal size directly to the OAP during characterization of the OAP (this is shown in the experimental procedure section). Therefore, the $PSF_{diffraction\ coll/OAP}$ is $PSF_{diffraction\ OAP}$. Now

$$PSF_{coll\&OAP}(x) = PSF_{aberr\ OAP}(x) * PSF_{diffraction\ OAP}(x) = PSF_{OAP}(x).$$

Returning to equations (2) and (3) and making appropriate substitutions we have an expression for the measured data:

$$v(x) = w(x) * p(x/dM) * PSF_{OAP}(x) \quad (4)$$

In a diffraction limited system, the point spread function would correspond in shape to the diffraction pattern produced by a point source. For a small pinhole and a small detector, this would be the case and $v(x)$ would equal PSF_{OAP} . But for a finite detector and a finite pinhole the affects of this convolution must be considered.⁴

Recall that a convolution in the spatial domain is a multiplication in the Fourier domain.⁸ If we take the Fourier transform of equation (4) and divide out the detector and pinhole affects we have the MTF of the OAP.

$$V(\xi) = W(\xi) \times P(dM\xi) \times MTF_{OAP}(\xi)$$

and

$$MTF_{OAP}(\xi) = V(\xi) / \{W(\xi) \times P(dM\xi)\}. \quad (5)$$

Calculation of the PSF of the OAP is performed by inverse Fourier transforming the MTF_{OAP} .

With this accomplished, we direct the IRSP to the detector by way of the OAP as in Figure 4 and recalling equation (1)

$$PSF_{measured} = PSF_{IRSP} * PSF_{OAP} * w(x) \quad (1)$$

and in the Fourier domain

$$MTF_{measured} = MTF_{IRSP} \times MTF_{OAP} \times W(\xi)$$

$$MTF_{IRSP} = MTF_{measured} / \{MTF_{OAP} \times W(\xi)\}$$

$$\mathfrak{F}^{-1}\{MTF_{IRSP}\} = PSF_{IRSP} \quad (6)$$

where $\mathfrak{F}^{-1}$ denotes the inverse Fourier transform.

III EXPERIMENTAL PROCEDURE

The objective of this experiment is to measure baseline performance of the infrared scene projector described above for three cases. The three cases are (1) without the scatter screen, (2) with the scatter screen in place, and (3) with the scatter screen in place and rotating. The measured data will then be used to determine point spread, and modulation transfer functions for all three cases.

i) Verify TOAN collimation.

A Telescopic Off Axis Newtonian (TOAN) Collimator is used during measurements to determine PSF_{OAP} as shown in Figure 5. It is necessary to confirm that the output beam is collimated.

Apparatus: SORL TOAN Collimator (focal length 30.059", F/# 5.01)
 Blackbody (temperature 800°K)
 Pinhole (1.2mm = 0.04724" diameter)
 Theodolite
 Penta prism

The effective field of view of the system is determined by the diameter of the pinhole aperture placed in the focal plane, and is calculated assuming the small angle approximation

$$\tan \Theta = \Theta = \phi_{\text{pin}}/f = \phi_{\text{pin}}/\{F/\# D\} \quad (7)$$

where ϕ is the diameter of the pinhole, f is the focal length of the primary mirror, and D (6") is the clear aperture of the system. The result is a 0.09° effective field of view and a 0.045° expected deviation from collimation.

Procedure: Alignment data was taken by means of a theodolite and translation of penta prism across the collimator output beam at 4 positions with one inch spacings.

Results: A vertical deviation of 0.025° and a horizontal deviation of 0.0072° confirms collimator quality.

ii) Determination of required aperture size.

To ensure that the OAP operates at the same F/# during its characterization as it does during IRSP data collection, IRSP collimated beam size must be determined.

A single pixel was used to generate an IRSP collimated beam output. The location and size of the pupil was determined with the use of a pyroelectric vidicon. The size of the pupil was measured to be 1.8 inches in diameter at 97.7 inches from the primary mirror. A pliable plastic material was fitted to the OAP aperture with a centered opening 1.8 inches in diameter.

iii) Data collection for PSF_{OAP} calculations.

Apparatus: SORL TOAN Collimator (focal length 30.059", F/# 5.01)
 SORL Off Axis Parabola (focal length 49.878", F/# 8.33)
 Blackbody (temperature 800°K)
 Pinhole (256μm = 0.01007874" diameter)
 Detector size 250μm

Procedure: The detector is placed in the focal plane of the OAP. We would like to have a diffraction limited collimator. With a blackbody temperature of 800°K we can assume a worse case resolution at 10μm and calculate the largest pinhole size to be considered a point source.

$$2.44\lambda F/\# = 2.44(10\mu\text{m})(5.01) = 122.2\mu\text{m} \quad (8)$$

Due to limited time and hardware availability we manufactured our own pinhole and had to settle for a 256μm diameter. The diameter of the pinhole was measured with an electron microscope which furnished photographs of the aperture and allowed for the identification of a flaw in the smoothness of the drilled hole. Now instead of the diffraction limited Airy disk at the detector, we have the image of the pinhole. This arrangement is equivalent to Figure 6.

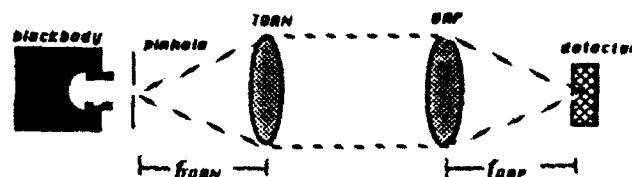


Figure 6. Equivalent magnification system for the OAP characterization.

Expected spot size at detector:

$$\begin{aligned} \text{diameter}_{\text{image}} &= \text{diameter}_{\text{pinhole}} \times (f_{\text{OAP}}/f_{\text{TOAN}}) \\ &= 424.8\mu\text{m} \end{aligned} \quad (9)$$

Results: On the first trial, the center of the spot was located using a Slow Scan Calibration System (SSCS) which moves the detector stage in vertical and horizontal directions taking any desired number of steps at $0.5\mu\text{m}$ per step. Each data point was taken by moving the stage and using an oscilloscope to measure the detector response at each position. The scanning process is a convolution of the detector with the image. The data collected will be larger than the actual image by the width of the detector. The results of a scan in the X and Y directions are shown in Figure 7.

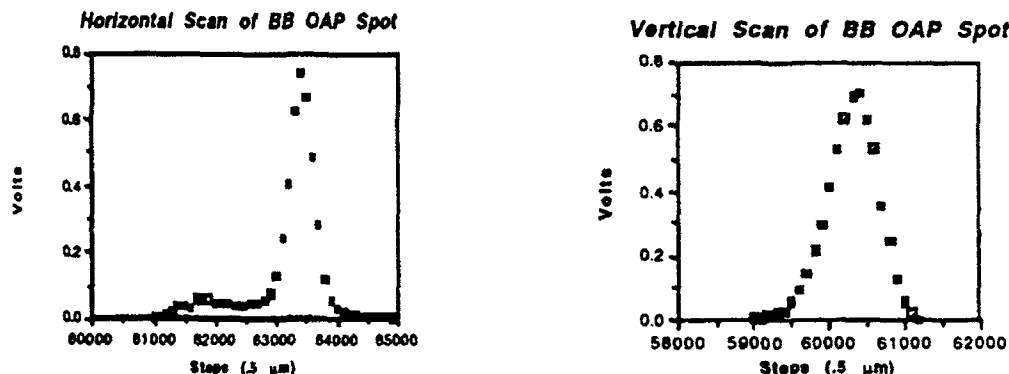


Figure 7. Original scan in the X and Y direction of image generated by the set up described by Figure 6.

In the horizontal direction 3700 steps taken at $0.5\mu\text{m}$ each means a $1850\mu\text{m}$ scan was taken. Subtracting the detector overlap of $250\mu\text{m}$ leaves an image size of $1600\mu\text{m}$. Horizontally we also notice an unexpected sidelobe, an anomaly with many possible causes. If we disregard the sidelobe which is $750\mu\text{m}$, the resulting image size is $850\mu\text{m}$ which is exactly twice as big as expected. Vertically the image appears symmetric and Gaussian, however after allowing for the detector overlap we still have an $850\mu\text{m}$ image, also twice the size predicted by equation (9).

Possible sources of error considered were coma, the flaw in the pinhole, internal detector reflections, and alignment. The possibility of coma was ruled out with the use of

$$\text{coma} = h' / 16(F/\#)^2 \quad (10)$$

where h' is the height of the image above the axis and $F/\#$ is the F-number of the system.⁹ Application of the above equation showed coma due to possible misalignment of one inch would be no greater than $23\mu\text{m}$. To rule out pinhole flaws it was rotated 90° with no improvement in results. Detector tilt was ruled out after an adjustment and more data collection showed no change.

Complete realignment and focusing solved the problem with nearly perfect results:

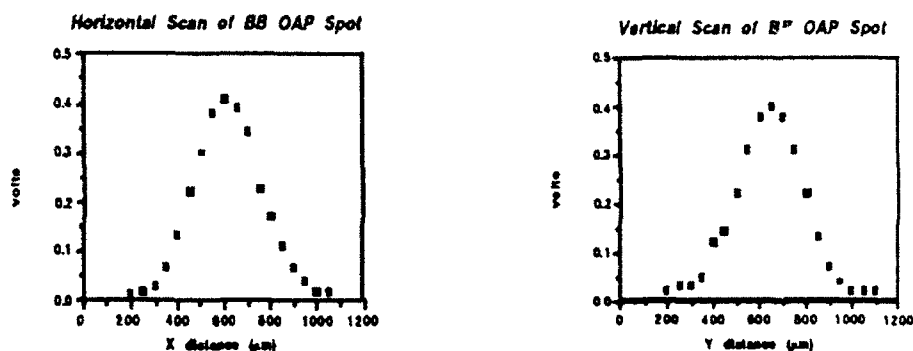


Figure 8. Image irradiance distribution after realignment of SSCS optical system.

Vertical and horizontal profiles are nearly perfect Gaussians and the image size in each case is $\approx 400\mu\text{m}$ after subtracting detector overlap as measured from the $1/e^2$ point.

This data and all that follows were collected using a scanning program written for the SSCS described above which now uses a digital oscilloscope and scans the entire image plane and displays the measured voltages on a CRT. Now instead of an X and Y slice of the irradiance distribution, a complete two-dimensional image sample is obtained.

Now with confidence in the OAP characterization data we may proceed with the IRSP characterization.

iv) Data collection for calculation of PSF_{IRSP}

Using the arrangement in Figure 4, and the scanning procedure described above, measurements were taken over the entire image plane of one pixel for each of the three desired configurations: without the scatter screen, with the scatter screen, and with the scatter screen rotating. An interactive data language (Precision Visuals: PV WAVE) was used at a VAX workstation to generate 3-dimensional plots of the measured data, MTFs, PSFs, and contour plots of the PSFs for each of the three cases. Hand calculations were performed for a horizontal dimensional analysis of the MTF and PSF of the IRSP without the scatter screen to verify computer results with excellent agreement and are included in the results section for comparison.

IV RESULTS

i) IRSP without the scatter screen.

The irradiance distribution of the image plane without the scatter screen is shown in Figure 9. The projected image is of one pixel and a Gaussian distribution is expected. Without the scatter screen the output is circular in the X-Y plane but does not appear Gaussian from the perspective of Figure 9. Although the fluctuations in the peak values rule out any saturation possibilities, it does appear to have an on/off quality. Possibly this is caused by a defocussing effect due to the removal of the scatter and the resulting shorter optical path length. The missing scatter screen could also be responsible for a mismatch in F/#s occurring between the acousto-optics and the

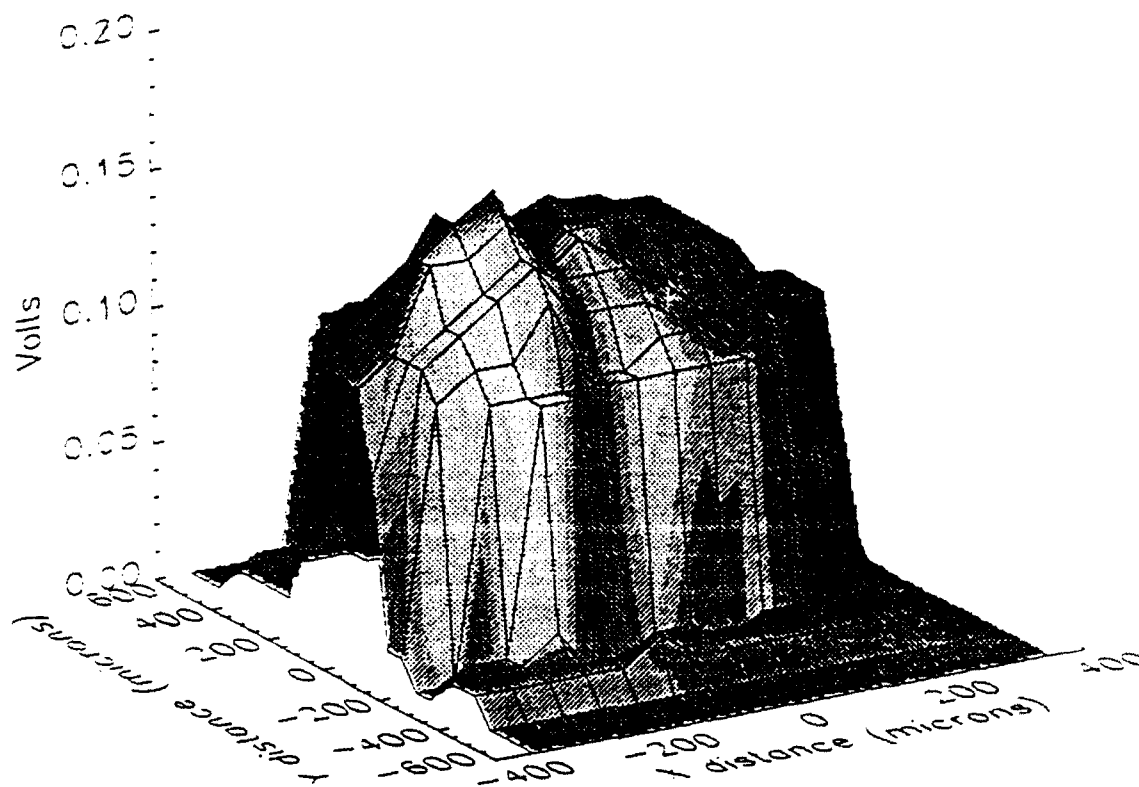


Figure 9. Image plane irradiance distribution of the IRSP without the scatter screen.

and the collimating optics. Further investigation shows the unusual appearance of this plot to be somewhat deceiving as shown in Figure 10. To confirm the validity of the data and obtain a more representative profile, the data were summed in the horizontal and vertical directions and plotted versus X and Y to verify a Gaussian distribution.

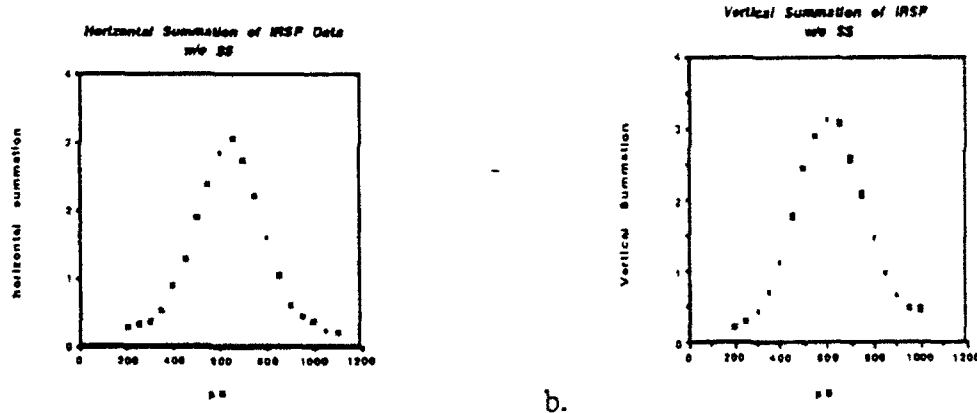


Figure 10. (a) Horizontal and (b) vertical summation profiles of the IRSP image plane without the scatter screen.

Figure 11 is a plot of the modulation transfer function of the IRSP without the scatter screen. The plot shows two curves of modulation transfer versus frequency in cycles per mm. The solid line represents the modulation transfer in a horizontal slice of the MTF. The dotted line represents the vertical MTF.

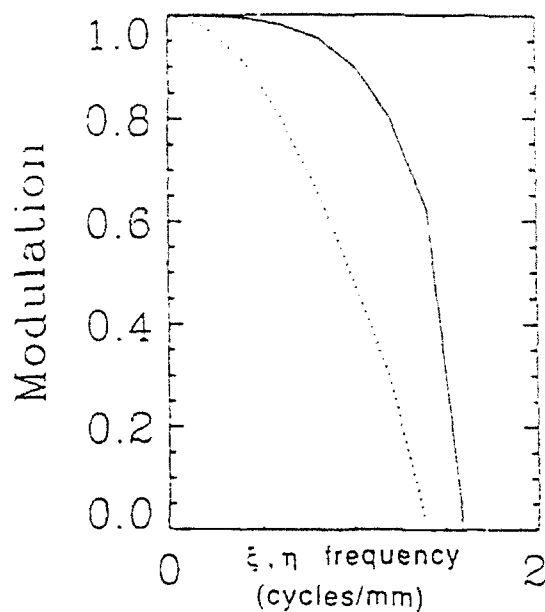


Figure 11. MTF of the IRSP without the scatter screen. The solid line is the horizontal MTF and the dotted line is the vertical MTF in cycles per mm.

Figures 12 and 13 are plots of the point spread function of the IRSP without the scatter screen. In Figure 12 a 3-dimensional plot is provided and Figure 13 is a contour plot of the PSF as it would look viewed from above. The center represents the maximum value of one and each line decreases in magnitude by one tenth.

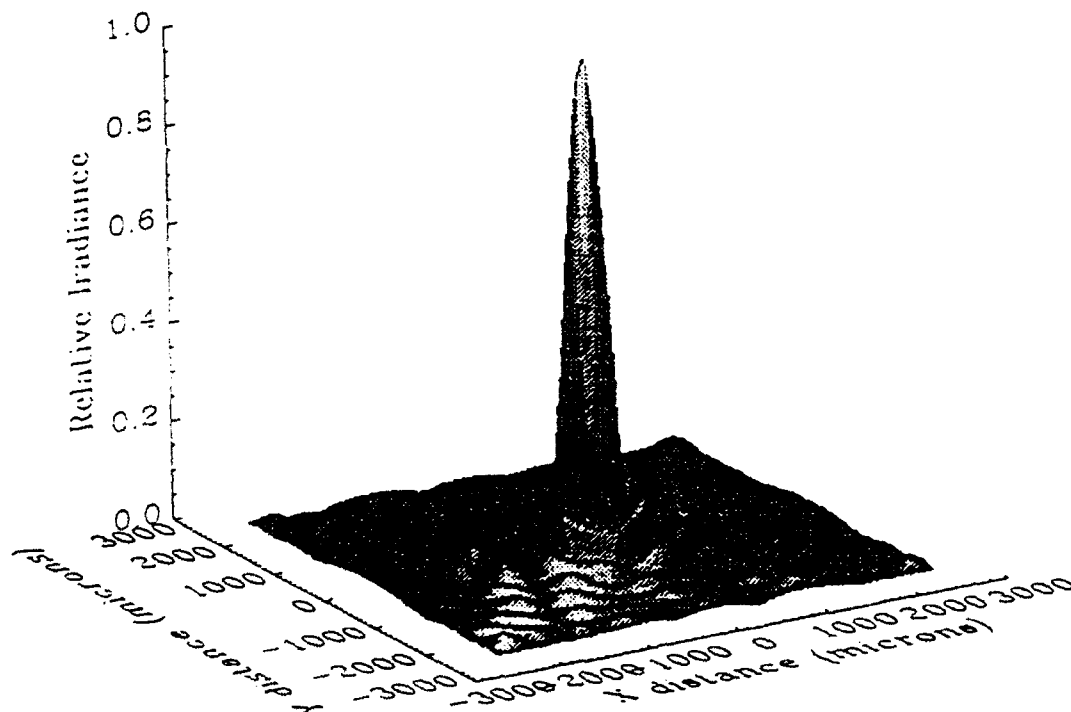


Figure 12. Point Spread Function of the IRSP without the scatter screen.

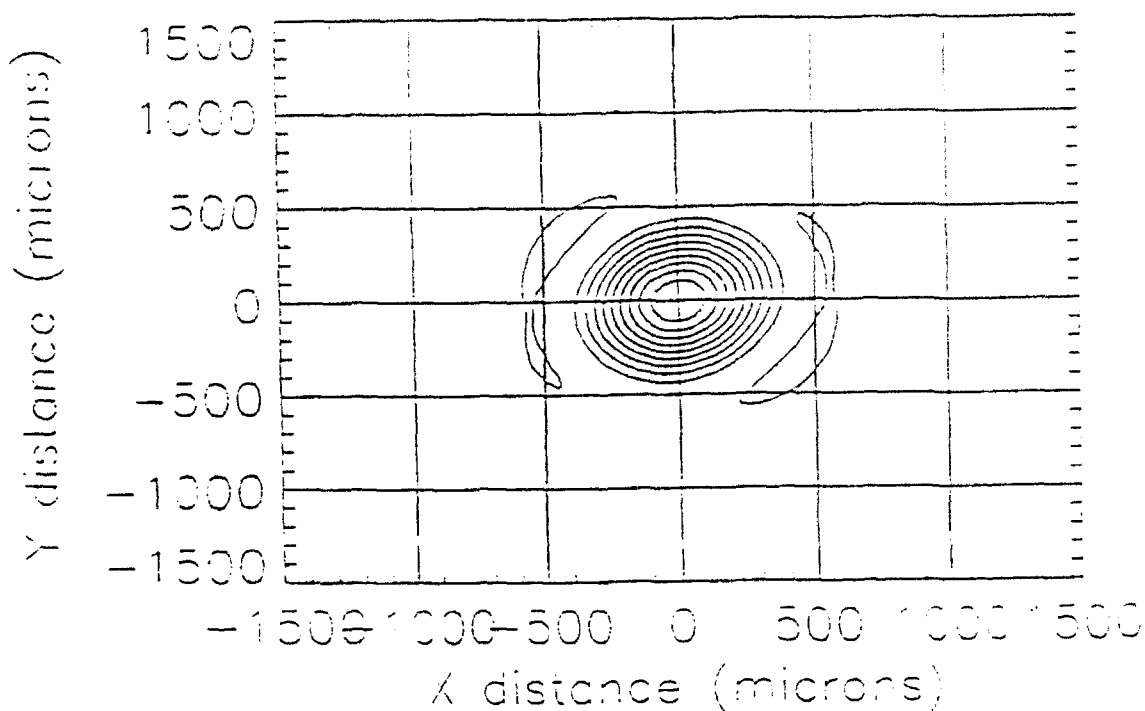


Figure 13. Contour of the IRSP PSF without scatter screen. The center magnitude is one and each level decreases by one tenth.

The modulation transfer function and the point spread function for the IRSP shown in Figure 14 were calculated by hand using the method of Fourier transformation, division of OAP and detector effects and inverse-Fourier transformation described in section II.

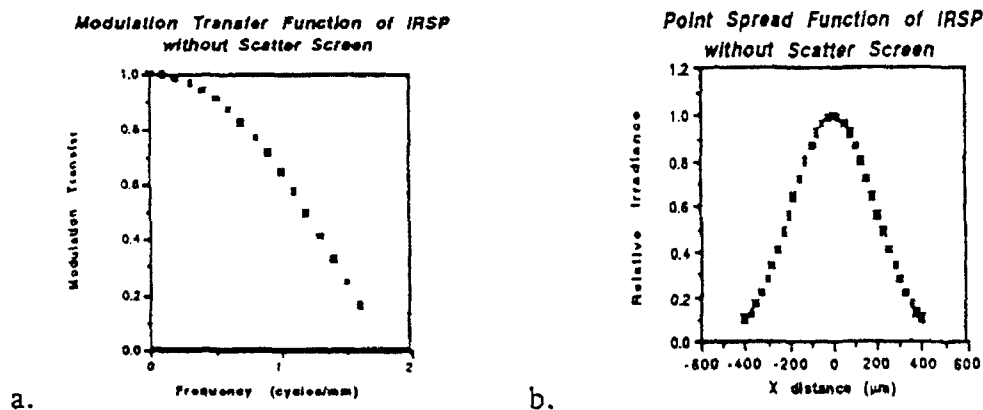


Figure 14 . Hand calculated (a) MTF and (b) PSF of the IRSP without the scatter screen.

Comparison of the hand calculated MTF and the FFT computer generated MTF shows excellent agreement and in fact the actual modulation transfer in the horizontal direction remains stronger at longer frequencies than the calculated MTF. Point spread functions compare ideally.

ii) IRSP with the scatter screen.

With the scatter screen in place, the measured irradiance distribution is quite smooth, uniform and Gaussian (Figure 15). Recall that rotation of the scatter screen destroys laser coherence and

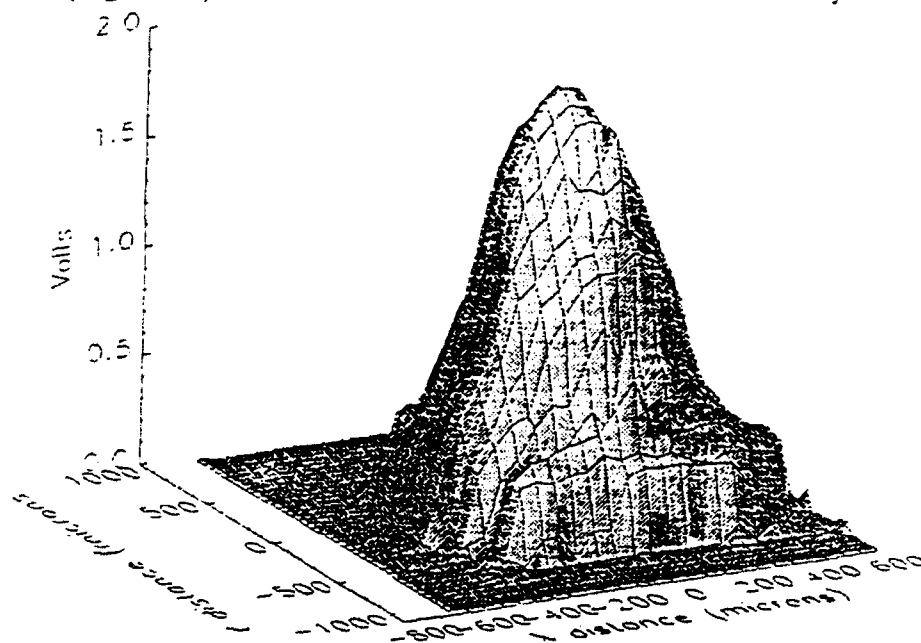


Figure 15. Image plane irradiance distribution for the IRSP with the scatter screen.

eliminates laser speckle. Here there are no signs of degradation due to laser speckle without rotation of the screen. It is possible that coherence has already been lost by the time the beam reaches the scatter screen. The slightly larger distribution can be attributed to the increased angle after passing through the scatter screen.

Figure 16 is the MTF in the horizontal and vertical direction of the IRSP without the scatter screen. A small decrease in modulation transfer can be expected due to the presence of the scatter screen.

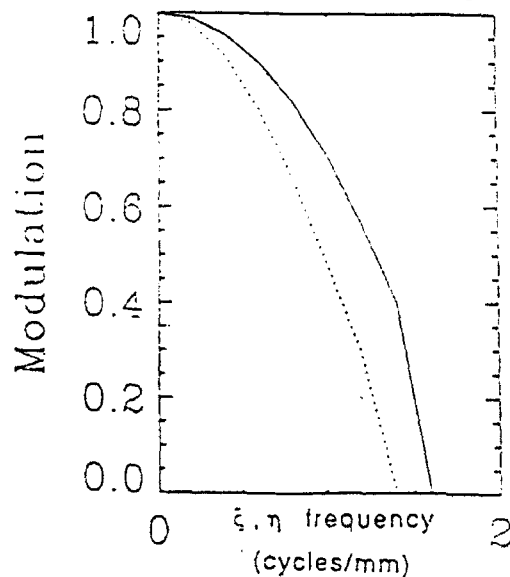


Figure 16. MTF in the X (solid line), and Y (dotted line) of the IRSP with scatter screen.

Figure 17 is the point spread function and Figure 18 is the contour of the PSF of the IRSP with the scatter screen.

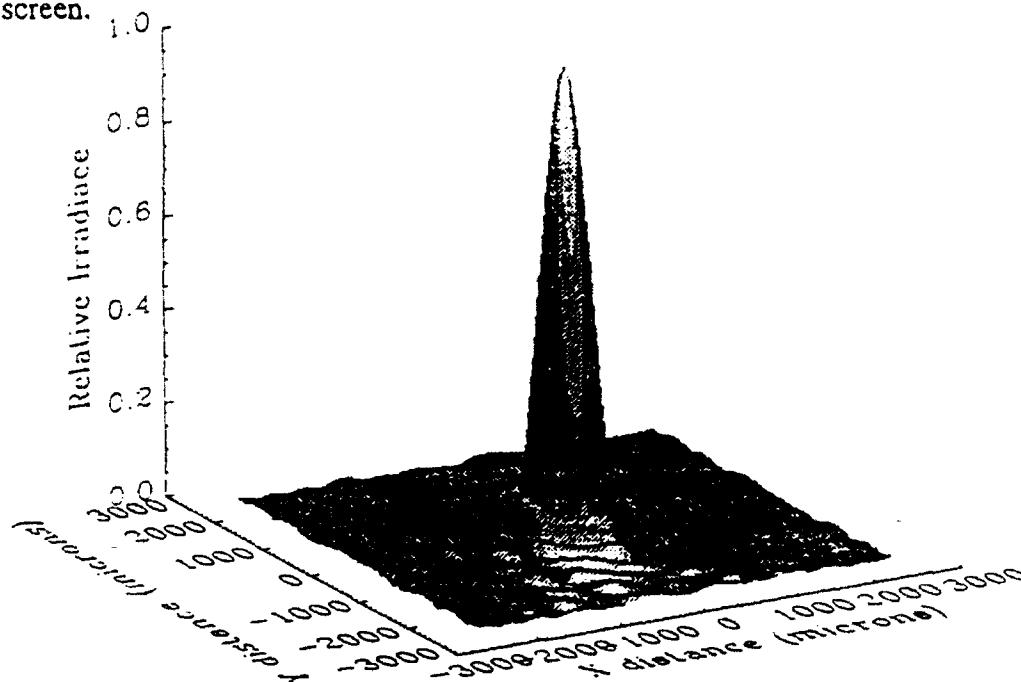


Figure 17. Point Spread Function of the IRSP with the scatter screen.

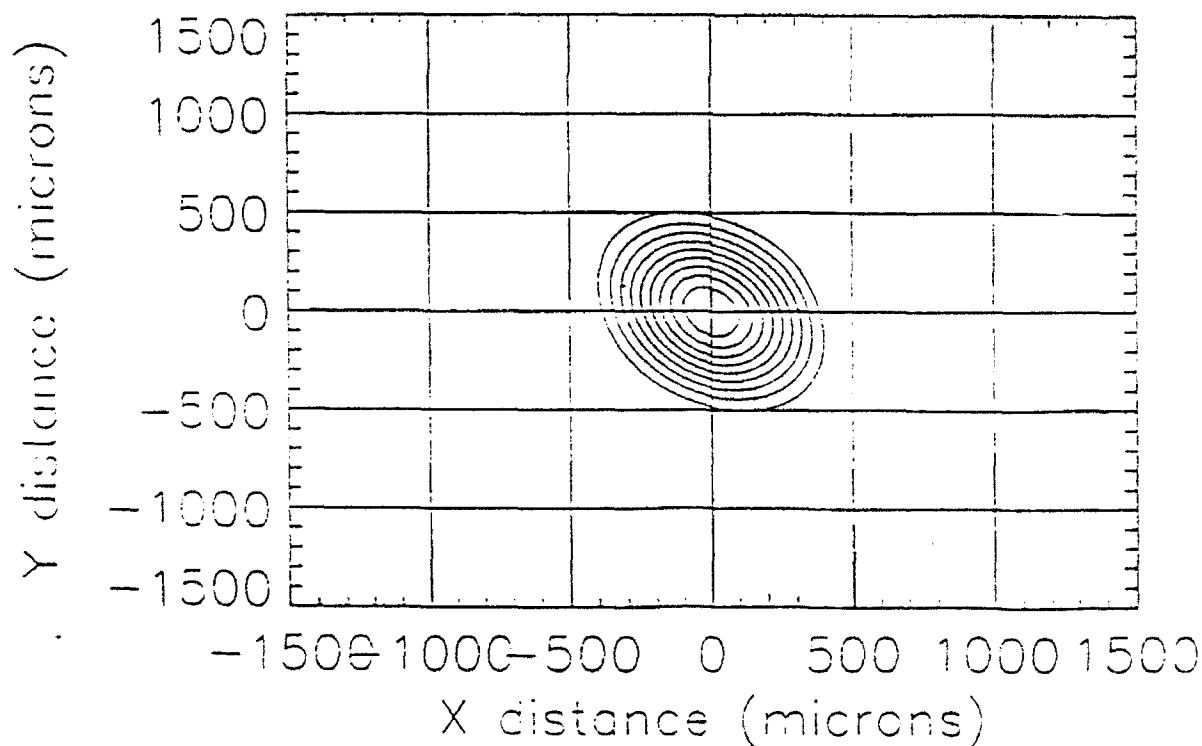


Figure 18. Contour of the IRSP PSF with the scatter screen. The center magnitude is one and each level decreases by one tenth.

iii) IRSP with the scatter screen rotating.

The most interesting result occurred with the scatter screen rotating. The data collection indicates a jagged and fluctuating distribution. As the scatter screen rotates the beam is continuously striking

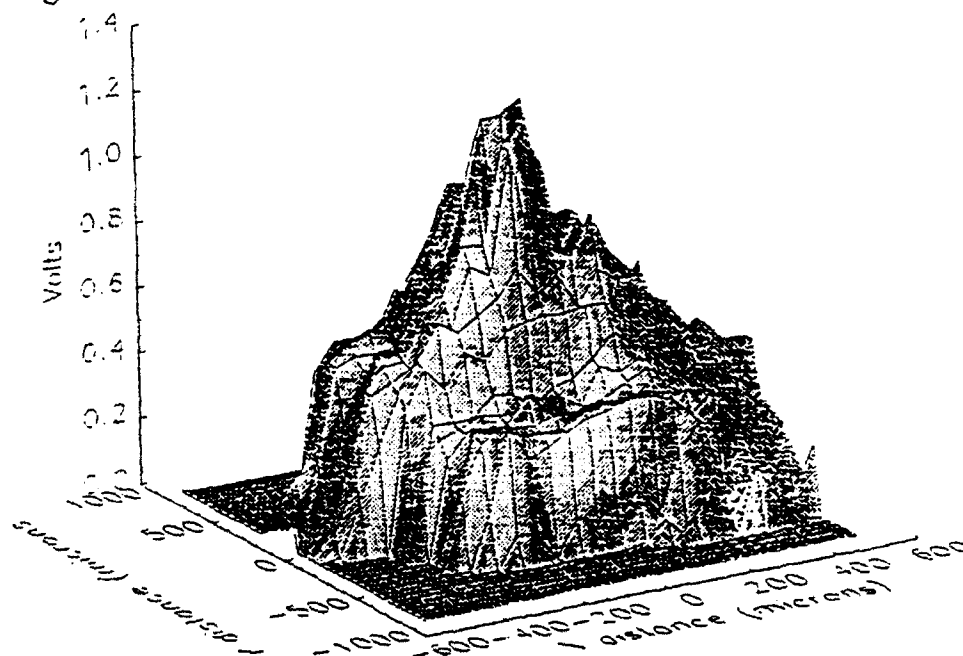


Figure 19. Image plane irradiance distribution of IRSP with scatter screen rotating.

a different position on the screen which has a rough surface. One possible explanation is the slow scan with respect to the frame time. As the detector scans the region, it collects data at approximately 1 sample per second. The IRSP is generating frames at 1 per $320\mu\text{s}$. Each data point is taken at a different frame and pixel dwell time at each sample is at a random position of the scatter screen. Rotation of the screen also introduces a deviation of line of sight and result in a larger spot. Figure 20 is the MTF in the X and Y directions.

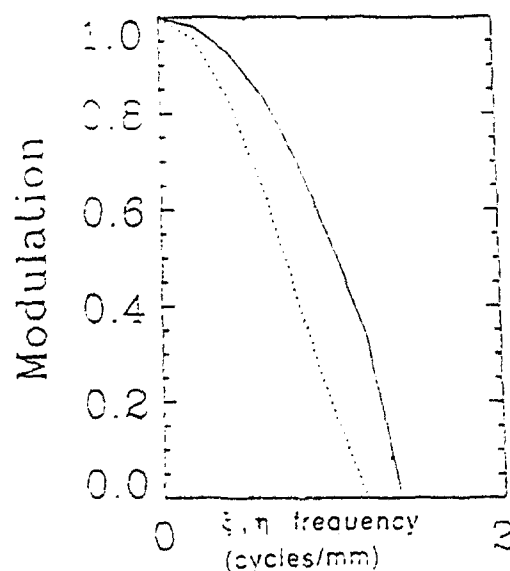


Figure 20. MTF of the IRSP in the X (solid line) and Y (dotted line) directions with the scatter screen rotating.

Figure 21 is the PSF of the IRSP and Figure 22 is the countour plot of the PSF with the scatter screen rotating.

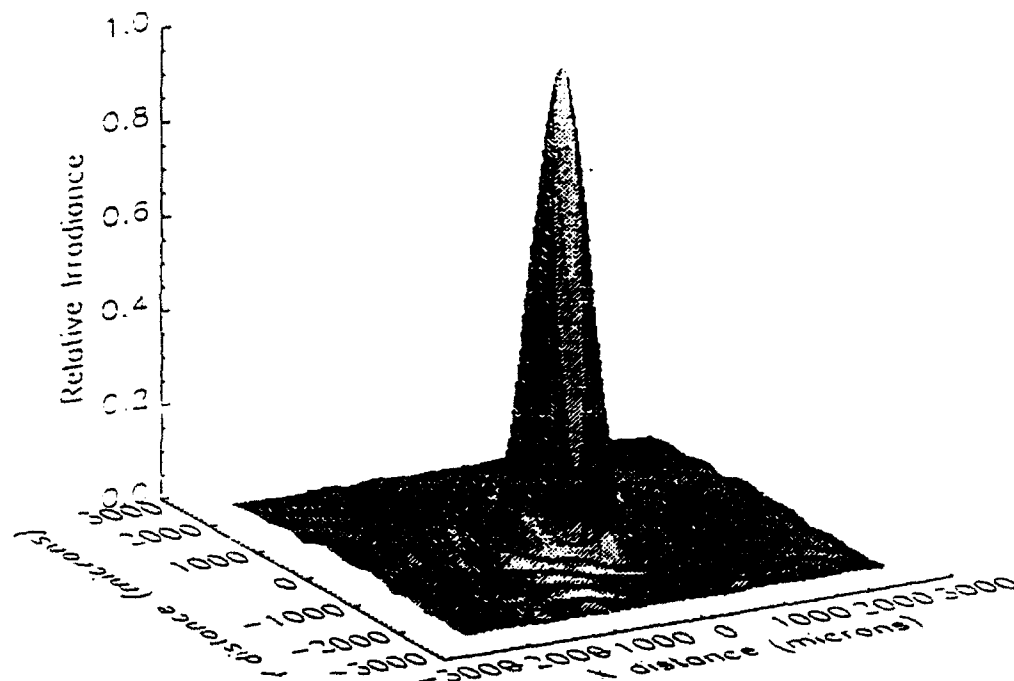


Figure 21. PSF of the IRSP with the scatter screen rotating.

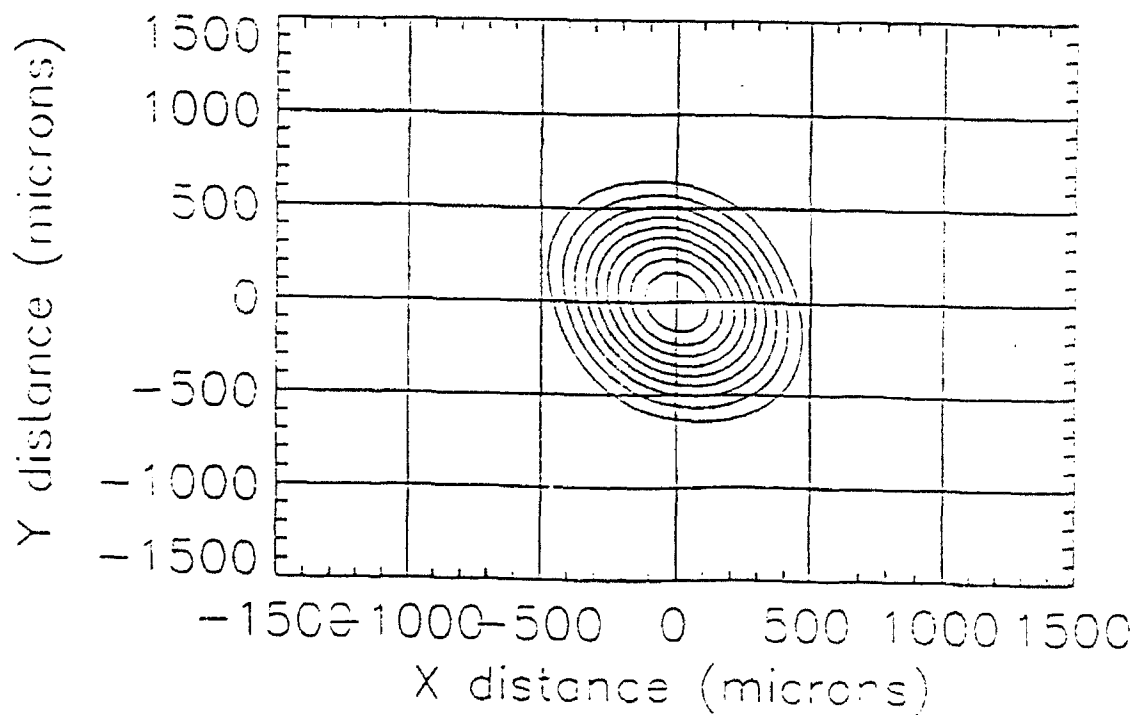


Figure 22. Countour plot of the IRSP PSF with the scatter screen rotating. Center magnitude is one and each level decreases by one tenth.

V CONCLUSION

The objective of this experiment was to determine the point spread and modulation transfer functions of a laser Scophony modulated KHILS infrared scene projector for three separate configurations: (1) without the scatter screen, (2) with the scatter screen in place and (3) with the scatter screen rotating.

The results are shown in section IV where for each configuration we have included graphical plots of 1) the original measured data, 2) an MTF curve in the X and Y directions, 3) a 3-dimensional plot of the PSF and 4) a contour plot of the PSF.

Noteworthy observations include the striking difference in the appearance of the image plane in each of the three cases. The effects of the scatter screen on the measured data, the MTFs, and the PSF are shown and discussed in the results section. It is shown that a defocusing and mismatch of F/#s possibly occurred while the scatter screen was removed and although the SSCS scanning stage is mechanized and calibrated to 0.5 μ m movements in the X and Y directions, there is no similar control for the Z direction. If a similar control was available for detector movement in the Z direction these possible sources of error could be ruled out. Laser speckle degradation was expected while the scatter screen was not rotating but was not evident in the results. Large fluctuations occurred in the measured image plane while the scatter screen was rotating. One possible explanation is the slow scan with respect to frame time.

VI REFERENCES

1. "KHILLS Facility Description and Test Article Interface Document ", Prepared by Guided Interceptor Technology Branch, Wright Laboratories Armament Directorate, Eglin AFB, FL, 1991.
2. JOHNSON, Richard V. "Scophony light valve ", Applied Optics, Vol. 18, No 23, 1 December 1979.
3. SCHILDWACHTER, Eric F., Glenn D. Boreman, "Modulation transfer function characterization and modeling of a Scophony infrared scene projector", Optical Engineering, Vol. 30, No. 11, November 1991.
4. BOREMAN, Glenn D., "Modulation Transfer Function in Optical and Electro-Optic Systems ", SPIE's 1989 Symposia on Aerospace Sensing.
5. SMITH, Warren J., "Modern Optical Engineering ", McGraw-Hill, Inc. 1966.
6. JOHNSON, R. J., "Diffusion and Throughput Measurements of Various Transmissive Screens", Calspan Corporation/AEDC Operations. AEDC-TMR-89-V5, March 1989.
7. HARVEY, James E., Richard A. Rockwell, "Performance characteristics of phased array and thinned aperture optical telescopes", Optical Engineering, Vol. 27, No. 9, pp. 762-768, September 1988.
8. GASKILL, Jack D., "Linear Systems, Fourier Transforms, and Optics ", John Wiley & Sons Inc., 1978.
9. KINGSLAKE, Rudolf, "Optical System Design", Academic Press, Inc. 1983.

THE EFFECT OF NONHOMOGENEOUS INTERPHASES
AND GLOBAL/LOCAL VOLUME FRACTION
ON THE MECHANICS OF A LAYERED COMPOSITE

Vernon T. Bechel
Graduate Student
Department of Mechanical Engineering

University of South Florida
4202 E. Fowler Ave., ENG 118
Tampa, FL 33620-5350

Final Report for:
Summer Research Program
Wright Laboratory

Sponsored by:
Air Force Office of Scientific Research,
Bolling Air Force Base, Washington, D.C.

September 1992

THE EFFECT OF NONHOMOGENEOUS INTERPHASES
AND GLOBAL/LOCAL VOLUME FRACTION
ON THE MECHANICS OF A LAYERED COMPOSITE

Vernon T. Bechel
Graduate Student
Department of Mechanical Engineering
University of South Florida

Abstract

A general layered composite model was developed which could be used to study three families of cracked composite problems. Normalized stress intensity factor (SIF), load diffusion, and stresses versus variation of mechanical properties in a nonhomogeneous interphase region, global volume fraction, or local volume fraction (simulating a defect) could be found.

Results for the global volume fraction parametric study were obtained. Normalized stress intensity factor for E_f/E_m ratios greater than unity decreased as global volume fraction increased in the perfect bond case. For the same configuration load diffusion improved with increasing volume fraction.

Preliminary results were also found for linear and quadratic variations of Young's Modulus in a nonhomogeneous interphase. For E_f/E_m ratios less than one, normalized stress intensity factor decreased as the concentration of the stiffer matrix material near the fiber increased.

THE EFFECT OF NONHOMOGENEOUS INTERPHASES
AND GLOBAL/LOCAL VOLUME FRACTION
ON THE MECHANICS OF A LAYERED COMPOSITE

Vernon T. Bechel

INTRODUCTION

In recent years it has become well-accepted that there exists a region (dubbed the interphase) between the pure fiber material and the bulk matrix material that has neither the fiber properties nor the matrix properties. This region occurs as a result of diffusion, coatings applied to the fiber, or chemical reaction between fiber and matrix. The interphase area is vitally important in the mechanics of a composite since it is through this region that the load is transferred from the matrix to the fiber. The lateral strength of the composite is also determined by the interphase properties. The interphase region and the interface bonds between the fiber, interphase, and matrix can be designed to toughen a composite by allowing the fiber to locally delaminate rather than cracking completely under axial tension loading. (Chamis, 1974; Kaw and Das, 1991) This paper considers the consequences of an interphase region with nonhomogeneous properties.

Another important characteristic of a composite is volume fraction. Often micromechanical fracture problems are solved using techniques such as allowing the portion of the composite away from the crack to have constant properties. This would simulate very low or "zero" volume fraction. The method described here can be used to solve the volume fraction problem for any volume fraction

by using alternating finite width strips of infinite lengths.

FORMULATION

The geometry of the problem is shown in Figure 1. A cracked layer of width $2h_1$ with a crack of length $2a$ that can extend up to the interface is attached to a nonhomogeneous strip whose Young's Modulus and Poisson's ratio are allowed to vary as any polynomial. An arbitrary number of homogeneous layers of differing mechanical properties and widths can be attached to the interphase. The problem is symmetric, and the composite is under a temperature load and remote uniform axial strain. The problem is broken into a cracked and uncracked problem. Only the pressurized crack problem is solved here for a pressure, p . All strips have isotropic properties.

As stated, three families of problems can be solved using this general model. (See Figure 2.) To study the effect of global volume fraction, the nonhomogeneous interphase is given a width of zero, and several layers with alternating properties and widths are attached to the cracked layer. Adding more than seven layers did not significantly alter the SIF or stresses near the cracked layer; therefore results here are for a total of eight layers - four fiber and four matrix.

A similar arrangement could be used to study the effect of a defective matrix layer that has a width larger or smaller than the other matrix layers. This is shown in Figure 2 as "local volume fraction".

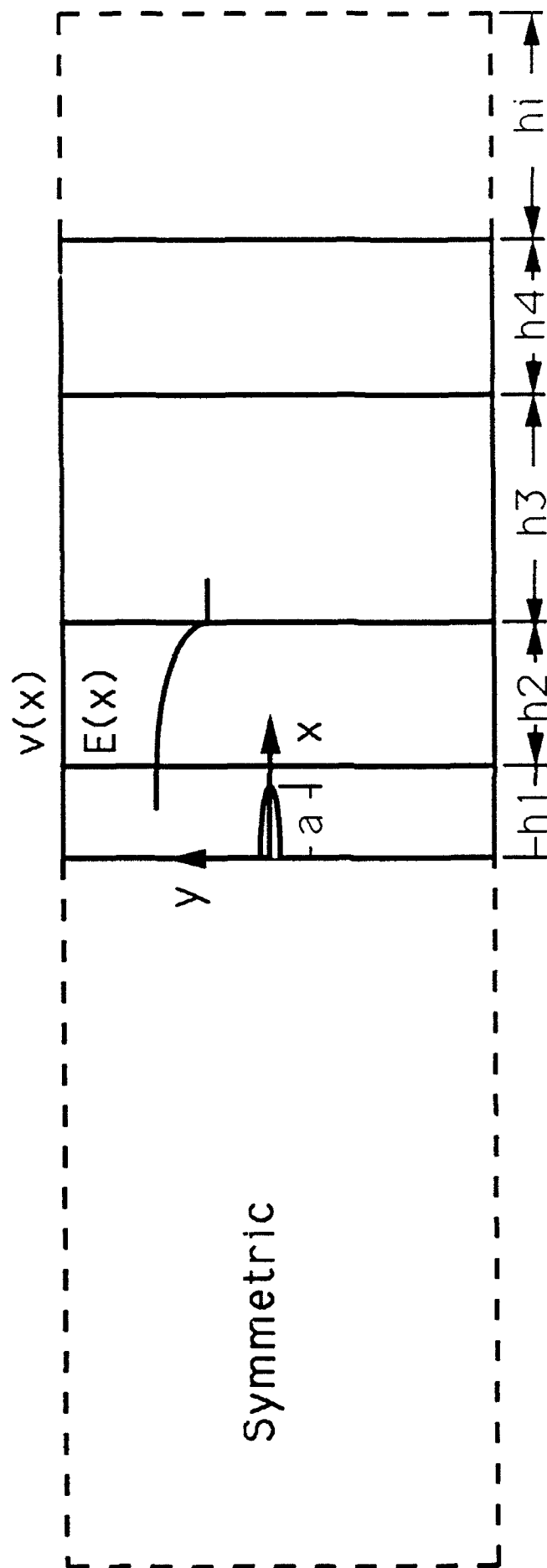


Figure 1 Geometry of the Crack Problem

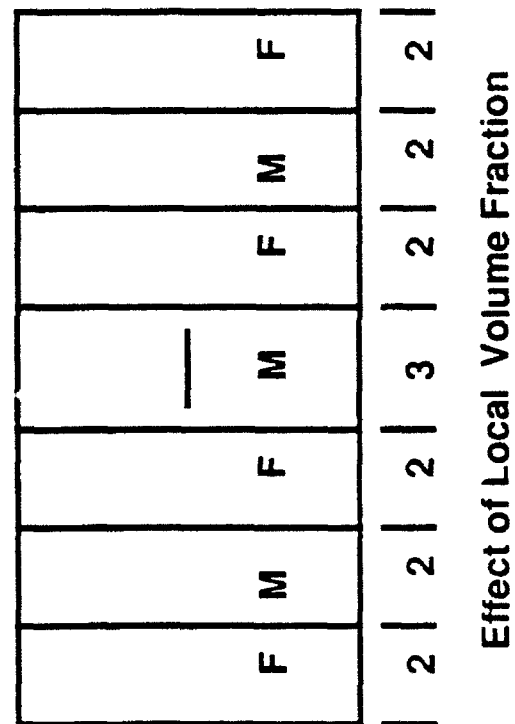
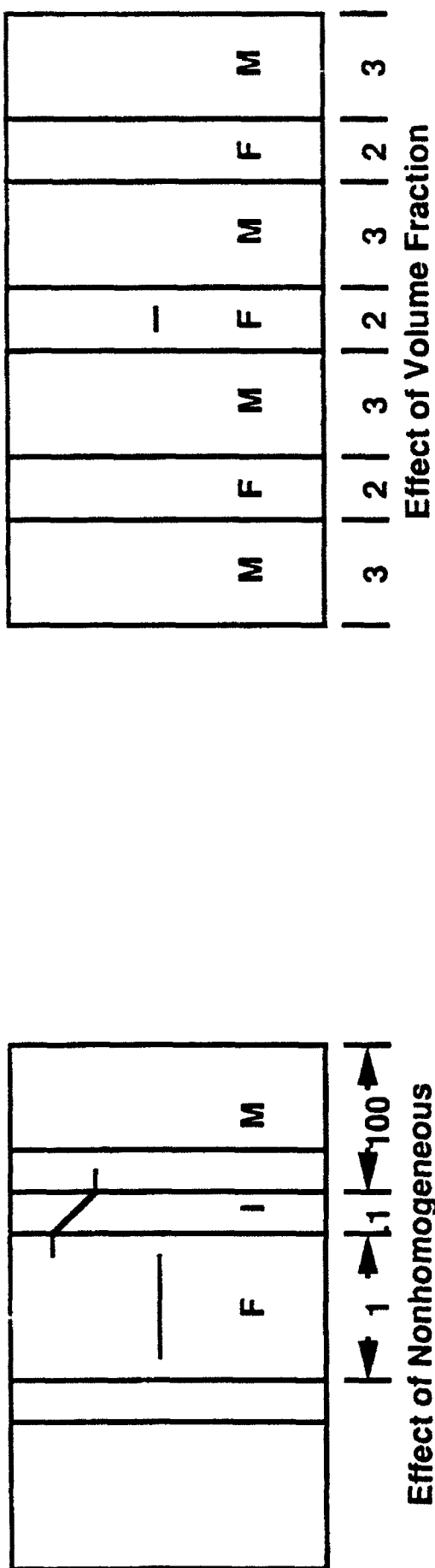


Figure 2 Applications of the General Model

When the width of the interphase is taken to be non-zero, the effect of a nonhomogeneous interphase can be studied. The interphase has to be broken into several exponentially varying strips, as shown in Figure 3, where the properties of each interphase strip are piecewise continuous and have the following form:

$$E_i(x) = E_0^i e^{\beta_i x}$$

$$v_i(x) = (A_0^i + B_0^i x) e^{\beta_i x}$$

where i represents the number of the strip. The E_0 's, A_0 's, B_0 's, and β 's are found by imposing continuity conditions. The exponential model was used because the stress and displacement field equations can be found for that type of nonhomogeneity (Delale and Erdogan, 1988). The number of strips that the interphase is broken into is governed by a prespecified tolerance between the actual value for the properties and the value given by the particular exponential model in each individual strip. The example in Figure 3 shows an interphase with linearly varying properties approximated by four narrower strips of exponentially varying properties. 0.01 was used as the prespecified tolerance.

Using this technique, several variations of the mechanical properties in the nonhomogeneous interphase and their effect on SIF and load diffusion could be investigated. Linear and quadratic functions for the Young's modulus were used.

$$E(x) = C \exp(\beta x)$$

$$v(x) = (A + Bx) \exp(\beta x)$$

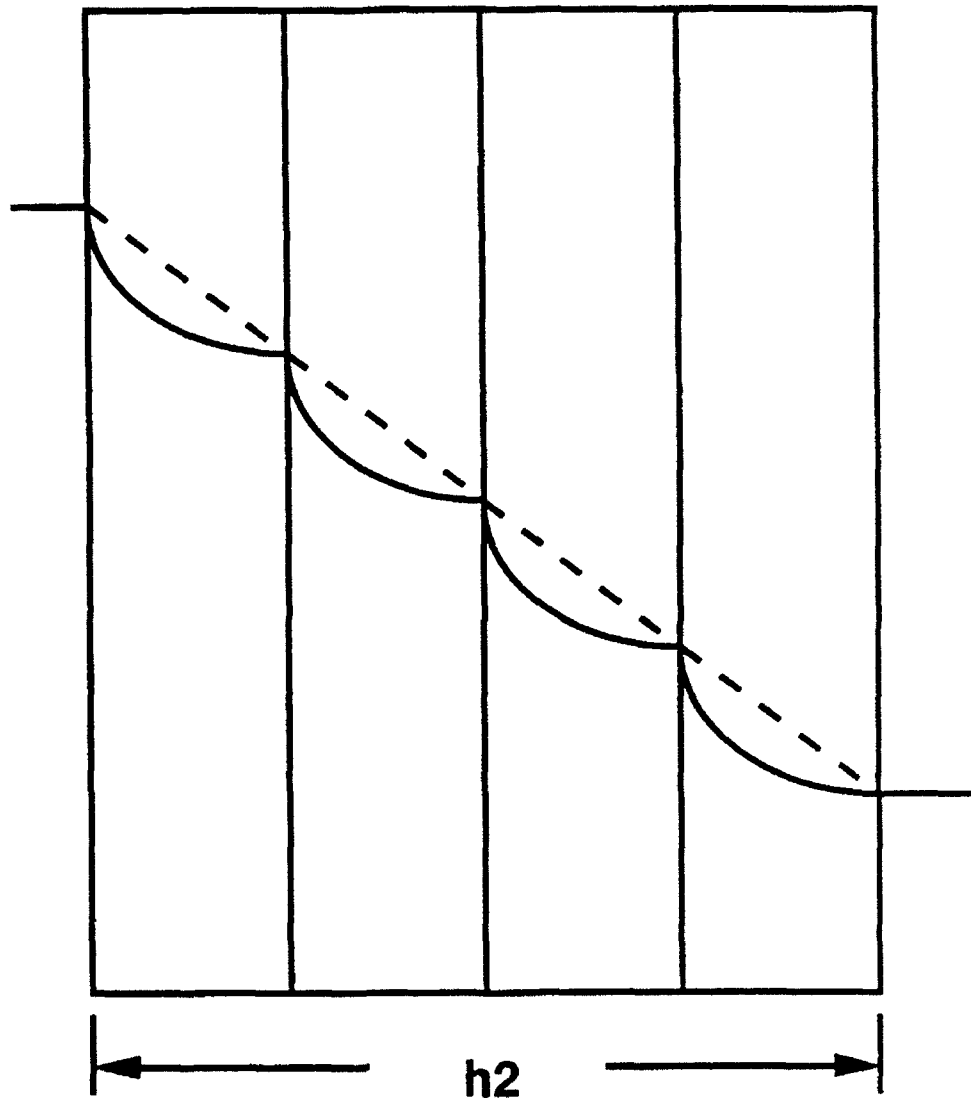


Figure 3 Breakdown of Interphase into Exponential Segments

Stress and Displacement Field Equations

The displacement and stress field equations for an infinitely long strip with a crack in it are given by Sneddon and Lowengrub, 1969.

$$\begin{aligned}
 u_1(x, y) = & - \frac{2}{\pi} \int_0^{\infty} \left(\frac{1}{\eta} \left[f_1(\eta) - \frac{\kappa_1 - 1}{2} g_1(\eta) \right] \sinh(\eta x) \right. \\
 & + \left. x g_1(\eta) \cosh(\eta x) \right) \cos(\eta y) d\eta \\
 & - \frac{2}{\pi} \int_0^{\infty} \frac{\phi_1(\xi)}{\xi} \left(\frac{\kappa_1 - 1}{2} - \xi y \right) e^{-\xi y} \sin(\xi x) d\xi, \\
 v_1(x, y) = & \frac{2}{\pi} \int_0^{\infty} \left(\frac{1}{\eta} \left[f_1(\eta) + \frac{\kappa_1 + 1}{2} g_1(\eta) \right] \cosh(\eta x) \right. \\
 & + \left. x g_1(\eta) \sinh(\eta x) \right) \sin(\eta y) d\eta \\
 & + \frac{2}{\pi} \int_0^{\infty} \frac{\phi_1(\xi)}{\xi} \left(\frac{\kappa_1 + 1}{2} + \xi y \right) e^{-\xi y} \cos(\xi x) d\xi.
 \end{aligned}$$

$$\begin{aligned}
 \sigma_{xx}^1(x, y) = & - \frac{4\mu_1}{\pi} \int_0^{\infty} \left[f_1(\eta) \cosh(\eta x) + \eta x g_1(\eta) \sinh(\eta x) \right] \cos(\eta y) d\eta \\
 & - \frac{2}{\pi} \int_0^{\infty} \phi_1(\xi) (1 - \xi y) e^{-\xi y} \cos(\xi x) d\xi, \\
 \sigma_{yy}^1(x, y) = & \frac{4\mu_1}{\pi} \int_0^{\infty} \left(f_1(\eta) + 2g_1(\eta) \right) \cosh(\eta x) \\
 & + \eta x g_1(\eta) \sinh(\eta x) \cos(\eta y) d\eta \\
 & - \frac{4\mu_1}{\pi} \int_0^{\infty} \phi_1(\xi) (1 + \xi y) e^{-\xi y} \cos(\xi x) d\xi \\
 \sigma_{xy}^1(x, y) = & \frac{4\mu_1}{\pi} \int_0^{\infty} \left([f_1(\eta) + g_1(\eta)] \sinh(\eta x) \right. \\
 & + \left. \eta x g_1(\eta) \cosh(\eta x) \right) \sin(\eta y) d\eta \\
 & - \frac{4\mu_1}{\pi} \int_0^{\infty} \xi y \phi_1(\xi) e^{-\xi y} \sin(\xi x) d\xi.
 \end{aligned}$$

The displacement and stress field equations for the nonhomogeneous interphase strips (exponential case) are given by Delale and Erdogan, 1988.

$$\begin{aligned}
 u_1(x, y) = & \frac{2}{\pi E_0} \int_0^{\infty} \left(\frac{\eta^2}{2} \left(\left[\frac{c_2(\eta)x}{m_2} + \frac{c_2(\eta)}{m_2^2} + \frac{c_1(\eta)}{m_2} \right] e^{-m_2 x} \right. \right. \\
 & + \left. \left[\frac{c_4(\eta)x}{m_1} + \frac{c_4(\eta)}{m_1^2} + \frac{c_3(\eta)}{m_1} \right] e^{m_1 x} \right. \\
 & - [b_0 c_2(\eta) m_1 x^2 + m_1 x (a_0 c_2(\eta) + b_0 c_1(\eta)) + a_0 c_1(\eta) m_1 \\
 & - b_0 c_1(\eta) + a_0 c_2(\eta)] \frac{e^{m_1 x}}{2} \\
 & - [b_0 c_4(\eta) m_2 x^2 + m_2 x (a_0 c_4(\eta) + b_0 c_3(\eta)) \\
 & + a_0 c_3(\eta) m_2 - b_0 c_3(\eta) + a_0 c_4(\eta)] \frac{e^{m_2 x}}{2} \left. \right) \cos(\eta y) d\eta,
 \end{aligned}$$

$$\begin{aligned}
 v_1(x, y) = & \frac{2}{\pi E_0} \int_0^{\infty} \left(\frac{1}{2\eta} (m_1^2 (c_1(\eta) + c_2(\eta)x) + 2m_1 c_2(\eta)) e^{-m_1 x} \right. \\
 & + (m_2^2 (c_3(\eta) + c_4(\eta)x) + 2m_2 c_4(\eta)) e^{-m_2 x} \\
 & + (a_0 + b_0 x) \eta [(c_1(\eta) + c_2(\eta)x) e^{m_1 x} \\
 & + (c_3(\eta) + c_4(\eta)x) e^{m_2 x}] \sin(\eta y) d\eta,
 \end{aligned}$$

$$\begin{aligned}
 \sigma_{xx}^1(x, y) = & -\frac{2}{\pi} \int_0^{\infty} \frac{\eta^2}{2} [(c_1(\eta) + c_2(\eta)x) e^{m_1 x} \\
 & + (c_3(\eta) + c_4(\eta)x) e^{m_2 x}] \cos(\eta y) d\eta,
 \end{aligned}$$

$$\begin{aligned}
 \sigma_{yy}^1(x, y) = & \frac{2}{\pi} \int_0^{\infty} \frac{1}{2} ([m_1^2 (c_1(\eta) + c_2(\eta)x) + 2m_1 c_2(\eta)] e^{m_1 x} \\
 & + [m_2^2 (c_3(\eta) + c_4(\eta)x) + 2m_2 c_4(\eta)] e^{m_2 x}) \cos(\eta y) d\eta,
 \end{aligned}$$

$$\begin{aligned}
 \sigma_{xy}^1(x, y) = & \frac{2}{\pi} \int_0^{\infty} \frac{\eta}{2} ((m_1 (c_1(\eta) + c_2(\eta)x) + c_2(\eta)) e^{m_1 x} \\
 & + [m_2 (c_3(\eta) + c_4(\eta)x) + c_4(\eta)] e^{m_2 x}) \sin(\eta y) d\eta.
 \end{aligned}$$

$$m_1 = \frac{\beta}{2} - (\eta^2 + \frac{\beta^2}{4})^{1/2}, \quad m_2 = \frac{\beta}{2} + (\eta^2 + \frac{\beta^2}{4})^{1/2},$$

Boundary Conditions

The above stress and displacement fields equations were then employed by applying appropriate boundary conditions. Continuity of displacements and tractions at the interface between each infinite strip was enforced as follows:

$$\begin{aligned}\sigma_{xx}^i(x_i, y) &= \sigma_{xx}^{i+1}(x_{i+1}, y) \\ \sigma_{yy}^i(x_i, y) &= \sigma_{yy}^{i+1}(x_{i+1}, y) \\ u^i(x_i, y) &= u^{i+1}(x_{i+1}, y) \\ v^i(x_i, y) &= v^{i+1}(x_{i+1}, y)\end{aligned}$$

==> 4(n-1) equations.

The free edge condition at the outer layer of the model was then accounted for as follows:

$$\begin{aligned}\sigma_{xx}^n(x_n, y) &= 0 \\ \sigma_{xy}^n(x_n, y) &= 0\end{aligned}$$

==> 2 equations.

Finally, mixed boundary conditions were applied at $y = 0$ in the cracked layer as follows:

$$\begin{aligned}\sigma_{yy}^1(x, 0) &= -p(x) & |x| < a & \quad (1) \\ v_1(x, 0) &= 0 & a < |x| < h & \quad (2)\end{aligned}$$

Define $\frac{\partial v_1}{\partial x}(x, 0) = G(x)$, invert, substitute into (1)

==> 1 equation.

This results in a total of $4(n-1)+3$ equations. Inspecting the number of unknown functions present in the stress and displacement field equations reveals the following:

3 in cracked strip 3 unknowns

4 in each attached strip 4(n-1) unknowns

4(n-1)+3 total unknowns.

Singular Integral Equation

A single singular integral equation (SIE) of the following form results

$$\int_{-a}^a \frac{G(t)}{t-x} dt + \int_{-a}^a G(t) K(t,x) dt = \frac{-\pi p(x) (1+\kappa_1)}{4\mu_1}$$

The asymptotic analysis for a crack located in a homogeneous material and with its tip impinging on a nonhomogeneous material was already done by Kaw and Selvarathinam (1990). Their investigation verified that the singularity is still of the type $G(t) = 1/(a^2 - t^2)^\gamma$ where γ can be found by solving the characteristic equation below.

$$2\cos\pi\gamma + 4\lambda_2(\gamma - 1)^2 - (\lambda_1 + \lambda_2) = 0$$

λ_1 and λ_2 are functions of the material properties of the two materials that the crack tip is embedded in.

The system of equations was then solved by solving the system (not including the SIE) at each value of η (required by the solution of the SIE) to get $f_1(\eta)$ and $g_1(\eta)$ in terms of integrals of the unknown function $G(t)$. The SIE was then solved by normalizing and converting it to a summation as follows

$$\int_{-1}^1 \frac{\psi(\tau)}{(1-\tau^2)^\gamma} \left[\frac{1}{\tau - y} + aK(a\tau, ay) \right] d\tau = \text{constant}$$

$$\sum_{j=1}^N A_j \psi(\tau_j) \left[\frac{1}{\tau_j - y_1} + aK(a\tau_j, ay_1) \right] = \text{constant}$$

where τ_j 's, y_1 's, and A_j 's were chosen by Gauss-Jacobi integration formulas. To be able to solve the equations in the above

summation, a final equation is required - the single-valuedness condition.

$$\sum_{j=1}^N A_j \psi(\tau_j) = 0$$

Stress Intensity Factor

The stress intensity factor was calculated based on the following formula

$$K = - 2\sqrt{2}\mu^* \left(\frac{a}{2}\right)' \psi(1)$$

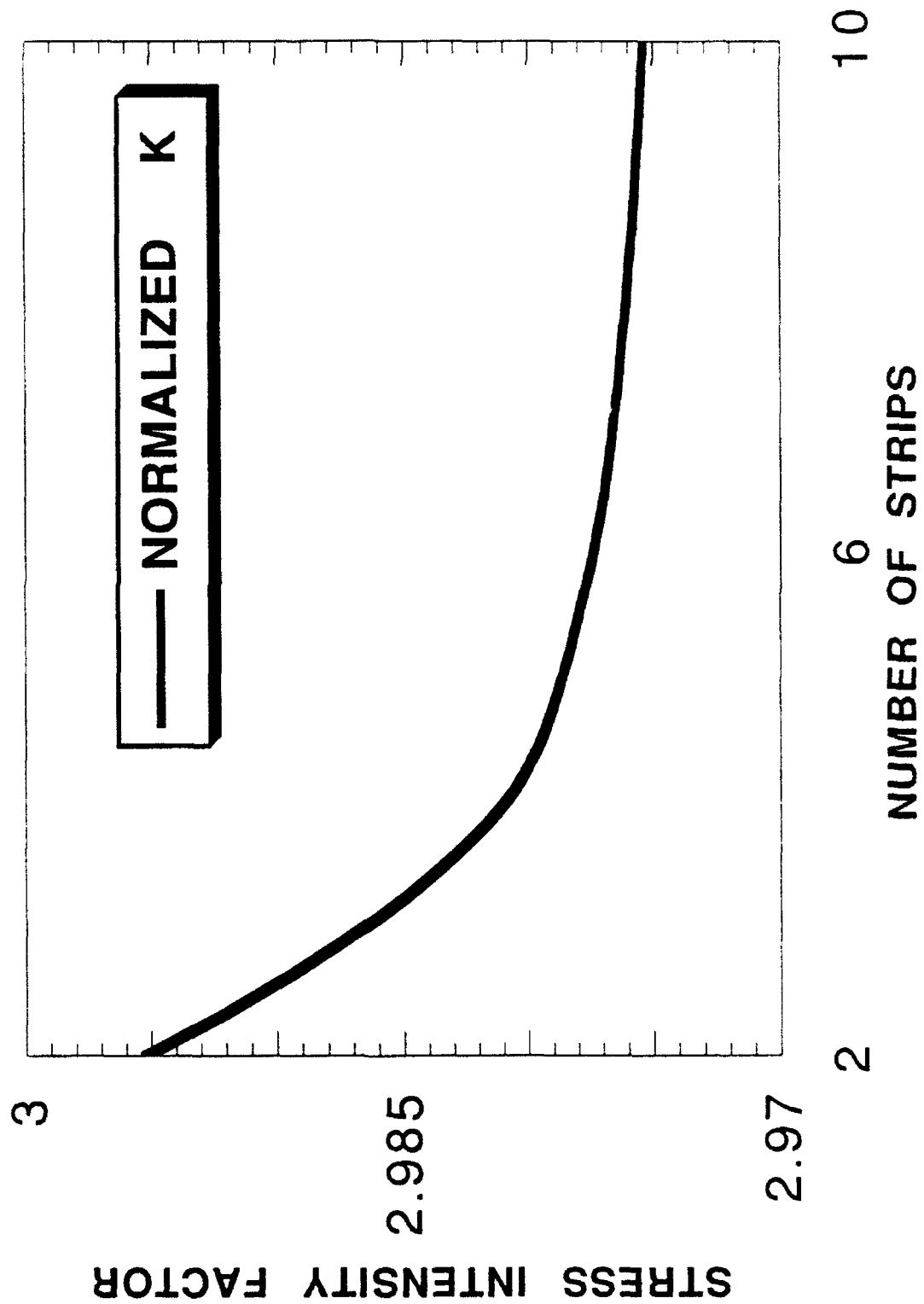
Gupta(1973), where μ^* is a function of material properties. The SIF was then normalized with respect to $pa^{1/2}$.

Results and Conclusions

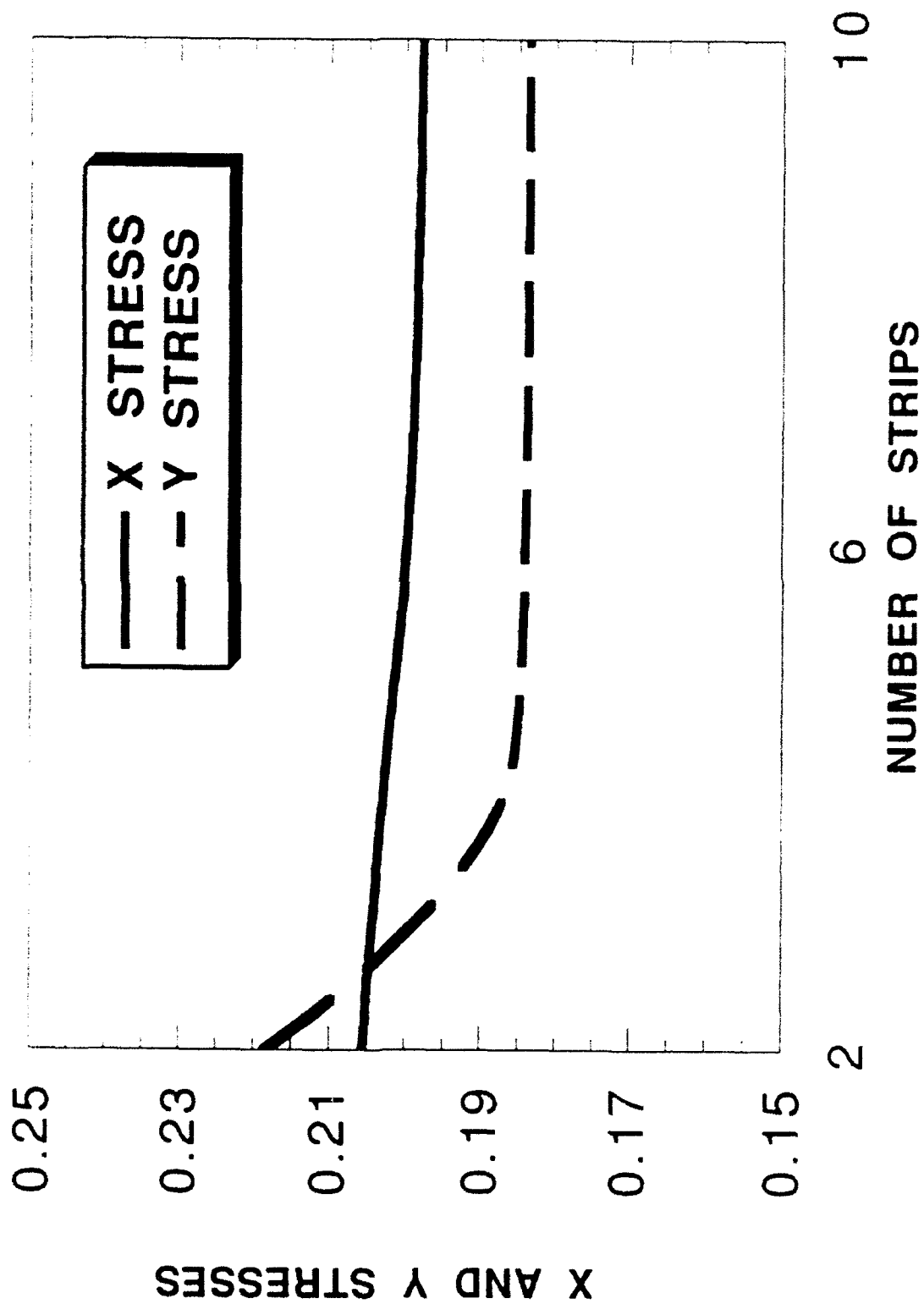
Because of time constraints, few results were obtained. As stated, a study was done to determine how many strips were required for consistent results for the global and local volume fraction results. The extreme case of volume fraction equal to .5 and fiber to matrix moduli ratio equal to 1/20 was used. The x and y normal stresses and the SIF had a changed negligibly when adding more than seven layers. (See Figures 4 and 5.) Therefore, eight total strips were chosen used.

Normalized stress intensity factor and load diffusion results were obtained for the global volume fraction problems. (See Figures 6 and 7.) E_f/E_m was calculated in all cases based on the cracked layer being considered to be a fiber layer. Stress intensity factor for arrangements with fibers stiffer than the matrix

**Figure 4 Normalized Stress Intensity Factor
($K/p\sqrt{a}$) vs. Number of Strips Used**



**Figure 5 X and Y Normal Stresses at $x=2$, $y=0$
vs. Number of Strips used.**



**Figure 6 Normalized Stress Intensity Factor
($K/K_{Vf=0}$) vs. Global Volume Fraction**

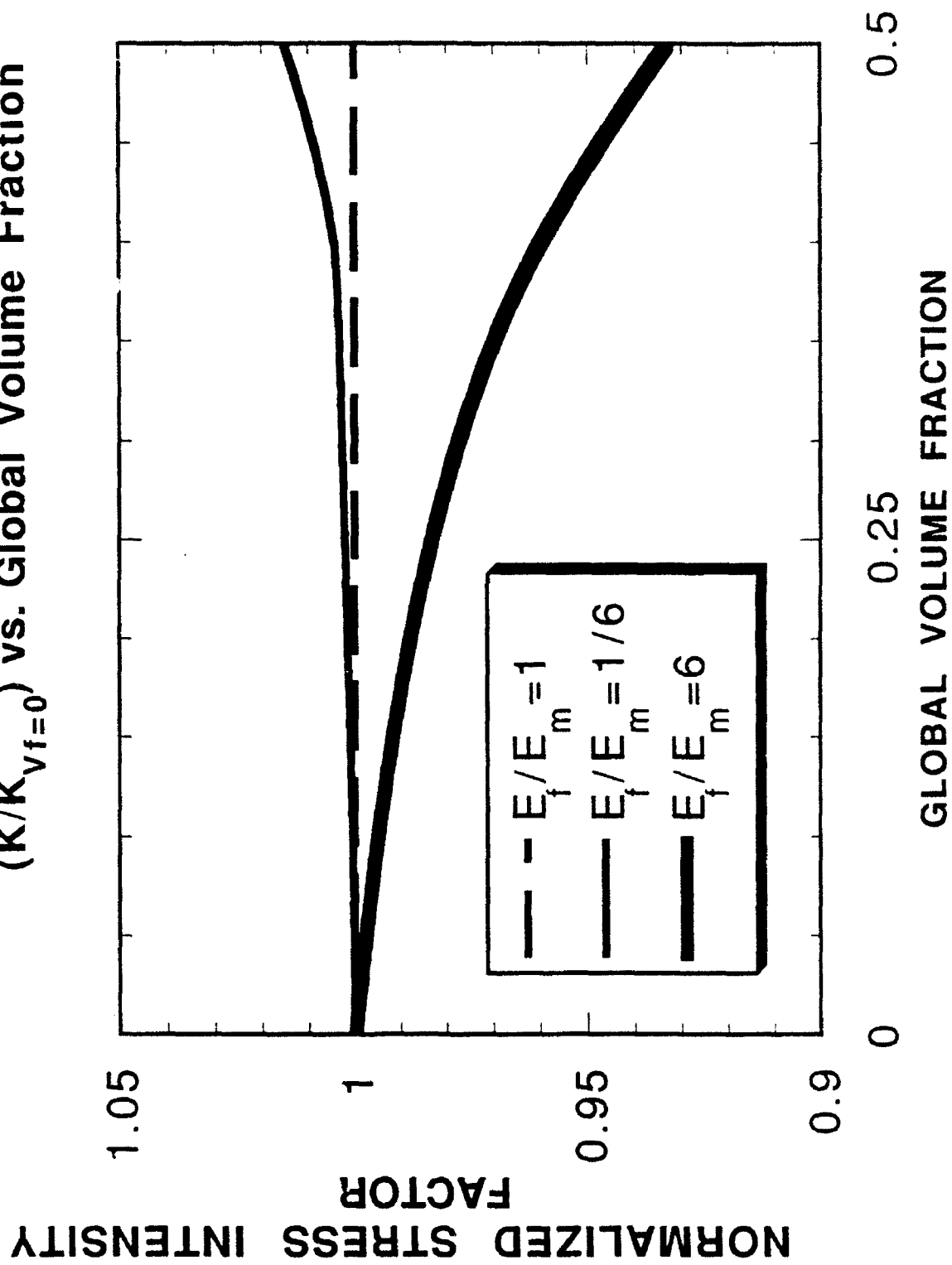
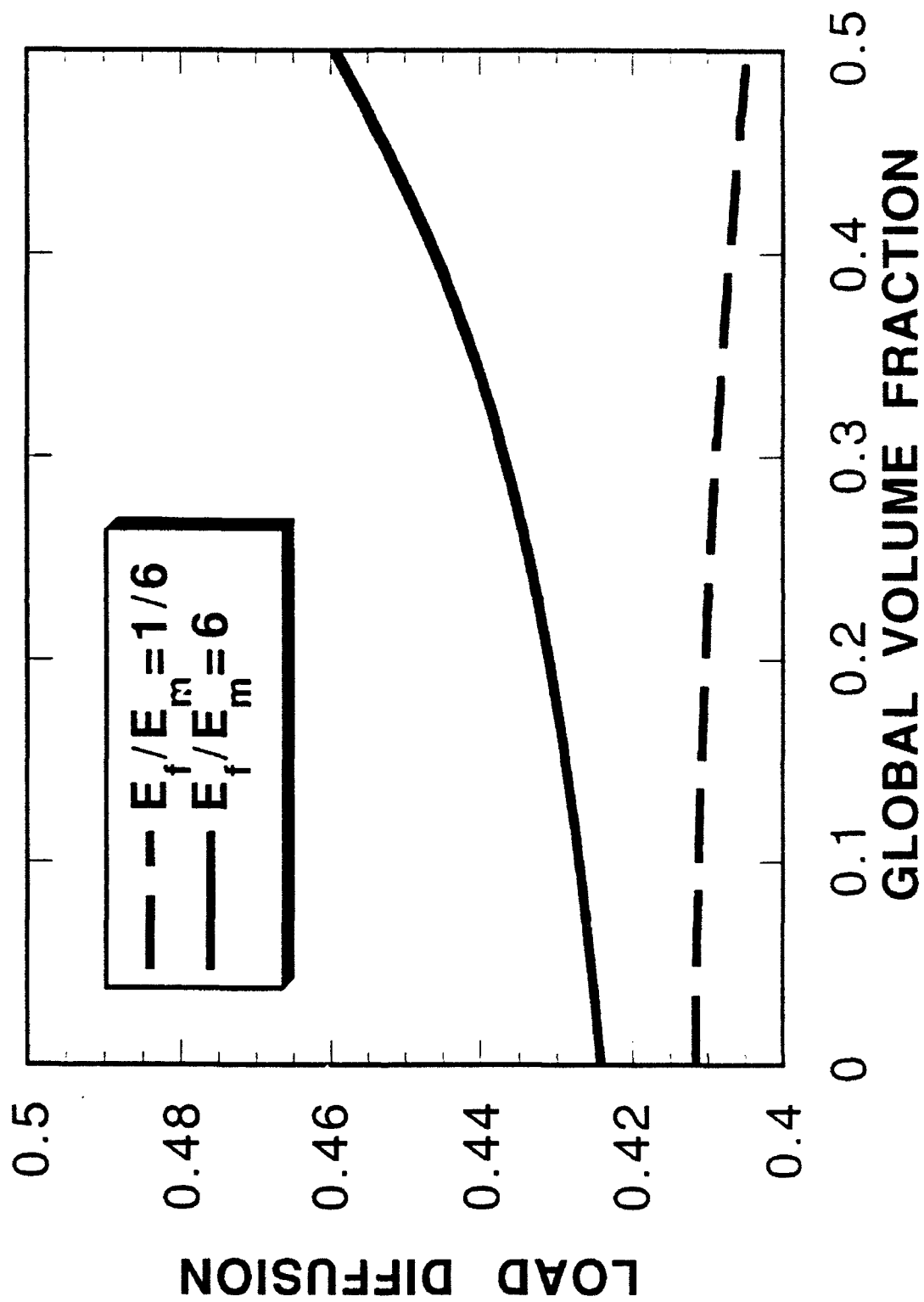


Figure 7 Load Diffusion (average y-normal stress in the cracked fiber at $y=1$) vs. Volume Fraction



decreased as volume fraction increased. For the same configuration load diffusion improved with increasing volume fraction. These results can be attributed to the cracked fiber sensing the second fiber layer moving closer as the volume fraction increases. As the second fiber layer moves closer, it carries more of the load that the cracked layer would have been required to carry.

Initial results were also acquired for the nonhomogeneous interphase family of problems (not tabulated here). The variations of Young's modulus in the interphase that were studied were the linear case and the two quadratic cases with the slope of the Young's modulus function zero at the either the right or left edge of the interphase. (See Figure 9.) For E_f/E_m ratios less than one, the linear case results in a greater average Young's modulus near the cracked fiber-matrix interface than the quadratic case with zero slope at the cracked fiber-matrix interface. The quadratic case with zero slope at the right edge of the interphase resulted in a still higher average Young's modulus near the crack tip. For E_f/E_m ratios less than one, the normalized stress intensity factor decreased as average Young's modulus near the crack tip increased. This can be attributed to the interphase near the crack tip becoming stiffer on the average and carrying more of the load that the cracked fiber would otherwise have to carry.

Results have not yet been obtained for the local volume fraction problems or load diffusion in the nonhomogeneous interphase problems.

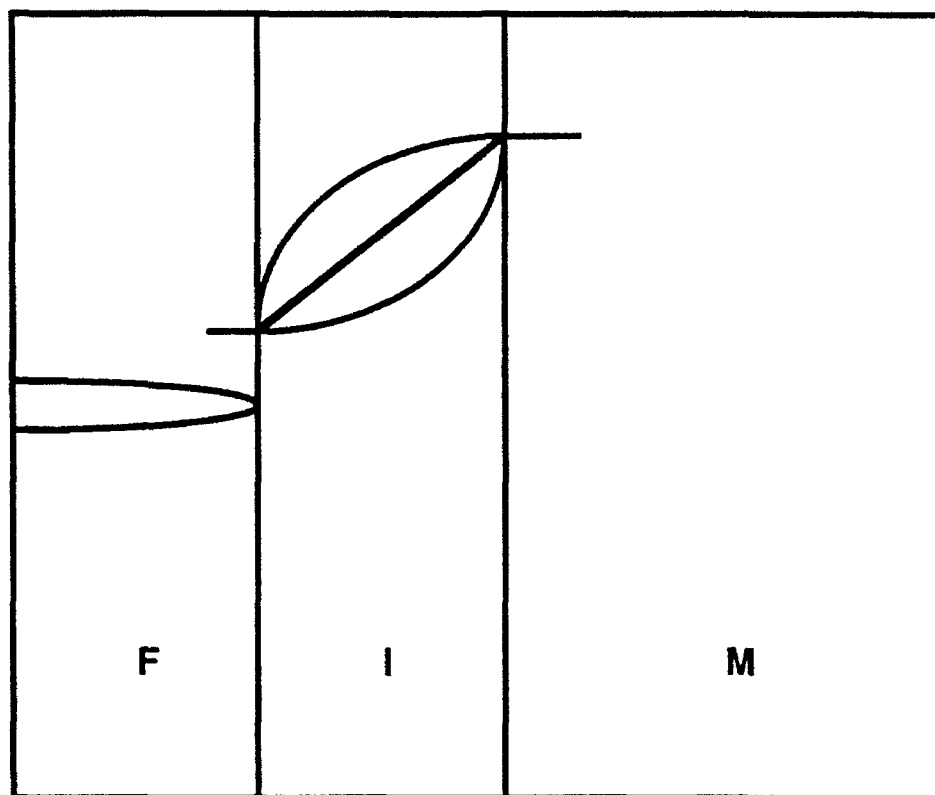


Figure 9 Interphase Models

References

Chamis, C.C. (1974). Mechanics of load transfer at the interface. E.P. Plueddemann, editor, Comp. Mats. 6, 32-77. Academic Press, New York.

Delale, F. and Erdogan, F (1988). On the Mechanical Modeling of the Interfacial Region in Bonded Half-Planes. ASME J. Appl. Mech. 55, 317-324.

Gupta, G.D. (1973). A Layered Composite with a Broken Laminate. Int. J. Solids Str. 9, 1141-1154.

Kaw, A.K. and Das, V.G. (1991). Crack in an imperfect interface in composites. Mech. of Mat. 11, 295-322.

Selvarathinam, A.S. (1991). Comparison of Interphase Models for a Fracture Problem in Fiber Reinforced Composite Materials. M.S. Thesis, University of South Florida, Tampa, FL. (A.K. Kaw advisor).

Sneddon, I.N. and Lowengrub, M (1969). Crack Problems in the Classical Theory of Elasticity. John Wiley and Sons, New York.

VELOCITY AND TEMPERATURE MEASUREMENTS
IN A HIGH SWIRL DUMP COMBUSTOR

Dr. Richard D. Gould and Mr. Lance H. Benedict

Mechanical and Aerospace Engineering
North Carolina State University
Raleigh, NC 27695

Final Report for:
Summer Research Program
Wright Laboratory

Sponsored by:
Air Force Office of Scientific Research
Bolling Air Force Base, Washington, D.C.

August 1992

VELOCITY AND TEMPERATURE MEASUREMENTS IN A HIGH SWIRL DUMP COMBUSTOR

Dr. Richard D. Gould and Mr. Lance H. Benedict
Mechanical and Aerospace Engineering
North Carolina State University
Raleigh, NC 27695

ABSTRACT

Successful two component laser Doppler velocimetry (LDV) and single point temperature measurements were made in the highly swirling flow field of a model dump combustor. A type S Platinum/10 % Platinum-Rhodium thermocouple probe was used to make temperature measurements. A lean propane-air diffusion flame with an overall equivalence ratio of $\Phi = 0.45$ was stabilized in the combustion chamber by the flow pattern of the Wright Laboratory/Rolls Royce (WL/RR) swirler. The reacting flow case was found to have higher axial and tangential mean velocities than the isothermal flow case throughout most of the flow field due to heat release. A shorter central recirculation bubble was found in the reacting flow case due to the pressure gradients produced by these higher mean velocities. Turbulent normal stresses were found to be larger while the maximum value of the $\overline{u'w'}$ turbulent shear stress was found to be less in the reacting flow case when compared to the isothermal flow case. Maximum mean temperatures occurred where maximum turbulent stresses occurred thus suggesting that gradient transport modeling may be successful in predicting this flow field.

**DEVELOPMENT OF AN ENHANCED
POST RUN DATA ANALYSIS PROGRAM FOR
THE INTEGRATED ELECTROMAGNETIC SYSTEM SIMULATOR
(IESS)**

**Benjamin F. Bohren
Graduate Student
Department of Computer Science**

**The University of North Carolina at Charlotte
Charlotte, NC 28223**

**Final Report for:
Summer Research Program
Avionics Directorate
Wright Laboratory
Aeronautical Systems Center**

**Sponsored by:
Air Force Office of Scientific Research**

July 1992

**DEVELOPMENT OF AN ENHANCED
POST RUN DATA ANALYSIS PROGRAM FOR
THE INTEGRATED ELECTROMAGNETIC SYSTEM SIMULATOR
(IESS)**

**Benjamin F. Bohren
Graduate Student
Department of Computer Science
The University of North Carolina at Charlotte**

Abstract

The ability for engineers to accurately evaluate the output signals generated by the Integrated, Communication, Navigation, and Identification Avionics(ICNIA) Advanced Development Models (ADMs) while in the IESS testing environment is a crucial factor in determining if ICNIA is operating correctly. Before this contract, IESS only generated an immense, difficult to evaluate text file. The enhanced post run data analysis program, called SLICK (Signal Listing IESS Critiquing Knowledge information processor), enables the engineer to select specific output data to be placed in a spreadsheet. Consequently, the engineer is able to graph selected data items; hence turning a chaos of numbers into a meaningful picture.

**DEVELOPMENT OF AN ENHANCED
POST RUN DATA ANALYSIS PROGRAM FOR
THE INTEGRATED ELECTROMAGNETIC SYSTEM SIMULATOR
(IESS)**

Benjamin F. Bohren

Introduction

The ability for engineers to accurately evaluate the output signals generated by the Integrated Communication Navigation and Identification Avionics (ICNIA) while in the IESS testing environment is a crucial factor in determining if ICNIA is operating correctly. ICNIA is a prototype avionics unit for military aircraft. The signals ICNIA generates include Global Positioning System (GPS), Tactical Navigation (TACAN), Identification Friend or Foe (IFF), UHF/VHF/HF radios. ICNIA's test environment, know as IESS, produces an analysis of a scenario execution in the form of a text file, called the post run dump. That file contains information about the signals ICNIA was programmed to produce, called truth data, and about the signals that were actually produced, called collected data. This file's size and layout makes data comparisons virtually impossible.

In an effort to remedy this problem, the engineers began a project, ZBOX, to generate a spreadsheet file by parsing the text file for specified information. The resulting spreadsheet could then be used to graph truth data versus collected data; thereby creating a visual representation of ICNIA's accuracy.

When this contract began ZBOX was partially complete. However, due to a design improvement, only the user's data selection process remains from ZBOX. The design improvement included integrating the spreadsheet file generation with the current text file generation program; thereby eliminating the parsing of the text file, a time demanding and format dependent step. Furthermore, the spreadsheet generation is now independent of the text file and can therefore be created without creating the text file.

The remainder of this document will describe the development process, the expanded analysis capability, how to use SLICK, and the Quattro Pro macros.

Development Process

By using the standard four phase development process of designing, coding, testing, and debugging, the project remained ahead of schedule. The following is a more detailed description of each step.

From a complexity standpoint, the design was straightforward. The task had few options and a definite goal, a spreadsheet containing the values associated with user selected record names. Yet, this is not to imply there were no obstacles. The two most important design decisions were

- 1) To integrate the data collection with the existing IESS post run analysis program. As mentioned in the introduction, this increased the stability of the program by eliminating its dependency on the format of the text file. Additionally, having only one program will save time for both the user and future enhancement programmers.
- 2) To use Quattro Pro as the spreadsheet program. Obviously, this meant the spreadsheet file generated by SLICK had to conform to Quattro Pro importing limitations, including 254 character lines, and empty quotes for missing data. However, by conforming to these restrictions, the data was not ready to use immediately after importing. Therefore, Quattro Pro macros were used to combine lines which had been split when they exceeded 254 characters and to replace empty quotes with nulls.

The coding was originally coded under VMS 5.0. However, the classified machine which runs IESS is limited to VMS 4.7 due to ORACLE code dependent on the older version. Thus minor adjustments were needed when the program was ported to its final destination. Furthermore, conforming to IESS standards on naming conventions and usage of include files added time to this phase.

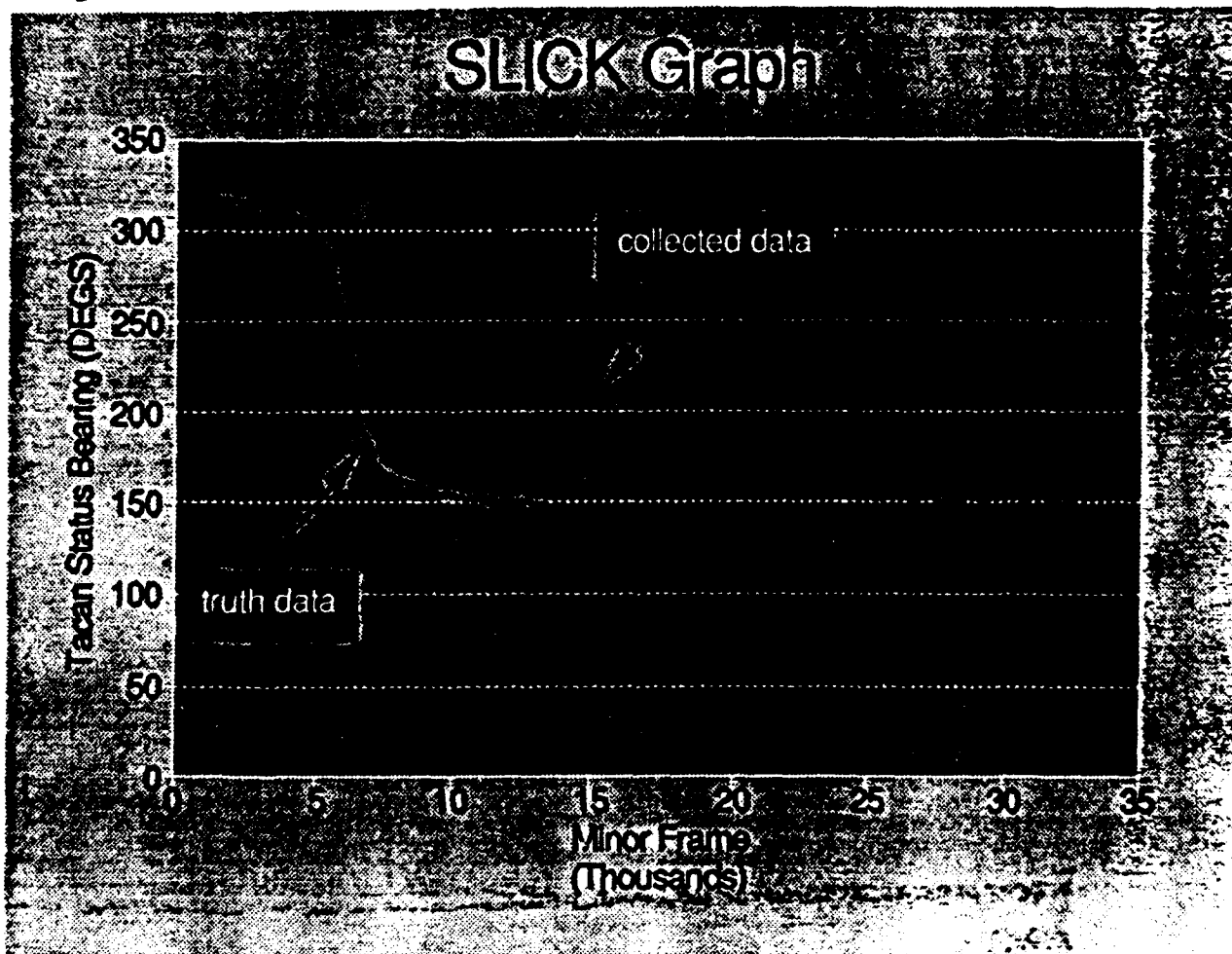
Testing revealed many minor problems due to a lack of knowledge about all the types of data IESS could produce. This led to the fourth stage, debugging, which entailed the usual fine tuning of the code and the clarification of the prompts. Fortunately, this stage was free of any major problems.

New Analysis Capability

The ability to graph the output data generated by IESS will help the engineers analyze ICNIA's capabilities. Below, figure 1, is an example graph made possible by SLICK. To the unknowing eye it is just a couple of lines; however, to the engineer running IESS such a graph will enable him/her to locate where ICNIA is malfunctioning.

In IESS truth data consist of signals IESS sends to ICNIA during a scenario execution and collected data consist of the corresponding signals the ICNIA produced. The minor frame number is IESS's method of keeping time during the run. Knowing this, one can see this graph shows the ICNIA responding late. Having this type of visual comparison will be invaluable to the engineer.

Figure 1



How to use SLICK:

Initiating the post run analysis has not been changed. Simply choose option 5 (post run analysis) from the IESS main menu. This is followed by a series of questions and selection screens, described below. However, not all the screens will appear if a SLICKly formatted file is not being generated. A SLICKly formatted file is either of the new file formats available with SLICK enhancements (see screen 4 for details).

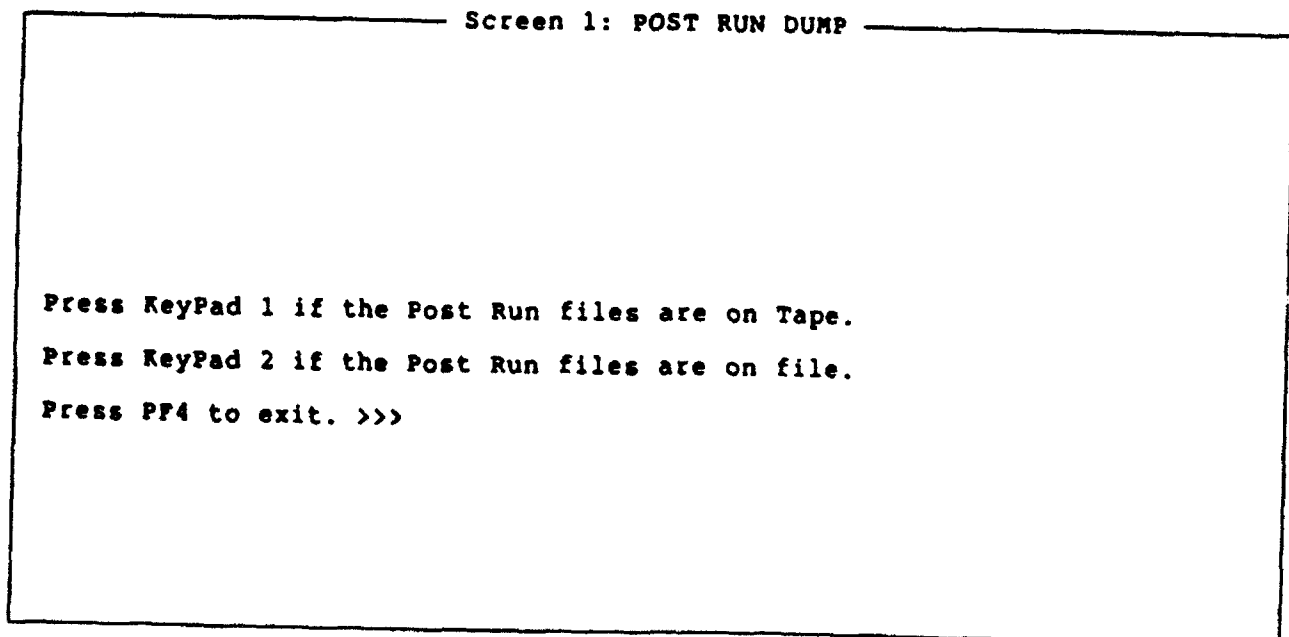
The screens are self explanatory; however, the term keypad should be clarified. Key pad refers to the numeric keys away from the alphabetic characters.

Screen 1:

Description: See figure 2 for actual screen.

Prompts user to find out if the input data files are on tape or disk. Tape refers to tapes which would need to be mounted. The majority of the time, the disk option is chosen.

Figure 2



Screen 2:

Description: See figure 3 for actual screen.

Prompts user to find out if he/she wants to use renamed files from past runs of IESS. The default names given by IESS are listed below. The prefix IESS_DIR_DATA is the logical name where IESS stores it's data files. Currently, this is the RP directory of the account running IESS. These are the names IESS gives to the data files upon each scenario run. The version number is the only unique identifier if the files are not renamed in the last screen of SLICK.

Default Names:

ieess_dir_data:coll_data.dat	The collected data information.
ieess_dir_data:coll_data_fmt.dat	The collected data variable specification records.
ieess_dir_data:header.dat	The header information.
ieess_dir_data:truth.dat	The truth data information and variable specification.

Figure 3

Screen 2: Old/New file Option Menu

Press Key Pad-1 to use old files - renamed files from previous run

Press Key Pad-2 to use files with the default names.

Press PF4 TO EXIT.

>>>

Screen 3:

Description: See figure 4 for actual screen.

Prompting user to find out which output files to generate.

A SLICKly formatted file refers to either of the new file formats available with SLICK (See screen 4 for details).

Option 1 generates the standard text dump that is not optimized for analysis.

Choosing option 2 or 3 will initiate further questions inquiring about the format of the SLICKly formatted file (See screens 4, 5, and 7).

Figure 4

Screen 3: Text/Spread Sheet Option

Press Key Pad-1 To generate text dump - STANDARD POST_RUN_DUMP

Press Key Pad-2 To generate a SLICKly formatted file.

Press Key Pad-3 To generate both types of files.

Press Key PF4 TO EXIT.

>>>

Screen 4:

Description: See figure 5 for actual screen.

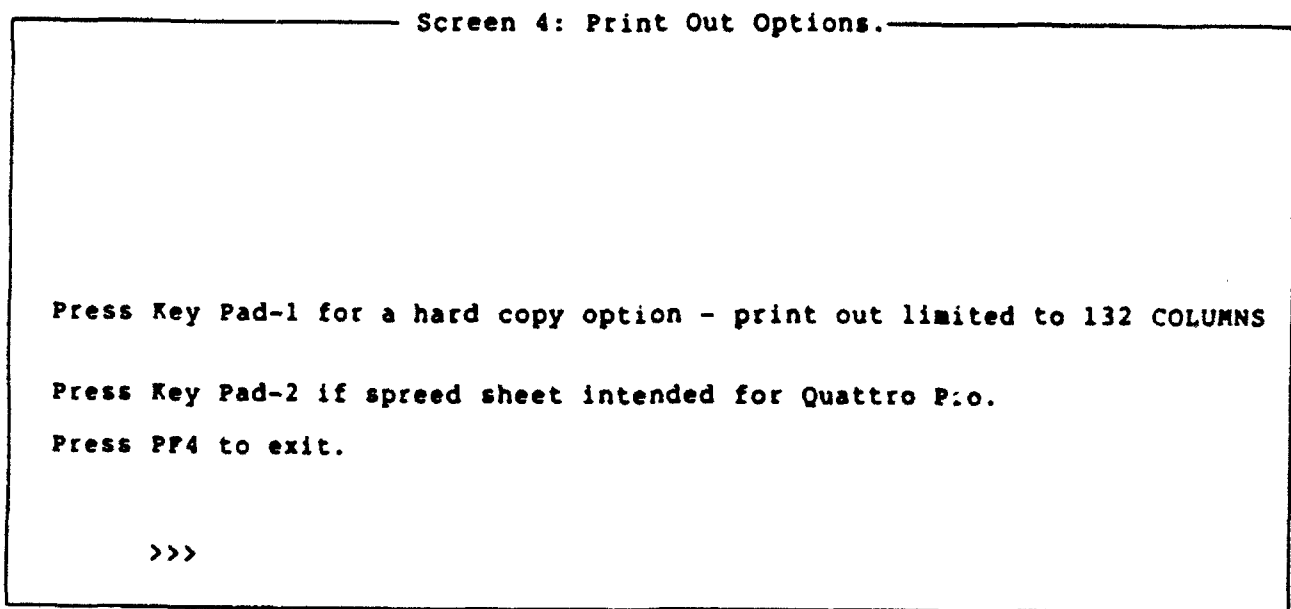
This screen only appears if a SLICK format is to be generated.

Prompts the user to find out if the results are intended for the printer or Quattro Pro.

Option one limits the file to 132 columns and formats the file to be more readable. This file is intended for printing.

Option two formats the data for importing into Quattro Pro. The file is then limited to the 254 columns with longer lines being wrapped. The titles are the most common occurrence of wrapping and Quattro Pro Macros have been written to unwrap them. While wrapped titles are easily identified by the existence of more than one title line, wrapped data lines may not be readily apparent. They can be identified by two or more consecutive lines with the same minor frame number, easily overlooked in the abundance of numbers.

Figure 5



Screen 5:

Description: See figure 6 for actual screen.

This screen only appears if a SLICK format is to be generated.

Prompt user for titles to appear on top of the SLICKly formatted file. File names and/or data type are good things to put here. Press return too leave titles blank.

Figure 6

Screen 5: Spread Sheet Titles

Enter Title of Graph Data >>> example title

Enter Subtitle of Graph Data >>> test

Screen 6:

Description: See figure 7 for actual screen.

Prompt user to find out if he/she wants the truth data to appear in the output files.

Option one means the user wants the truth data to appear in the output file(s).

Option two guarantees truth data will not appear in the output file(s).

Figure 7

Screen 6: Process Truth

Truth data options:

Hit Key-Pad 1 to process truth data.

Hit Key-Pad 2 to skip truth data. >>>

Screen 7:

Description: See figure 8 for actual screen.

This screen only appears if a SLICK format is to be generated.

The selection menu allows the user to select the data records which will appear in the spreadsheet. The menu appears once if truth data is processed and once if collected data is processed. Furthermore, it contains self explanatory usage on the bottom half of the screen.

One point of clarification should be mentioned about the user search. The search string the user enters is a substring that can occur anywhere in the record name. It is the equivalent of using the wildcards as follows: *user input*.

Figure 8

RECORDED DATA (RD) MENU		ID #
SELECT ALL RECORDED DATA	0	
HOST_AIRSPPEED	8	
HOST_ALTITUDE	8	
HOST_HEADING	8	
HOST_LATITUDE	8	
HOST_LONGITUDE	8	

USER KEY PAD LAYOUT :		
Return - Select	7 - Scroll Top	8 - Up One
<---- or		5 - Scroll Middle
Deselect	1 - Scroll Bottom	2 - Down One
PF4 - EXIT	0 - User Search	9 - Page Up
		3 - Page Down
		. - Done

Screen 8:

Description: See figure 9 for actual screen.

Prompts user to find out if he/she wants the collected data to appear in the output files.

Option one means the user wants collected data to appear in the output file(s). This will cause screen 7 to appear with the list of collected data.

Option two guarantees collected data will not appear in the output file(s).

Figure 9

Screen 8: Search of Collected Data option.

Completed Truth Data portion.
TRUTH DATA written to the output file/files.

Collected IEEE-488 and UUT Data Options:
Press Key Pad-1 to process Collected Data.
Press Key Pad-2 to skip Collected Data.

>>>

Screen 9:

Description: See figure 10 for actual screen.

Displays messages informing user of which files were generated. Then the user is prompted with the option of renaming the input files. If the user chooses to rename the files another screen will ask for the new names. This is followed by the same process for renaming the output files. In the example below the user chose not to rename the files.

This option is given so the next run of IESS will not overwrite the files.

Figure 10

Screen 9: Rename Files

SELECTED RECORDS WRITTEN TO SPREAD.TAB

PRESS KEYPAD 1 TO RENAME INPUT FILES, ANY OTHER KEY TO CONTINUE.
>>>

PRESS KEYPAD 1 TO RENAME OUTPUT FILES, ANY OTHER KEY TO CONTINUE.
>>>

The Quattro Pro Macros

As mentioned in the design decisions, the Quattro Pro importing restrictions caused a need for macros to prepare spreadsheets for analysis. Three macros have been prepared for various needs described below. Each is a series of Quattro Pro commands and is not intelligent enough to know when it needs to be executed. Thus, the user must know how they work and when to use them if they are going to be effective.

To use the macro library simply open the file containing the macros, currently `maclib.wq1`. Then import the SLICKly formatted file into a new Quattro Pro worksheet and call the macros.

Macros:

1) unwrap

Why needed: Quattro Pro only allows 254 characters per line when importing. Therefore, SLICK splits longer lines into two.

Description: unwraps the titles associated with the collected data. This is accomplished by moving lines 3 and 4 (where line 1 is the first title line after the heading 'collected data') onto the end of lines 1 and 2. If there are more than 4 title lines, repeat the macro until all titles are on the same line. A single title line consist of the titles on line 1 and their units on line 2.

2) unwrap_truth

Why needed: Quattro Pro only allows 254 characters per line when importing. Therefore, SLICK splits longer lines into 2.

Description: unwraps the titles associated with the truth data. This is accomplished by moving lines 3 and 4 (where line 1 is the first title line after the heading 'truth data') onto the end of lines 1 and 2. If there are more than 4 title lines, repeat the macro until all titles are on the same line.

3) nulls

Why needed: Quattro Pro has no method of importing missing values. Thus, SLICK uses a quote followed by a single space whenever data is missing. However, when graphed this causes spikes wherever the data is missing.

Description: Replace all quotes followed by a space with a null value. This macro should be run whenever there is missing data values.

Conclusion

Although this program is called the summer *research* program, this particular project was entirely development oriented. To administrators reading this, do not take that the wrong way. I still learned a lot while I was here, but had expected a different type of work. Thus, for a conclusion I can only offer the expanded analysis capabilities and the answers it will help find in the future.

I would like to take this opportunity to point out problems I observed while working in the lab. First, there is a high turn over of technical expertise. Second, the engineers spend a lot of time doing diagnostic work on ICNIA's hardware. While these two facts alone would not constitute building an expert system to aid in diagnosing these problems, I believe an expert system could be beneficial in persevering knowledge over time and could reduce *diagnostic time when problems occur*.

Two questions must first be answered before seriously considering such a project. First, could the engineers knowledge be transformed into rules? Second, is there an expert who can not only solve the problems efficiently, but be able to explain how he solved them? If both of these can be answered yes, then an expert system is potentially valuable for this problem and would warrant further research.

HARD TARGET CODE ASSESSMENT
AND A QUALITATIVE STUDY OF SLIDE LINE
EFFECTS IN EPIC HYDROCODE

Thomas C. Byron
Graduate Student
Aeronautical engineering, Mechanics and Engineering Sciences

University of Florida
231 Aero Building
Gainesville FL, 32611

Final Report for:
Summer Research Program
Wright Laboratory

Sponsored by:
Air Force Office of Scientific Research
Bolling Air Force Base, Washington, D.C.

August 1992

HARD TARGET CODE ASSESSMENT
AND A QUALITATIVE STUDY OF SLIDE LINE
EFFECTS IN EPIC HYDROCODE

Thomas C. Byron
Graduate Student
Aeronautical engineering, Mechanics and Engineering Sciences
University of Florida

ABSTRACT

As part of an ongoing Wright Laboratories experimental program to investigate the validity of hydrocodes for massive concrete structures subjected to an internal explosion, the problem was modeled and then analyzed using the research version of EPIC91. Two different models with substantially different mesh shapes and sizes were generated to study the effects on the computed results as well as the required CPU time. The concrete and explosive material properties were obtained from the existing EPIC materials library.

A study was also conducted into the effects of slide lines versus no slide lines in EPIC calculations where two or more materials of vastly different densities come into contact with one another. This was accomplished using a simple one dimensional impact setup to analyze the pressure wave as it moved through the model.

HARD TARGET CODE ASSESSMENT
AND A QUALITATIVE STUDY OF SLIDE LINE
EFFECTS IN EPIC HYDROCODE

Thomas C. Byron

INTRODUCTION

USAF Wright Laboratory is conducting an experimental program to investigate the validity of hydrocode results for massive concrete structures subjected to an internal explosion (ref 1). This simulation of a warhead detonation following penetration is also being used to generate improved engineering models for future research and development.

The experimental models, constructed by Denver Research Institute (DRI), are solid concrete cylinders measuring six feet tall with a four foot radius. There is a three inch diameter core from the top to the center of the cylinder. A one pound sphere of explosive is placed at the geometric center of the cylinder and the remainder of the core is packed with moist sand. Various guages are preset in the cylinder prior to the concrete pour. The cured concrete cylinders are kept in their galvanized steel culvert forms for the test. See Fig 1.

METHODOLOGY

Two distinctly different models were constructed using

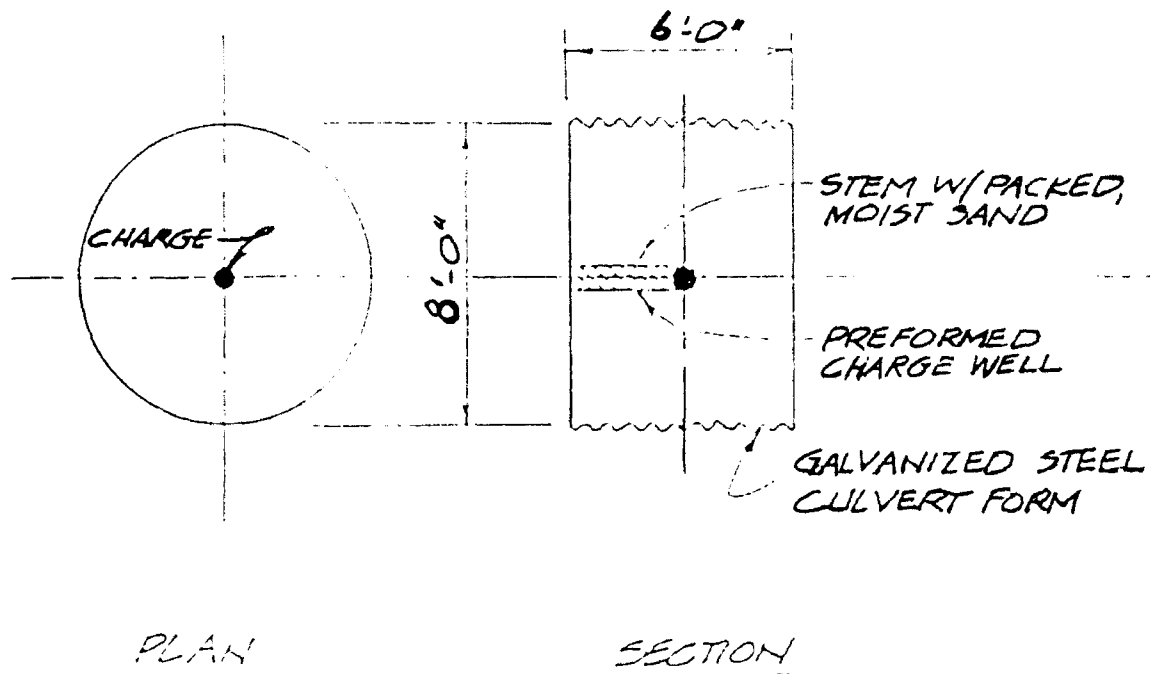


Fig 1. Plan and section view of test cylinders from
DRI "TEST PLANNING STATUS AND REVISIONS"

PATRAN (ref 2). These were then input into the research version of EPIC91 (ref 3). One of the many benefits of PATRAN is that it will write the data from a generated model into a neutral file which can be read directly into EPIC. In both cases, the core of moist packed sand was modeled as concrete. This allowed for the elimination of slide lines and for a much simpler geometric model. Use of symmetry and axisymmetry reduced the models to one quarter of the section view of the cylinder (Fig 2).

In both cases, the explosive and the concrete material

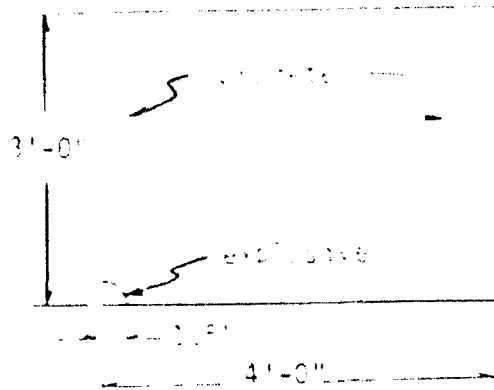
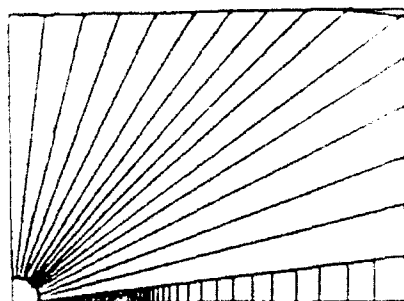


Fig 2. Model dimensions

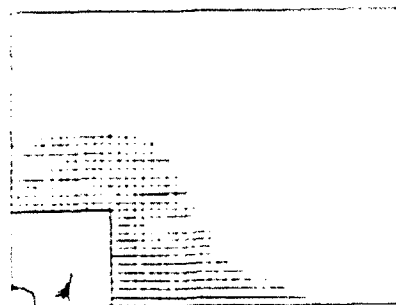
properties were taken from the materials library of EPIC. Also in both cases, the mesh for the explosive region is identical. This region consists of 64 triangular elements in a quarter circle area of 1.9 sq in.

Since it is the effects in the concrete that are of interest and since the explosive material will have little effect on the concrete after detonation, the boundary between the explosive and the concrete was handled without slide lines, further simplifying the model.

The near radial symmetry of the model and the expected radial symmetry of the post detonation waves through the model were taken into account for Model 1. The mesh was generated by dividing the concrete area into 16 wedges of equal angle from the center point and eliminating the area



typical, not to scale



triangle elements

Fig 3. Model 1

Fig 4. Model 2

Each four sided PATRAN element is transformed to 4 crossed triangles by EPIC.

for the explosive (Fig 3). This allows for a very tight mesh in close to the center to eliminate any impedance mismatch with the small triangular explosive elements. It also allows for increasingly large mesh as the distance from the center increases and the effects of the explosion are diminished. All but one of the elements in the concrete region were constructed using PATRAN's quad\5 elements. The quad\5 is simply a four sided element with a center node. EPIC transforms the quad\5 into four crossed triangle elements when the PATRAN neutral file is read into EPIC. The final element in the far corner is a single triangular element due to the fact that the model is rectangular vice square. This model resulted in 2401 elements.

Model 2 was constructed with the intent of looking at the bounce back of the pressure wave off the free surfaces of the cylinder. With that objective in mind, the element size was kept small throughout the entire region. From the quarter circle region for the explosive, there is a triangular element transition area out to six inches with the same size elements as the explosive. From there, the elements are half inch square quad\5's which are again transformed to crossed triangles by EPIC (Fig 4). These elements are roughly the same size as the smallest elements in Model 1. This type construction resulted in over 27,000 elements.

Model 1, due to its relatively small size, was run on a SiliconGraphics IRIS 4D\340GTX. The problem used just over six hours of CPU time. Model 2, due to its relatively large size, was run on a CRAY Y-MP supercomputer. It took approximately 1700 seconds of CPU time which equates to roughly 14-15 hours of SiliconGraphics time.

RESULTS

The data from the numerical analysis were dumped for every 10 microseconds of time step in the computation. These data dumps were written to PATRAN neutral files. Once all the neutral files have been collected, they can be

displayed and saved individually. The saved screens are then put together in a movie type presentation which allows one to "see" the pressure wave or the desired output wave move through the model after detonation. This is a very powerful tool to help in understanding what goes on in a complex system. In this case, it can also be very instrumental for comparison purposes by running the movies of both models simultaneously on the screen and stepping through frame-by-frame.

In comparing the output of both models, it can be seen that the results are virtually identical for the first 90 milliseconds after detonation. By this time, the initial wave front is very near the free surface and has begun to separate into a distinct wave. For Model 1, as the leading wave continues to separate and strike the free surface, noise begins to be generated in the results where the elements tend to be much larger. The bounce back anticipated in Model 2 is readily visible. The bounce back from Model 1 is also visible but, for the most part, is lost in the generated noise.

The results will eventually be compared with the experimental data collected by DRI to assess the validity of the models and the assumptions made. Adjustments and

corrections will be made after that analysis.

CONCLUSIONS

It is readily apparent that if the failure of the concrete is the only concern, then Model 1 is quite sufficient for analyzing that region at a significantly reduced cost in CPU time. However, if bounce back of the wave is of concern, which it is for this problem, then the CPU price will have to be paid by using Model 2 since any useable information gets lost in the generated noise for Model 1 away from the center, near the free surfaces.

Actual assessment of the validity of the models, as stated earlier, is a future project.

A QUALITATIVE STUDY OF SLIDE LINE EFFECTS IN EPIC

INTRODUCTION

Often a soft, low density material is used to buffer a hard, high density material from shock waves and vibration. A particular application of this is the placing of a thin plastic layer between two thick ceramic armor plates. A problem of this nature was run with EPIC giving some unexpected results. The question asked was how do slide lines, or the lack of slide lines, effect the results of EPIC calculations for this type of problem.

METHODOLOGY

To get a feel for the effects of slide lines on the results of a problem of this type, a simple one-dimensional model was used. The model consisted of a projectile and target type setup. The projectile was given the same initial velocity for all tests, even those where it was physically attached to the target (no slide lines between target and projectile).

The target was made up as a sandwich type composite. The high density material used was OFHC copper found in the existing EPIC material library. The low density material has the exact same properties as the OFHC copper except that

the density has been changed to one tenth its original value. The sandwich consisted of two long pieces on the outside and a short piece in between. Both long pieces were of equal length and four times longer than the short piece. Several variations to the composite were examined. In all cases, the outside pieces of the sandwich are OFHC copper. The variations to the center piece included: 1) using the same material as the outside pieces but reducing element size by a factor of ten so that the total mass of each center element is one tenth that of the outer pieces; 2) using the low density copper with the same element size all the way through the model; 3) using low density copper with an element size 10 times that of the outer pieces to equate masses in all the elements; 4) using copper all the way through. All of these variations were run with and without slide lines between sandwich components and also between the projectile and the composite target.

RESULTS

The only variation which showed no difference in results between slide lines and no slide lines is the all copper model. All other variations, using low density copper, large element size changes or both, showed differences in the results. The result differences were not in the overall shape of an output curve over the length of

the model but amounted to very localized spikes. For instance, the pressure curve for each variation of the model is virtually identical over time for both slide line and no slide line calculations except for those spikes. The slide line results in all cases gave spikes in the output at any material interface in the model as well as occasional spikes throughout the model. These spikes do not represent any physical phenomenon and are simply a result of the numerical calculations.

CONCLUSIONS

The option of eliminating all slide lines in EPIC is not always available, e.g. in the case of large deformation problems, however, the results of this simple test would indicate that elimination of slide lines where ever possible would be advantageous. Removing slide lines should prove even more critical for higher dimension problems as the numerical analysis becomes more involved and the iterative process becomes much larger.

REFERENCES

1. Hard Target Assessment Code
Contract F08630-91-C-005
2. PATRAN, a division of PDA Engineering, is a 3-D
Mechanical Computer-Aided Engineering software system.
3. EPIC research code, 1991 version - a computer program
for impact and explosive detonation computations in 1,
2, or 3 dimensions. Authors: Gordon Johnson and Robert
Stryk, Alliant Techsystems Inc.

LASER IMAGING AND RANGING (LIMAR) PROCESSING

Jack S.N. Jean
Assistant Professor
Department of Computer Science and Engineering
Wright State University
Dayton, Ohio 45435

Louis A. Tamburino
Avionics Directorate
Wright Laboratory
Wright-Patterson AFB, Ohio 45433

Ahmet A. Coker
Department of Computer Science and Engineering
Wright State University
Dayton, Ohio 45435

Final Report for:
Summer Research Program
Wright Laboratory

Sponsored by:
Air Force Office of Scientific Research
Bolling Air Force Base, Washington, D.C.

September 1992

LASER IMAGING AND RANGING (LIMAR) PROCESSING

Jack S.N. Jean
Assistant Professor
Department of Computer Science and Engineering
Wright State University

Louis A. Tamburino
Avionics Directorate
Wright Laboratory

Ahmet A. Coker
Department of Computer Science and Engineering
Wright State University

Abstract

The LIMAR (Laser IMaging and Ranging) project is a Wright Laboratory effort to develop an advanced imaging and ranging system for robotics and computer vision applications. LIMAR embodies a concept for the fastest possible three-dimensional camera. It eliminates the conventional scanning processes by producing a registered pair of range and intensity images with data collected from two video cameras. The initial prototype system was assembled and successfully tested at Wright Laboratory's Avionics Directorate in 1992. This prototype LIMAR system used several frame grabbers to capture the demodulated LIMAR image signals from which the range and intensity images were subsequently computed on a general purpose computer. The prototype software did not address the errors which are introduced by differential camera gain, misalignment, and distortion. The tasks performed during this Summer Research Program include (1) modeling and developing algorithms to correct the distortion introduced by using two cameras and (2) design of special purpose hardware to convert, in real-time, the outputs from the two cameras into a fully registered range and intensity image.

NEURAL ON-LINE LEARNING IN MISSILE GUIDANCE

Jeffrey S. Dalton
Graduate Student
Department of Electrical Engineering

Advisor:
S. N. Balakrishnan
Associate Professor
Department of Mechanical and Aerospace Engineering and Engineering Mechanics

University of Missouri-Rolla
Rolla, MO 65401

Final Report for:
Summer Research Program
Wright Laboratory

Sponsored by:
Air Force Office of Scientific Research
Eglin Air Force Base, FL

September 1992

NEURAL ON-LINE LEARNING IN MISSILE GUIDANCE

Jeffrey S. Dalton
Graduate Student
Department of Electrical Engineering
University of Missouri-Rolla

Advisor:
S. N. Balakrishnan
Professor
Department of Mechanical and Aerospace Engineering and Engineering Mechanics
University of Missouri-Rolla

Abstract

In this work we investigate the use of neural networks in providing control signals to solve the target intercept problem. The approach taken here is based on an architecture that contains an adaptive critic network which evaluates previous control actions and produces a complementary control to counteract target acceleration. In previous work we used a linear optimal control law to produce the primary missile command accelerations. In this work we replace the optimal control law with a neural network approximation of the optimal control law and modify network weights on-line to react to target acceleration. We show a series of simulations which compare current results with those obtained previously. Results of this study are encouraging, however, they show that proper network training is a key issue. Analysis of the proposed control system with respect to stability, convergence, and robustness remain to be done and further work is in progress.

NEURAL ON-LINE LEARNING IN MISSILE GUIDANCE

Jeffrey S. Dalton and S. N. Balakrishnan

INTRODUCTION

The problem under investigation in this work is the intercept problem, that is the determination of proper control signals which guide a missile to the point of intercepting a target. Numerous approaches to the solution of this problem exist ranging from line-of-sight, pursuit, and proportional navigation to more advanced techniques based on optimal control theory and the theory of differential games [1]. One challenge in all of these techniques is in detecting and responding to unknown target accelerations.

Many researchers have studied the use of neural networks in control systems. Introductory material for the use of neural networks in controls may be found in [2]. Our approach to the intercept problem is based on the use of neural networks. The architecture described in this report is similar to that used in our own previous work [3]. The control applied to the plant is decomposed into two components. The first component is taken from an optimal control law which assumes that the target is not accelerating. The second component is taken from a critic network that monitors plant inputs and outputs to detect and respond to target accelerations. There are two primary contributions in this work. First, we demonstrate that a feed-forward neural network controller can be used to approximate the optimal control law in the control system. Second, we show a method of performing on-line adaptation of the controller to respond to target accelerations in flight.

The remainder of this report proceeds as follows. First we present the methodology used in our approach to the intercept problem. We discuss the function of each of the blocks of the proposed control system and describe the procedures used in both on-line and off line training of the neural networks. Next a series of simulations are presented which demonstrate an implementation of the method. Finally, we discuss the results obtained and describe work that is in progress.

METHODOLOGY

Figure 1 shows a block diagram of the control system to be considered. As shown, there are three primary components to the system. The plant implements the dynamics of the intercept problem. The controller network provides the primary missile acceleration commands to the plant. The critic network observes control commands and plant responses to detect target acceleration and produces a complementary control signal. The details of the blocks and interactions between them are discussed in the remainder of this section. We also describe training procedures for

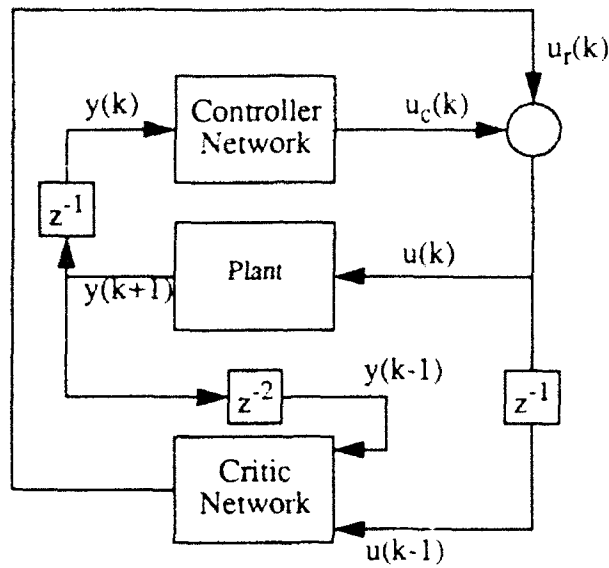


Figure 1. Control System Block Diagram

both the controller and critic blocks.

The plant for the intercept problem is modeled using discrete time state equations. States for the plant are relative positions, relative velocities, and target accelerations each in two directions:

$$X = \begin{bmatrix} x & y & \dot{x} & \dot{y} & a_{tx} & a_{ty} \end{bmatrix}^T \quad (1)$$

The state transition matrix is given by:

$$\Phi(k, k+1) = \begin{bmatrix} 1 & 0 & \Delta T & 0 & \frac{1}{2}\Delta T^2 & 0 \\ 0 & 1 & 0 & \Delta T & 0 & \frac{1}{2}\Delta T^2 \\ 0 & 0 & 1 & 0 & \Delta T & 0 \\ 0 & 0 & 0 & 1 & 0 & \Delta T \\ 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 \end{bmatrix} \quad (2)$$

There are two plant input signals, the missile commanded accelerations in each direction and the input matrix is given by:

$$\Gamma(k) = \begin{bmatrix} -\frac{1}{2}\Delta T^2 & 0 & -\Delta T & 0 & 0 & 0 \\ 0 & -\frac{1}{2}\Delta T^2 & 0 & -\Delta T & 0 & 0 \end{bmatrix}^T \quad (3)$$

In our work we assume that the relative positions and velocities are perfectly known to the missile control system but that the target acceleration states are not known. Thus the plant output matrix is given by

$$C = [I_{4 \times 4} \quad 0_{4 \times 2}] \quad (4)$$

where I and O are appropriately sized identity and zero matrices, respectively. The plant output vector is

$$Y = [x \ y \ \dot{x} \ \dot{y}]^T \quad (5)$$

In the plant model of equations (1)-(5) for the intercept problem constant target accelerations are propagated based on their initial values. In a real engagement a target maneuver would result in changes in the direction of the velocity vector but the magnitude of velocity would remain constant. In our model of the engagement, on the other hand, velocity magnitude does change in response to target accelerations. In a sense, this makes our problem more difficult.

The controller network in the block diagram of Figure 1 marks a point of departure from previous work. Previously the control signal u , was calculated using an optimal control law from [4] which is based on the performance index

$$J = \frac{1}{2} Y_f^T S_f Y_f + \frac{\gamma}{2} \int_{t_0}^{t_f} u^T u dt \quad (6)$$

where

$$S_f = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix}, \quad (7)$$

$\gamma = 10^{-4}$ is the control weighting, and u is the missile command acceleration. In the solution of the optimal control problem we assume that the target acceleration states internal to the model are zero and that the controller is the only source of control in the system. The optimal guidance law is also based on an estimate of time-to-go calculated by assuming that the relative range rate is constant.

$$T_{go} = \frac{|R|^2}{V \cdot R} \quad (8)$$

where R is the relative range and V is relative velocity. The optimal missile command accelerations are then obtained

by state feedback with a time varying gain matrix,

$$C(t) = \begin{bmatrix} C_1 & 0 & C_2 & 0 \\ 0 & C_1 & 0 & C_2 \end{bmatrix} \quad (9)$$

where

$$C_1(t) = \frac{N(t)}{T_{go}^2} \quad (10)$$

$$C_2(t) = \frac{N(t)}{T_{go}} \quad (11)$$

and

$$N(t) = \frac{3T_{go}^3}{(3\gamma + T_{go}^3)} \quad (12)$$

The control is then calculated by

$$u(t) = C(t)Y(t). \quad (13)$$

When the target maneuvers the control produced by this method must be complemented by an additional control signal which offsets the target acceleration.

In the current work we replace the optimal control law with a neural network controller trained to produce the mapping (13). We make this replacement to facilitate on-line adaptation of the controller in the presence of unknown target acceleration. The actual training data for the controller network was taken from plant inputs and outputs along a specific trajectory with no target acceleration present. The network was trained using the backward error propagation method of [5]. After training, the controller network approximates the mapping of relative positions and velocities to missile commanded accelerations to be fed into the plant.

In the presence of unknown target accelerations it is necessary to add an additional control signal which counteracts the target acceleration. This additional control is provided by the critic network. The function of the critic is to observe previous plant inputs and outputs and evaluate the performance of the controller. Based on this evaluation a correction signal is provided to complement the action of the controller.

The critic, like the controller network is trained on data obtained from a nominal simulation in which there is no target acceleration. The network is presented plant output vectors and the corresponding optimal control signals, i.e. $y(k)$ and $u_c(k)$. It is trained to produce the vector $[0 \ 0]^T$ for each point along the trajectory. The motivation for this

type of training comes from the fact that the derivative of the Hamiltonian with respect to plant inputs should be zero along the optimal trajectory. Thus we are asking the critic network to assess the correspondence between plant outputs and controller commands. When there is a miss match the critic produces a signal which is amplified to produce the complementary control, $u_c(k)$. Further details of the operation of the critic network may be found beginning on page 55 of [3].

The critic network block also contains an internal model of the plant which is used to predict current plant outputs based on the previous output and current control. This prediction is compared with the actual output of the plant. When the norm of the difference between actual and predicted output exceeds a threshold the critic network signal is switched on-line. The switching logic reduces the sensitivity of the control system to errors in the critic network when operating close to the nominal trajectory.

During the course of a simulation we wish to adapt the controller network so that it learns to produce the complementary control that is being provided by the critic network. The procedure used to do this is straight forward. We use the control signal, $u_c(k)$, produced by the critic network at each sample interval as an error signal for the controller. This error is back propagated through the controller network to modify the controller's weights. For a given time step, k , the training proceeds until either a maximum number of iterations is exceeded or until the following error criteria is satisfied

$$\left(\left[u(k) - u_c(k) \right]^T \left[u(k) - u_c(k) \right] \right)^{\frac{1}{2}} < \epsilon \quad (14)$$

where ϵ is a preselected small positive constant.

In the following section we present results of an implementation of the control system which has been described and compare these results with those obtained previously using the optimal control law.

RESULTS

In this section we discuss experimental results. A series of four simulations are described. In the first three cases we compare results obtained using the neural network controller with those resulting from use of the optimal control law. The fourth simulation illustrates on-line learning in the controller network.

Four types of graphs are shown in the figures. These are plant outputs, command accelerations, differences between signals in the neural network control system and the system which uses the optimal control law, and training error curves. For graphs showing plant outputs or differences in the outputs solid lines and dashed lines represent relative position in the x and y directions respectively. Dot-dash and dotted lines show relative velocities in the x and y

directions respectively. Control signals resulting from the controller, the critic, or both and respective differences are shown with solid lines representing control in the x direction and dashed lines representing y direction commands. Differences between configurations are calculated by subtracting signals in the network controller configuration from corresponding signals in the optimal control law configuration. In all simulations the sampling interval is 0.2 seconds.

Our first goal is to compare the use of the neural network controller with that of the optimal control law for the case when there is no target acceleration. In this case the critic network is not needed and is removed from the control loop. We expect a close correspondence between the two configurations since data from the optimal control law configuration was used to train the neural network controller. Results are shown in Figure 2 for initial conditions

$$x_0 = [1000 \text{ ft. } 200 \text{ ft. } -100 \text{ ft./s } -50 \text{ ft./s } 0 \text{ ft./s}^2 \ 0 \text{ ft./s}^2]^T \quad (15)$$

The top row of graphs compare plant output states for the two configurations. As expected there is a close correspondence between the two configurations with the largest difference being approximately 0.6 feet in the y direction of relative position. The bottom row of the graph shows that the optimal control law and network controller are also closely matched.

In the second set of simulations we added target acceleration in the x direction with the initial states

$$x_0 = [1000 \text{ ft. } 200 \text{ ft. } -100 \text{ ft./s } -50 \text{ ft./s } 5 \text{ ft./s}^2 \ 0 \text{ ft./s}^2]^T \quad (16)$$

As before, the critic network remains out of the control loop. The purpose of this simulation is to examine controller network performance along a trajectory different from that used to train the network.

Figures 3 and 4 show a comparison between optimal control law and controller network configurations. Figures 3a-3c show comparisons in the plant output states. Despite the presence of target acceleration, the optimal control law configuration reaches the point of intercept at $k=74$ with a negligible miss distance. The network controller configuration, on the other hand, misses by 33.5 feet. The simulation terminates when the relative range rate changes sign at $k=73$. As shown in Figure 3c the miss is largely due to a difference in the y-component of relative position. Figures 3d-3f illustrate the reason for the deviation. The network controller produces less control in the y-direction.

We also examine critic network output signals even though they do not contribute to the total plant input. Recall that the critic network produces supplemental control based on plant inputs and outputs delayed by a single time step. Figures 4a-4c show a comparison of critic network output signals. The large difference of Figure 4c illustrates that the critic is very sensitive to changes in its inputs.

Using the initial states of (16) we consider the case where the critic is allowed to contribute to the total plant input. Figures 5a-5c show plant outputs and differences between the two configurations. Figures 5d-5f show plant inputs and their differences. The optimal control law configuration reaches intercept in 53 time steps with negligible miss. The controller network configuration terminates after 56 time steps with a miss distance of 27.8 feet. The plant

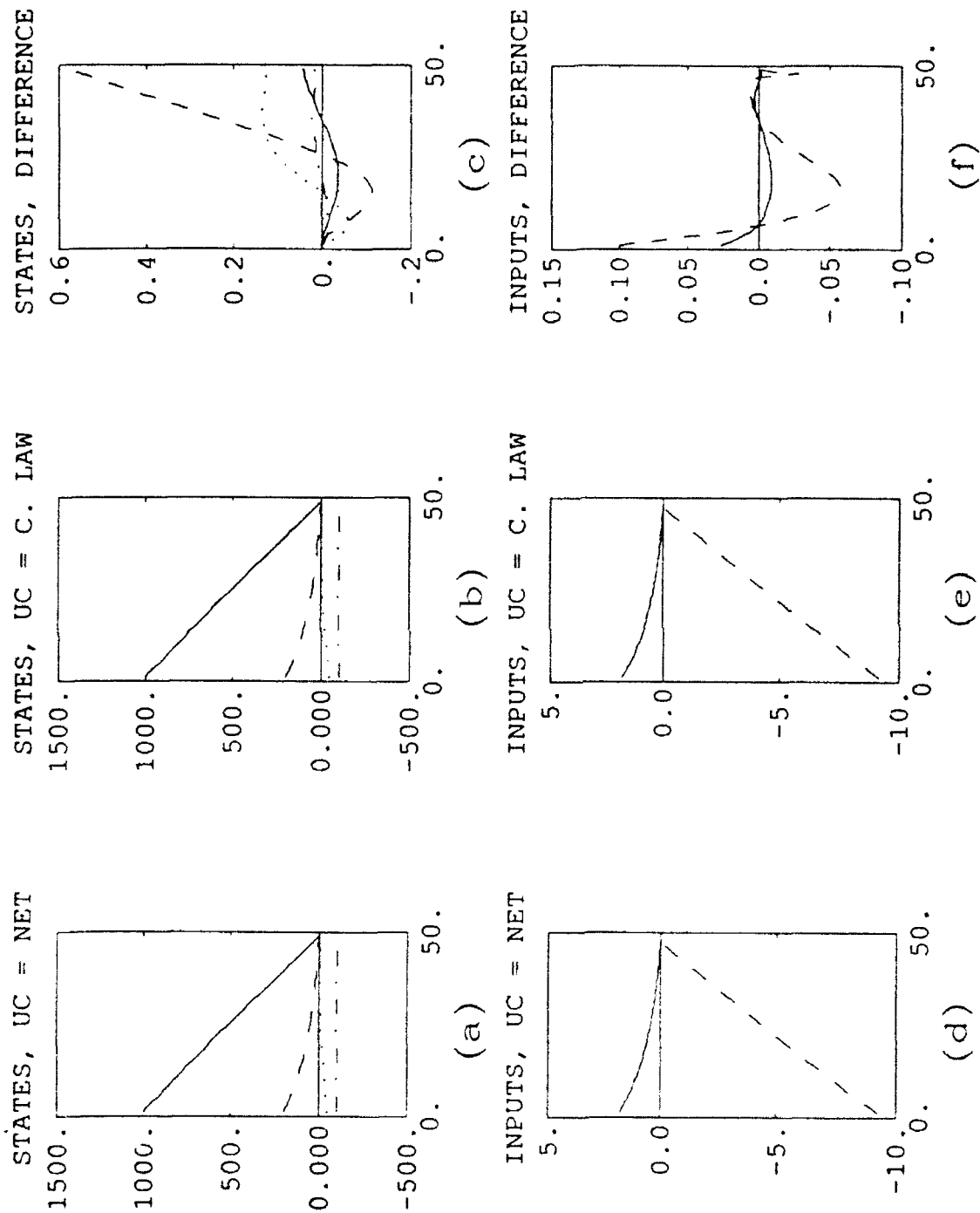


Figure 2. Comparison between optimal control law and neural network controller configurations when there is no target acceleration and no on-line learning.

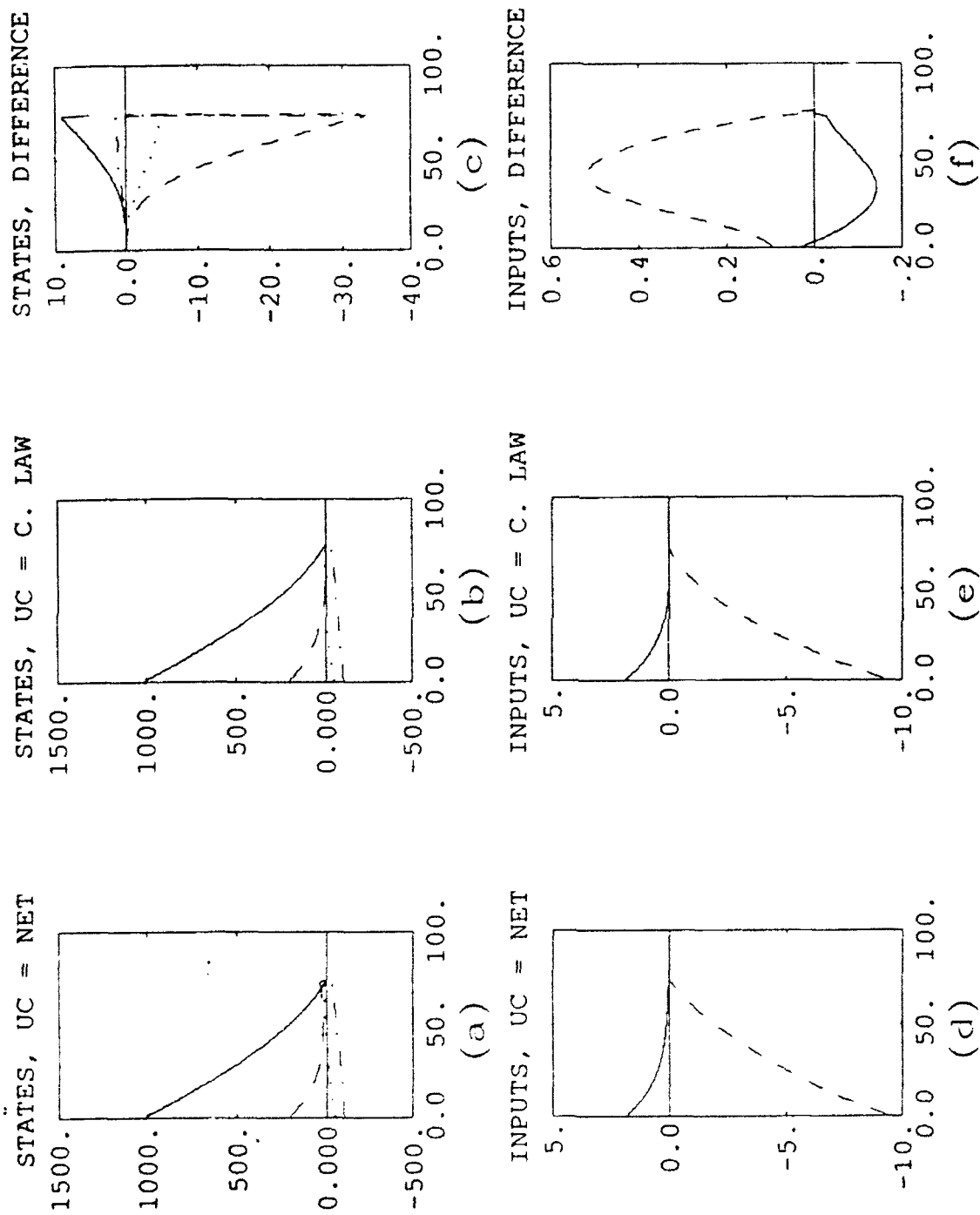


Figure 3. Comparison between optimal control law and neural network controller configurations in the presence of target acceleration with the critic out of the control loop and no on-line learning - Plant inputs and outputs.

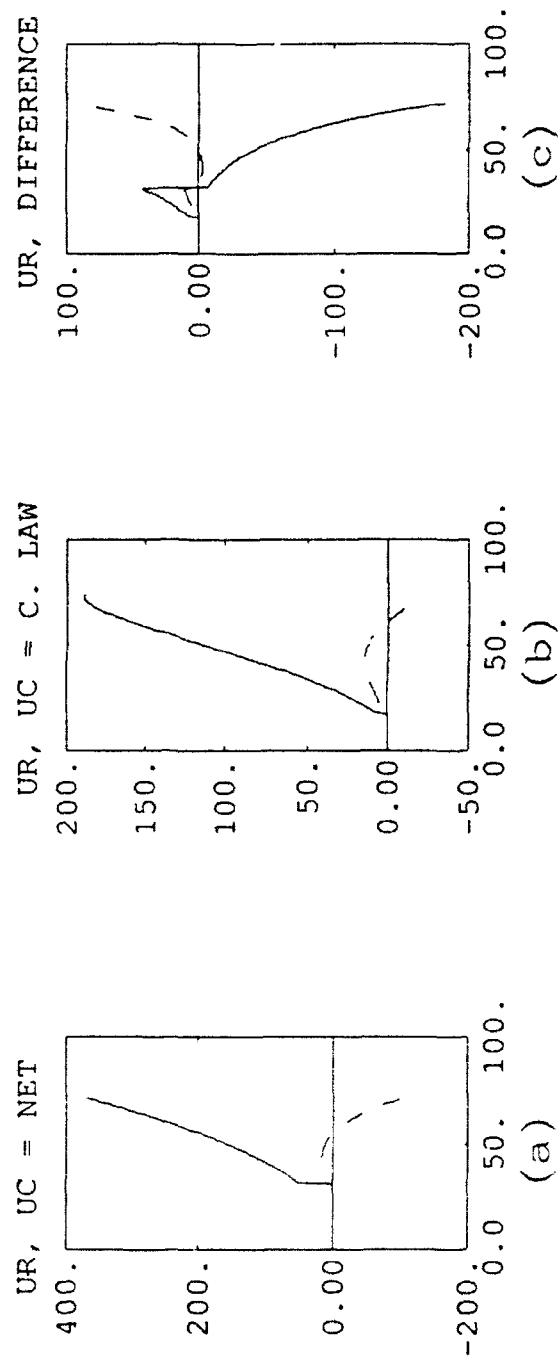


Figure 4. Comparison between optimal control law and neural network controller configurations in the presence of target acceleration with the critic out of the control loop and no on-line learning - controller outputs.

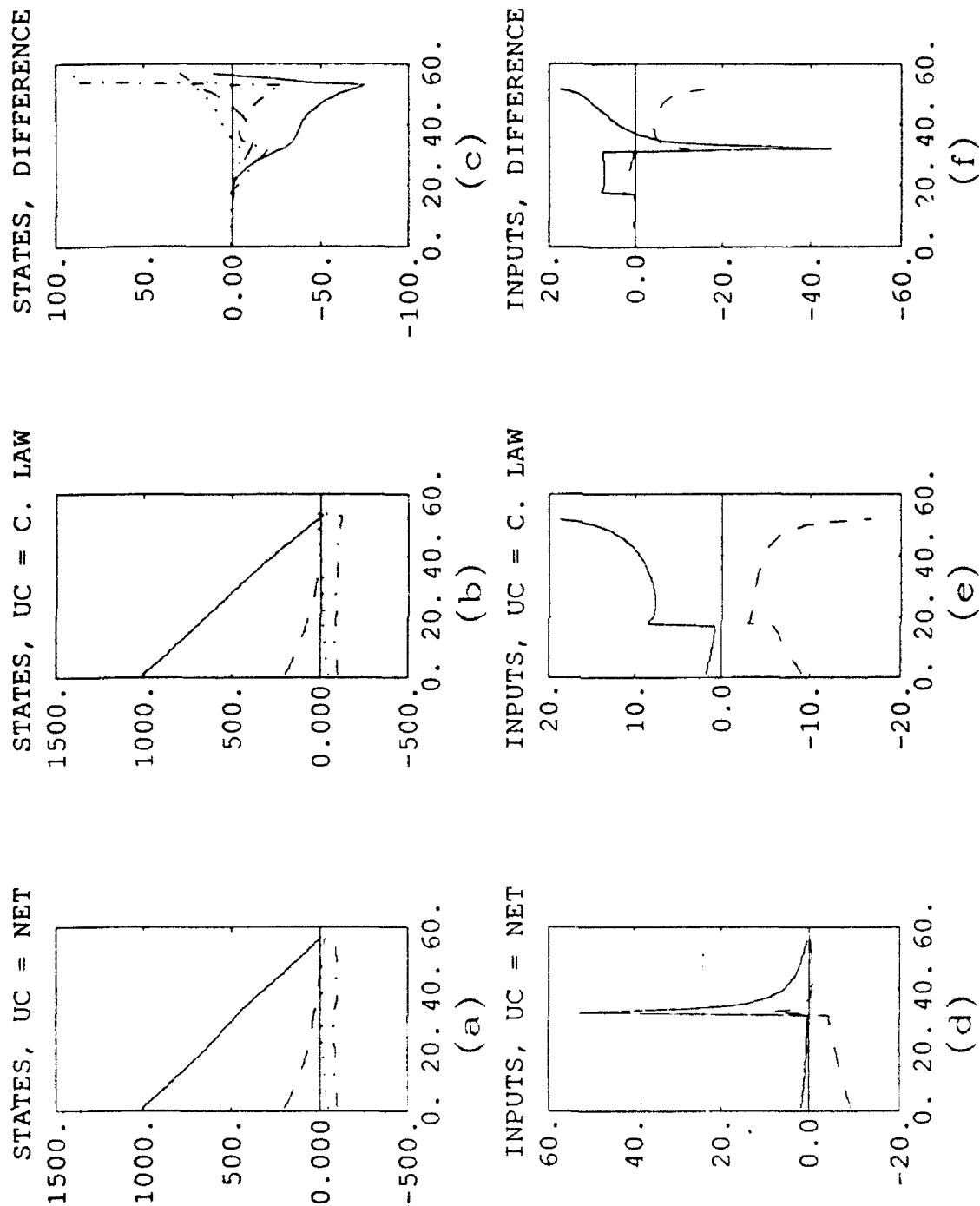


Figure 5. Comparison between optimal control law and neural network controller configurations in the presence of target acceleration with the critic in the control loop and no on-line learning - plant inputs and outputs.

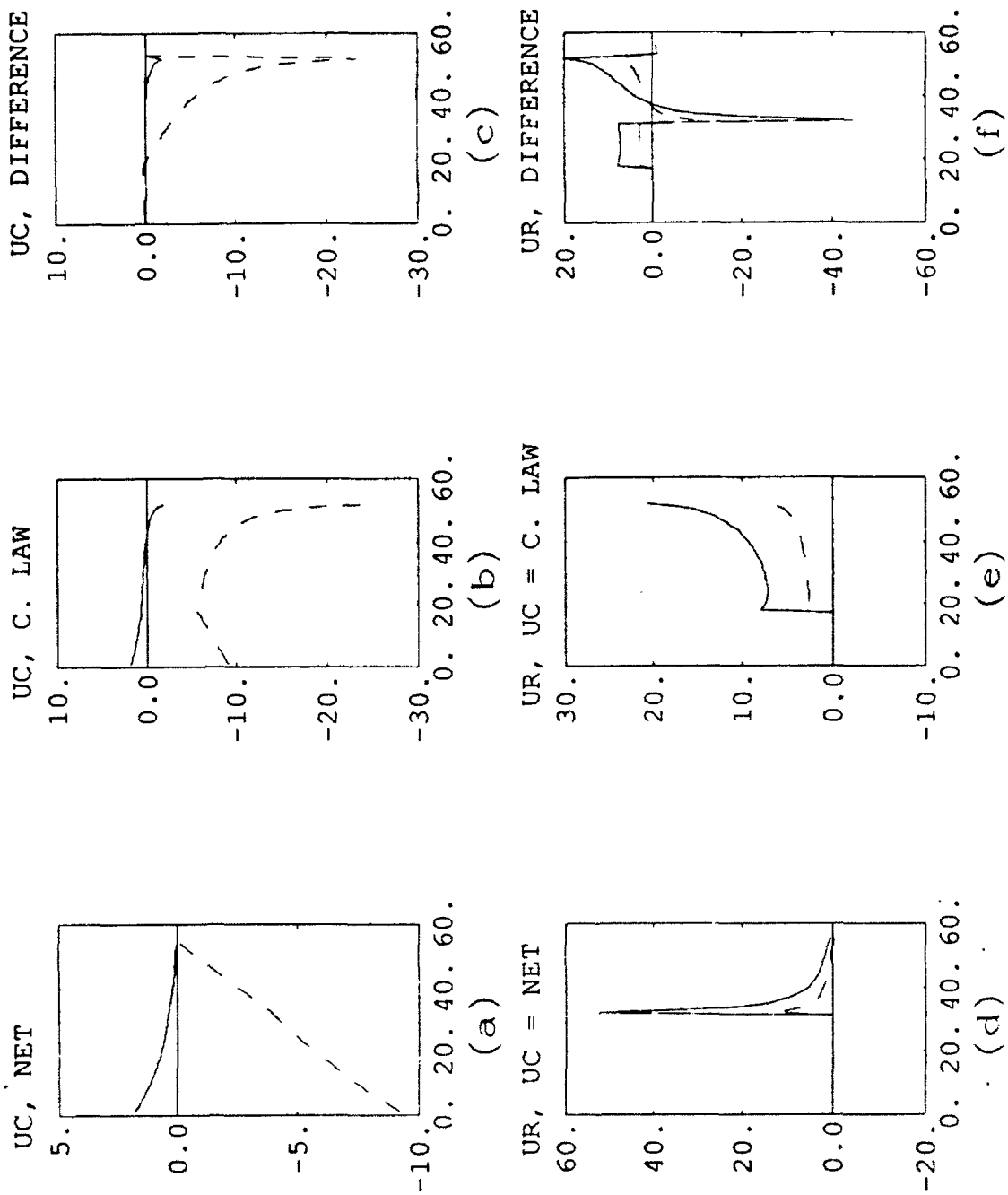


Figure 6. Comparison between optimal control law and neural network controller configurations in the presence of target acceleration with the critic in the control loop and no on-line learning - controller and critic outputs.

inputs are radically different for the two configurations as seen in Figures 5d-5f. A decomposition of the plant inputs into their critic and controller components is shown in Figures 6a-6f. The results show an unstable interaction between controller and critic in the optimal control law configuration. The neural network controller, on the other hand, produces decreasing signals as it was trained to do and the critic produces a complementary control which also decays to zero.

In our final simulation we demonstrate the use of on-line adaptation. Here again we used the initial states as in (16). At each sampling instant beyond the point where the critic switches on-line we use backward error propagation to modify controller weights. Figure 7 shows typical error curves as calculated for training at the indicated sample instances, k , using equation (14). The required number of learning iterations ranges from six to a maximum allowed number of 100 during the time when the critic network is switched on-line. For the first five sampling instants beyond the critic switching point the controller error decreases during the course of training, however, it does not get below the specified error threshold of $\epsilon = 0.01$. For all sampling instances after these, however, the controller network does converge in less than 35 iterations.

Figure 8 shows the results of the simulation. It is interesting to note that the point of minimum range is very nearly the same as that shown in Figure 2a where there was no target acceleration. Moreover, we again observe a significant miss in the y-component of relative position. In this case the miss distance is 45.1 feet. Figure 8b shows that the decrease in flight time has an associated increase in control energy. The effect of on-line learning can be seen in Figure 8c beyond the critic switching point. The on-line weight adjustment also modifies the critic's operation as seen in Figure 8d.

DISCUSSION AND CONCLUSIONS

In this report we have outlined an approach to on-line learning in a missile guidance problem. The approach imbeds two feed-forward neural networks into a control structure which approximates linear optimal guidance. A controller network initially provides command accelerations based on relative positions and velocities. A critic network monitors these relative positions, velocities, and also plant inputs. Using this information the critic responds to unknown target accelerations by producing a correction command which complements the controller command. At each time step the correction command from the critic is used to modify the controller network. In addition, we have compared neural network performance with performance resulting from the use of a linear optimal control law.

There are two components of this work which distinguish it from previous work. The first is the fact that in this work we replaced a linear optimal guidance law (the controller) with a neural network approximation. This replacement was made in order to implement the second distinguishing feature of this work, that is, the process of on-line

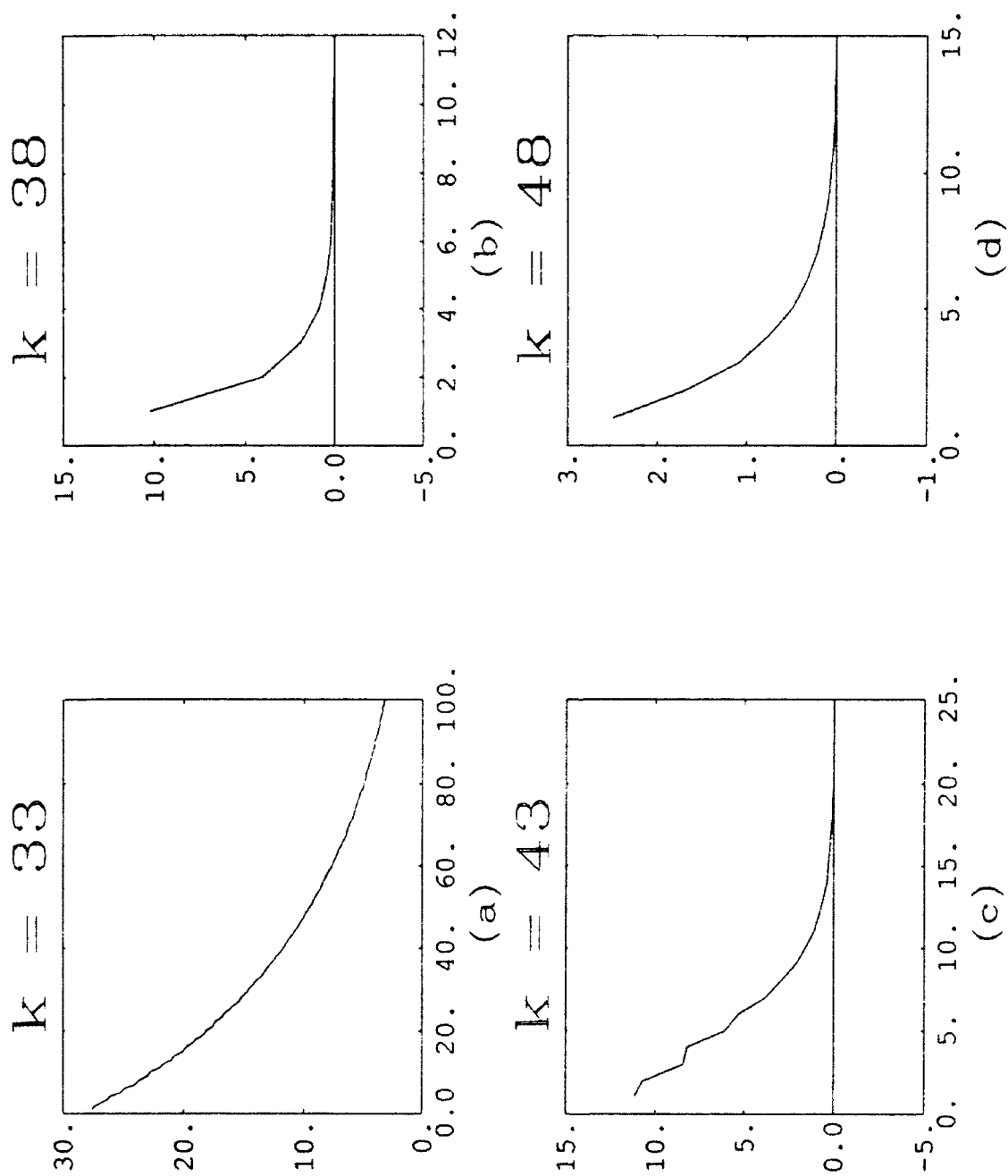


Figure 7. Typical controller network training error convergence

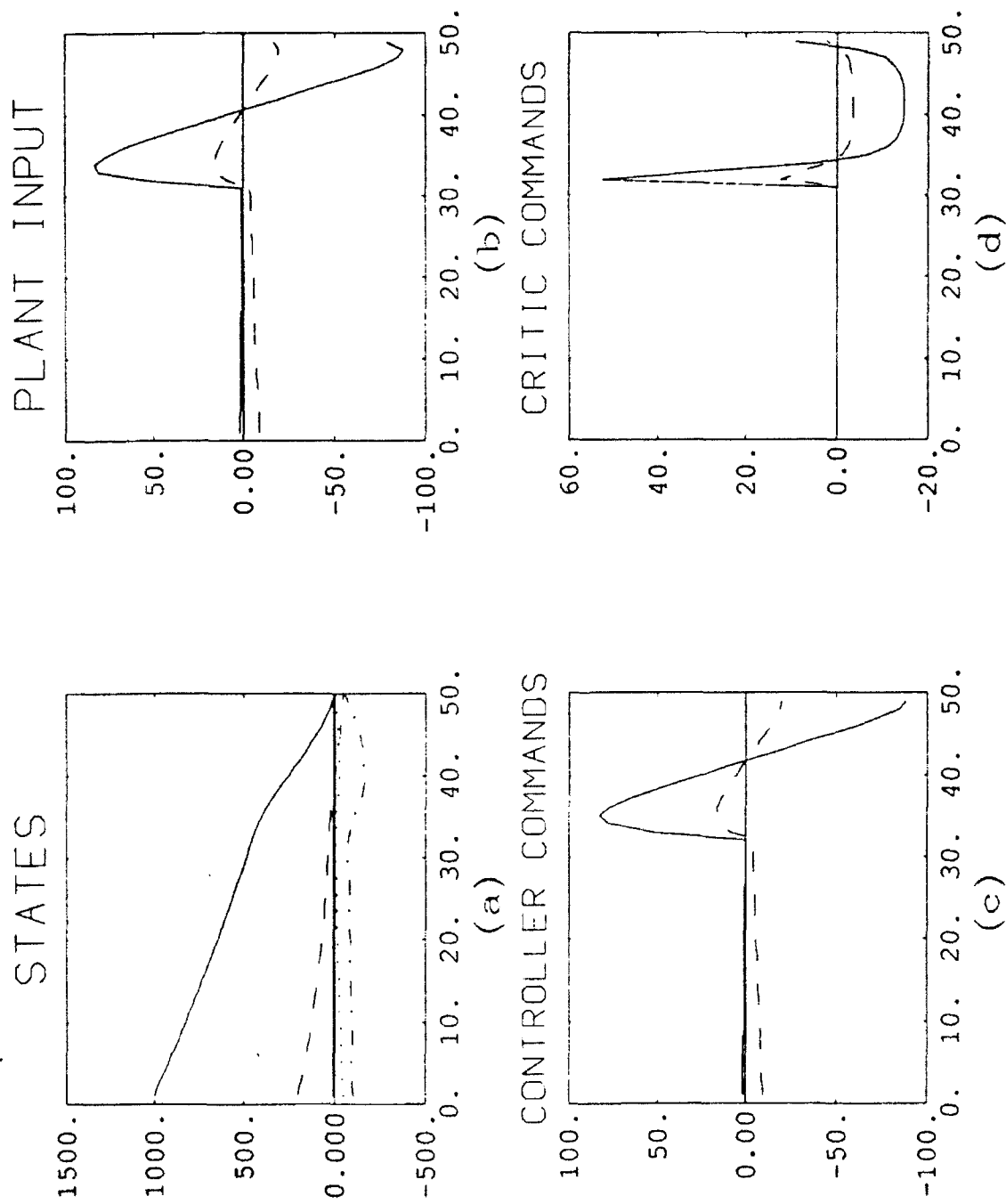


Figure 8. Simulation results with both network controller and critic in the control loop and on-line learning

adaptation.

Simulation results included here demonstrate the method and compare its performance with previous results. We first show that the neural network approximation for the controller is capable of producing control signals which drive the missile to intercept for a non-maneuvering target. Next, in following the method of previous work, we show that the controller produces expected results by itself in the presence of a target maneuver. In a third scenario, we add the critic to the control loop and observe performance improvements consistent with previous results despite differences between the interaction between controller and the critic. In the final step we show further improvements by implementing on-line adaptation of the controller.

The results shown here are encouraging, however, they do show that initial training of the neural networks must be done with care. In our simulations we examined results based only on a single set of initial conditions. The controller and critic networks are trained based on a nominal (no target acceleration) trajectory issuing from these assumed initial conditions. Thus the controller and critic networks are trained along paths in their input spaces rather than on regions. For the initial conditions that were selected in this study the x component of the relative position dominates that of the y component. As a result of this dominance the training sets for the controller and critic networks are biased in the x direction causing a greater focus of control attention in this direction. We observed that control in the y direction failed to bring the y component of relative position completely to intercept. We believe however that this problem can be corrected by more judicious initial training of the networks.

It is also important to note that initial training becomes an issue in generalizing these results to different initial states in the intercept problem. Training the controller is straight forward. The mapping to be performed by the controller network is well defined away from the point of intercept. A feed-forward neural network which performs this mapping to a prescribed degree of accuracy can be determined for a given operating region in the state space. Also note that this initial training of the neural network need only be done once. The controller may then be used as a starting point for initial conditions within the selected operating region.

We are currently pursuing several extensions of this work. These include the ability to operate from different initial points in the state space while simultaneously increasing the accuracy of neural network mappings. In particular, we are looking at alternate ways of training the controller and critic networks. Our first goal is to train the networks to correctly map from *regions* of their input spaces to corresponding regions in their output spaces. This will enable us to solve the intercept problem for more general initial conditions.

We will also investigate the partitioning of the critic and controller networks each into separate networks for the two direction axes. For example, a single controller network could be trained to take relative position, relative velocity and time-to-go for one axis only to produce the corresponding control command. There are three advantages

to doing this. The first is that smaller networks may be used to produce a simpler mapping. Second, the control produced for the x-axis is not directly dependent upon positions and velocities from the y-axis and vice versa (note that an indirect connection still exists in the use of time-to-go as an input variable). The third benefit of this type of structure is that the same trained network can be duplicated for each axis at the beginning of the engagement.

ACKNOWLEDGEMENTS

We would like to thank Dr. Jim Cloutier and the other members of WL/MNAG for their useful discussions and encouragement during the summer. We would also like to thank the Department of the Air Force and staff of RDL for their support of the summer research program.

REFERENCES

- [1] Pastric, H., Seltzer, S., and Warren, M., "Guidance Laws for Short-Range Tactical Missiles," *Journal of Guidance and Control*, Vol. 4, No. 2, pp 98-108, March-April 1981.
- [2] Miller, T. W., Sutton, R. S., and Werbos, P. J., eds., Neural Networks for Control, MIT Press, Cambridge, MA, 1990.
- [3] Balakrishnan, S. N., and Dalton, J., "Neural Networks for Homing Missile Guidance," Technical Report, Department of Mechanical and Aerospace Engineering and Engineering Mechanics, University of Missouri-Rolla, Rolla, MO, 1992.
- [4] Fiske, P. H., "Advanced Digital Guidance and Control Concepts for Air-to-Air Tactical Missiles," Armament Development and Test Center, Eglin Air Force Base, FL, Rept. FAFTL-TR-77-130, 1977.
- [5] Rumelhart, D. E., Hinton, G. E., and Williams, R., "Learning Internal Representations by Error Propagation," in Parallel and Distributed Processing: Explorations in the Microstructure of Cognition, Rumelhart, D. E., McClelland, J. L., and PDP Research Group, eds., MIT Press, Cambridge, MA, 1986.

Optimal Detection of Targets in Clutter using
an Ultra-Wideband, Fully-Polarimetric SAR

Ronald L. Dilsavor
Graduate Research Associate
Department of Electrical Engineering

The Ohio State University
2015 Neil Avenue
Columbus, Ohio 43210

Final Report for:
AFOSR Summer Research Program
Wright-Patterson Air Force Base

Sponsored by:
Air Force Office of Scientific Research
Bolling Air Force Base, Washington, D. C.

September 10, 1992

Optimal Detection of Targets in Clutter using an Ultra-Wideband, Fully-Polarimetric SAR

Ronald L. Dilsavor
Graduate Research Associate
Department of Electrical Engineering
The Ohio State University

Abstract

This report presents the results of work accomplished during the 8-week AFOSR summer research program at the AARA lab of Wright Patterson Air Force Base. The goal of this research is to design and analyze optimal detectors for targets in a clutter-filled environment using a low-frequency, ultra-wideband, fully-polarimetric synthetic aperture radar (SAR). In this report, we focus on optimal techniques for combining information across polarization to detect point targets. It is realistic in many cases to assume that the target of interest will have an unknown amplitude, unknown orientation about the radar line-of-site, and unknown absolute phase. These unknown parameters lead to a detection problem with composite hypotheses and nuisance parameters. We develop a GLRT detector which is designed to accomodate these unknown parameters. The first phase of this research is to obtain a realistic model of the clutter statistics by analyzing clutter data which was measured in the field. This clutter statistical analysis and some preliminary results in the GLRT design were accomplished during the AFOSR summer research program and are presented here. The clutter statistics were found to be well-modeled by the K-distribution and the maximum likelihood (ML) estimates of the target amplitude, orientation, and absolute phase are obtained through a bounded two-dimensional search.

Optimal Detection of Targets in Clutter using an Ultra-Wideband, Fully-Polarimetric SAR

Ronald L. Dilsavor

I. Introduction

The goal of this research is to develop, analyze, and compare optimal radar signal processing techniques for detecting targets of interest in a clutter-filled environment. The radar data is fully-polarimetric and is collected at short range over an ultra-wide bandwidth by a low-frequency SAR. The diverse nature of the measurements provides the detector a substantial amount of target scattering information and the low-frequency band provides improved penetration of small-scale clutter.

We approach this general detection problem by subdividing it into three subproblems which concern each diverse aspect of the data in turn. Namely, we are interested in optimal techniques for combining information across 1) frequency, 2) angle, and 3) polarization. In this report, we focus on techniques which combine information across polarization for the purpose of detecting point targets in clutter. We assume that the frequency diverse measurements are used to obtain high resolution, complex in-phase and quadrature (IQ), time-domain pulse responses at each position along the synthetic aperture. Thus, we have three complex responses (HH, HV, and VV) at each aperture position. These angle diverse responses may then be combined using a SAR or tomographic imaging technique to create two-dimensional HH, HV, and VV images of the observed area [1]. The polarimetric target detection techniques described in this report may be applied to the polarimetric signals before the imaging step or to the polarimetric images after imaging.

Note that the SAR or tomographic imaging techniques mentioned above may not be optimal with regard to detecting certain targets of interest in certain clutter environments; hence, we reserve the right to redefine "image" to suit the purpose of detection [2]. Imaging techniques which are optimal for our detection purposes are the subject of further research and are not addressed in this report. Subsequently, we use the term "image" to refer to the two-dimensional result of combining information across aperture position using some chosen technique.

In many problems of practical interest, the detection problem is complicated by a target which is incompletely specified. We assume that the target signal may be specified in terms of a number of unknown parameters. In particular, we make the realistic assumption that the target amplitude, orientation about the line-of-sight, and absolute phase may be unknown. We may choose to treat an unknown parameter as random with a known probability density function (pdf) or as an unknown nonrandom parameter. The

random parameter case leads to the LRTI, a likelihood ratio test which is integrated over the known pdf. The nonrandom parameter case leads to two accepted approaches: the uniformly most powerful invariant (UMPI) test, if such a test exists, and the generalized likelihood ratio test (GLRT) which necessarily exists [3, 4].

The goal of this research is to derive and compare the polarimetric LRTI, GLRT and UMPI tests for detecting point targets of unknown amplitude, orientation about the line-of-sight, and absolute phase. It is expected that at high signal-to-clutter ratio (S/C) the GLRT will outperform the UMPI and LRTI tests. As the S/C decreases, however, there may come a point at which the ML estimates of the unknown parameters are poor enough that the GLRT performance will drop below that of the UMPI and LRTI tests.

In order to design these optimal detectors, we need models for target and clutter scattering. In our application, the clutter consists of scattering from a forest and the targets of interest are, at this stage of the research, canonical scattering centers. The full-polarization scattering characteristics of many canonical point targets such as dipoles, flat plates, dihedrals, and trihedrals are well-documented [5] [7]. In our application, we require clutter pixel statistics at 1.0×1.0 ft resolution in the low frequency range 0.2 - 1.5 GHz. These clutter statistics are not widely available in the literature [8]. Hence, this research must begin with a clutter statistical analysis at the resolution and frequency range given above. The goal of the analysis is an analytical expression for the clutter probability density function (pdf) in terms of a small number of parameters which can be estimated from the clutter data. Once a realistic pdf is chosen we can continue with the detector design.

This report contains the results of the first phase of the research that was completed during the 8-week AFOSR Summer Research Program. Namely, the clutter analysis results and some preliminary results in the GLRT design are presented here. The next section briefly describes the SAR system and forest clutter environment and presents the SAR forest clutter images. Section III compares the ability of the Gaussian and K-distributions to model forest clutter scattering. Section IV provides a brief introduction to the Huynen parameterization of the scattering matrix along with histograms of the Huynen parameters of the SAR images. The Huynen parameters lend valuable insight to the detection problem. Section V presents current approaches to polarimetric detector design and discusses how our approach compares with them. Section VI presents some preliminary results in the GLRT design process. The final section presents our conclusions.

II. SAR System Description and Clutter Imagery

Forest clutter data was measured by a proof-of-concept linear aperture SAR. The SAR has an aperture length of 120 ft while the aperture spacing between successive measurements is 4 inches giving a total of

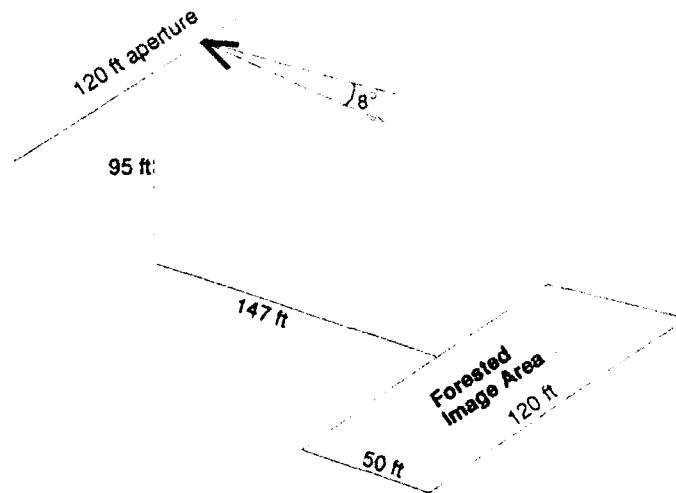


Figure 2.1: Radar geometry.

361 measured signals. The aperture is approximately 95 ft above ground level and the radar views the forest at a declination angle of 8° . The forest clutter image area extends the length of the aperture (120 ft), has a downrange dimension of 50 ft, and begins at a range of 147 ft. Figure 2.1 shows the radar geometry. This central Ohio site consists of gently rolling, glaciated upland that supports a mixed deciduous forest dominated by northern red oak, black cherry, and hickory. Tree sizes range primarily from 4-12 inches dbh (diameter at breast height) with relatively few larger trees exceeding 20 inches dbh. The site is fully-forested with a basal area of 106 sq ft/acre in trees exceeding 2 inches dbh.

Measurements were made in the ultra-wide, low frequency range 0.22 - 1.56 GHz. These frequency measurements were used to synthesize the complex IQ time domain pulse responses [9, p.218]. Complex HH, HV, and VV images were formed using a modified convolution backprojection algorithm [1, 10, 11] applied to the pulse responses.

Figure 2.2 shows the magnitude of the complex HH, HV, and VV images. The horizontal direction corresponds to crossrange and the vertical to downrange. The downrange resolution is 8.9 inches and the crossrange resolution at the maximum frequency of 1.56 GHz is 10.4 inches. Darker pixels represent larger magnitudes. The dark spots correspond to tree trunks and major tree limbs.

III. Comparison of Gaussian and K-Distributions with Clutter Histograms

In this section we compare the ability of the Gaussian and K-distributions to adequately model forest clutter scattering. We begin by computing the polarimetric mean and covariance matrices. We use

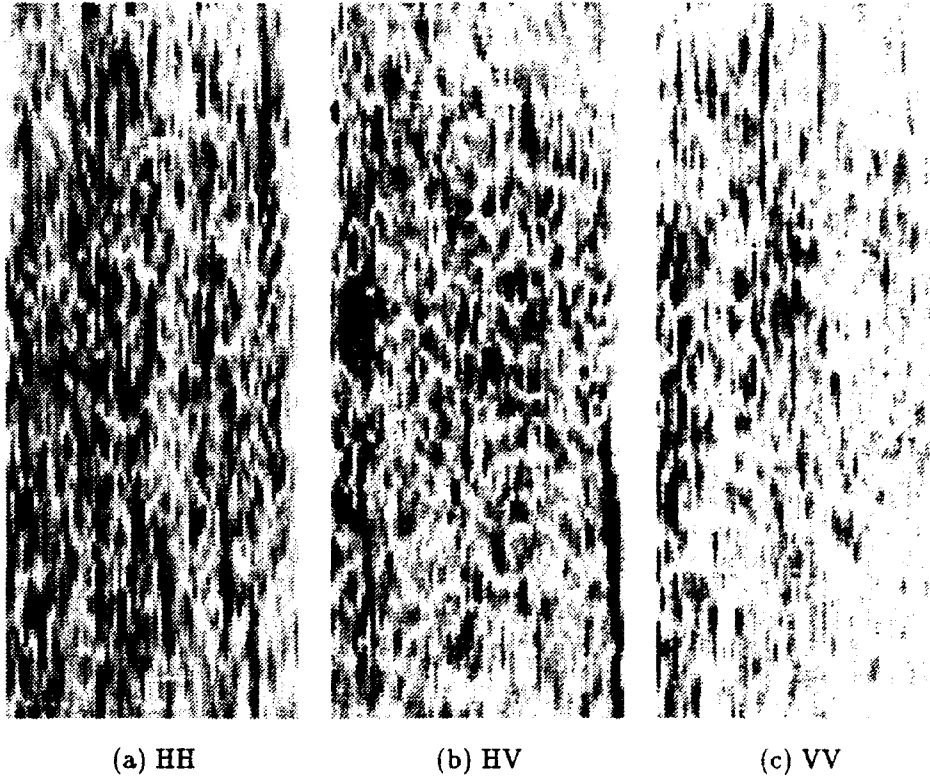


Figure 2.2: Magnitude of the (a) HH image, (b) HV image, and (c) VV image.

these statistics to generate the Gaussian and K-distributed pdfs. These pdfs are overlaid with the clutter histograms for comparison purposes.

A. Mean and covariance of forest clutter scattering

Let X_{Ci} be the complex vector containing samples of the HH, HV and VV images at the i^{th} pixel and let X_{Ri} be the real vector containing the real and imaginary parts of the elements of X_{Ci} :

$$X_{Ci} = [HH_i \quad HV_i \quad VV_i]^T \quad 3.1$$

$$X_{Ri} = [Re(HH_i) \quad Im(HH_i) \quad Re(HV_i) \quad Im(HV_i) \quad Re(VV_i) \quad Im(VV_i)]^T. \quad 3.2$$

The mean and covariance of the X_{Ri} were computed as

$$\begin{aligned} \bar{X}_R &= \frac{1}{N} \sum_{i=1}^N X_{Ri} \\ cov(X_R) &= \frac{1}{N} \sum_{i=1}^N (X_{Ri} - \bar{X}_R)(X_{Ri} - \bar{X}_R)^\dagger \end{aligned} \quad 3.3$$

where N is the number of pixels and $\dagger$ denotes complex conjugate transpose. The mean and covariance of X_C were computed similarly. The experimental results presented in polar format with angle in degrees are

$$\bar{X}_C = 10^{-5} [0.675\angle -144.9 \quad 0.158\angle -118.1 \quad 0.360\angle -143.4]^T \quad 3.4$$

$$\text{cov}(X_C) = 10^{-5} \begin{bmatrix} 0.651\angle 0 & 0.0246\angle -132.5 & 0.0722\angle -82.3 \\ 0.0246\angle 132.5 & 0.1196\angle 0 & 0.0074\angle -120.5 \\ 0.0722\angle 82.3 & 0.0074\angle 120.5 & 0.2371\angle 0 \end{bmatrix} \quad 3.5$$

$$\bar{X}_R = 10^{-5} [-0.552 \quad -0.388 \quad -0.075 \quad -0.140 \quad -0.289 \quad -0.215]^T \quad 3.6$$

$$\text{cov}(X_R) = 10^{-5} \begin{bmatrix} 0.3308 & 0.0110 & -0.0101 & 0.0076 & 0.0108 & 0.0393 \\ 0.0110 & 0.3203 & -0.0106 & -0.0065 & -0.0323 & -0.0011 \\ -0.0101 & -0.0106 & 0.0586 & -0.0028 & 0.0016 & 0.0025 \\ 0.0076 & -0.0065 & -0.0028 & 0.0610 & -0.0039 & -0.0053 \\ 0.0108 & -0.0323 & 0.0016 & -0.0039 & 0.1066 & -0.0021 \\ 0.0393 & -0.0011 & 0.0025 & -0.0053 & -0.0021 & 0.1305 \end{bmatrix} \quad 3.7$$

Notice from 3.4–3.7 that the pixel values are largely zero mean and that for any polarization the real and imaginary parts of a pixel value are largely uncorrelated. These observations agree with findings in [12]. The magnitudes and phases of the covariances in 3.5 are different than found elsewhere. Novak and others [13] estimated the forest clutter covariance matrix for a radar operating at 33 GHz and a similar resolution to be

$$\text{cov}(X_C) = k \begin{bmatrix} 1\angle 0 & 0 & 0.61\angle 0 \\ 0 & 0.16\angle 0 & 0 \\ 0.61\angle 0 & 0 & 0.89\angle 0 \end{bmatrix}.$$

Our measurements show a weaker VV response relative to HH and a smaller correlation coefficient between HH and VV than do those of Novak.

B. Gaussian and K-distributions and clutter scattering histograms

Several families of distributions have been used to model clutter scattering. Gaussian statistics have been assumed in many instances to classify terrain types and to reduce speckle while enhancing target components [14]–[16]. But experimental evidence has shown that many terrain clutter distributions typically have a shape that is “heavy-tail” when compared to that of a Gaussian distribution. Product models such as Weibull, lognormal, and the K-distribution (or gamma distribution) have been found to more accurately model these heavy-tail distributions and have been used in target enhancement and clutter suppression techniques [12, 13], [17]–[21]. We compare the ability of the Gaussian and K-distributions to adequately model forest clutter

scattering at our low frequencies and wide bandwidth. The comparison is made by overlaying plots of best fit pdfs with histograms of clutter data.

The pdfs for a real Gaussian distributed N-vector X (denoted $X : N(m_x, \Sigma_x)$) and a real K-distributed N-vector Y (denoted $Y : K(\alpha, m_y, \Sigma_y)$) are

$$\begin{aligned} p_X(X) &= \frac{1}{(2\pi)^{N/2} |\Sigma_x|^{1/2}} \exp\left[-\frac{1}{2}(X - m_x)^T \Sigma_x^{-1} (X - m_x)\right] \\ p_Y(Y) &= \frac{(2\alpha)^{N/4 + \alpha/2}}{(2\pi)^{N/2} |\Sigma_y|^{1/2} 2^{\alpha-1} \Gamma(\alpha)} [(Y - m_y)^T \Sigma_y^{-1} (Y - m_y)]^{\alpha/2 - N/4} \times \\ &\quad K_{N/2 - \alpha}(\sqrt{2\alpha} [(Y - m_y)^T \Sigma_y^{-1} (Y - m_y)]^{1/2}). \end{aligned} \quad 3.8$$

Here K_ν is the modified Bessel function of order ν . α is a shape-adjusting parameter for the K-distribution, and (m_x, Σ_x) and (m_y, Σ_y) are the mean and covariance of X and Y , respectively. In 3.8 we have used the one-parameter, non-zero mean version of the K-distribution that is derived using the homodyned approach [12, 16, 20, 22]. It has been shown that $Y : K(\alpha, 0, \Sigma_y)$ may be represented as $Y = \sqrt{g}X$ where g is a gamma distributed random variable with parameter α and $X : N(0, \Sigma_y)$ [12]. In fact, the distribution of Y approaches that of a Gaussian random variable as $\alpha \rightarrow \infty$.

Figure 3.3 shows the Gaussian and K-distributed densities overlayed on the histograms of the real and imaginary parts of HH. To generate the analytical pdf curves, we assume that the clutter is zero-mean ($m_x = m_y = 0$), as supported by experimental evidence in the previous section. The clutter covariance ($\Sigma \triangleq \Sigma_x = \Sigma_y$) is estimated using 3.3 and is given by 3.7. However, to generate these one-dimensional curves we only required the 1-1 and 2-2 elements of Σ . Using the technique of [12, pp.242, 248] we estimated the parameter α of 3.8 to be $\alpha = 2.56$. Qualitatively, we see that the narrower central lobe of the K-distribution more accurately models the histogram of the clutter scattering. The plots for HV and VV polarizations are similar.

IV. Histograms of Clutter Huynen Parameters

In the previous section we modeled the statistics of the data by fitting a distribution to the real and imaginary parts of the clutter scattering coefficients. Unfortunately, using this representation of the data, it is difficult to gain physical insight into the types of scattering mechanisms present in the clutter data. In order to provide a more intuitive representation of the data, we transform the raw scattering coefficient data to the Huynen parameter representation. In this section, we briefly review the Huynen parameterization of the scattering matrix and present Huynen parameter histograms of the clutter data.

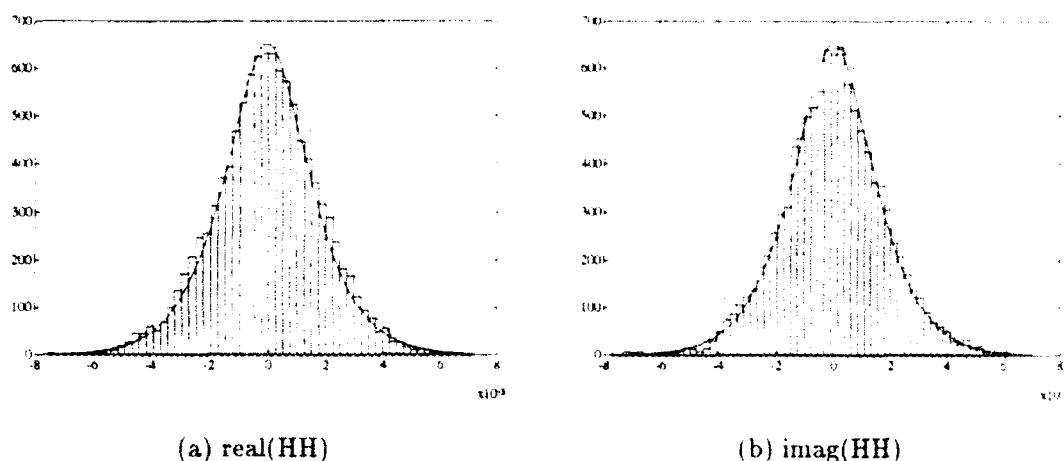


Figure 3.3: Overlay of pdfs with clutter histograms: (a) real(HH), (b) imag(HH). Solid line = K-distribution. Dashed line = Gaussian distribution.

Huynen [5] expresses a general scattering matrix in terms of several geometrically relevant descriptors

$$\begin{aligned}
 \mathbf{S} &= e^{j2\rho} \mathbf{U}^*(\psi, \tau_m, \nu) m \begin{bmatrix} 1 & 0 \\ 0 & \tan^2 \gamma \end{bmatrix} \mathbf{U}^H(\psi, \tau_m, \nu) \\
 \mathbf{U}(\psi, \tau_m, \nu) &= e^{j\psi} \mathbf{J} e^{j\tau_m} \mathbf{K} e^{j\nu} \mathbf{L} \\
 e^{j\psi} \mathbf{J} &= \cos \psi \mathbf{I} + \sin \psi \mathbf{J} \\
 e^{j\tau_m} \mathbf{K} &= \cos \tau_m \mathbf{I} + \sin \tau_m \mathbf{K} \\
 e^{j\nu} \mathbf{L} &= \cos \nu \mathbf{I} + \sin \nu \mathbf{L} \\
 \mathbf{I} &= \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}, \quad \mathbf{J} = \begin{bmatrix} 0 & -1 \\ 1 & 0 \end{bmatrix}, \quad \mathbf{K} = \begin{bmatrix} 0 & i \\ i & 0 \end{bmatrix}, \quad \mathbf{L} = \begin{bmatrix} -i & 0 \\ 0 & i \end{bmatrix}.
 \end{aligned}$$

The positive-valued descriptor m , called the radar target magnitude, gives an overall electromagnetic measure of target size. The angle ψ , called the target rotation angle, gives the orientation of the target about the line-of-sight. The ellipticity of the polarization state that leads to maximum power received from the target is given by τ_m which satisfies $-45^\circ \leq \tau_m \leq 45^\circ$. The angle ν is the target skip angle and is related to the number of bounces of the reflected signal, $-45^\circ \leq \nu \leq 45^\circ$. The angle γ which satisfies $0^\circ \leq \gamma \leq 45^\circ$ is called the characteristic angle of the target. It can be shown that the angle 4γ separates the two null-polarization vectors $\mathbf{p}_{N1}$ and $\mathbf{p}_{N2}$ which satisfy $\mathbf{p}_{N1}^T \mathbf{S} \mathbf{p}_{N1} = 0$. The last descriptor needed to fully define the scattering matrix is the absolute phase $\rho \in [-90^\circ, 90^\circ]$ of the target. The absolute phase is a mixed target parameter in that it depends upon the distance from the radar to the target.

In addition, Huynen, Boerner, and others [5, 23, 6] have derived, from basic symmetry arguments, the

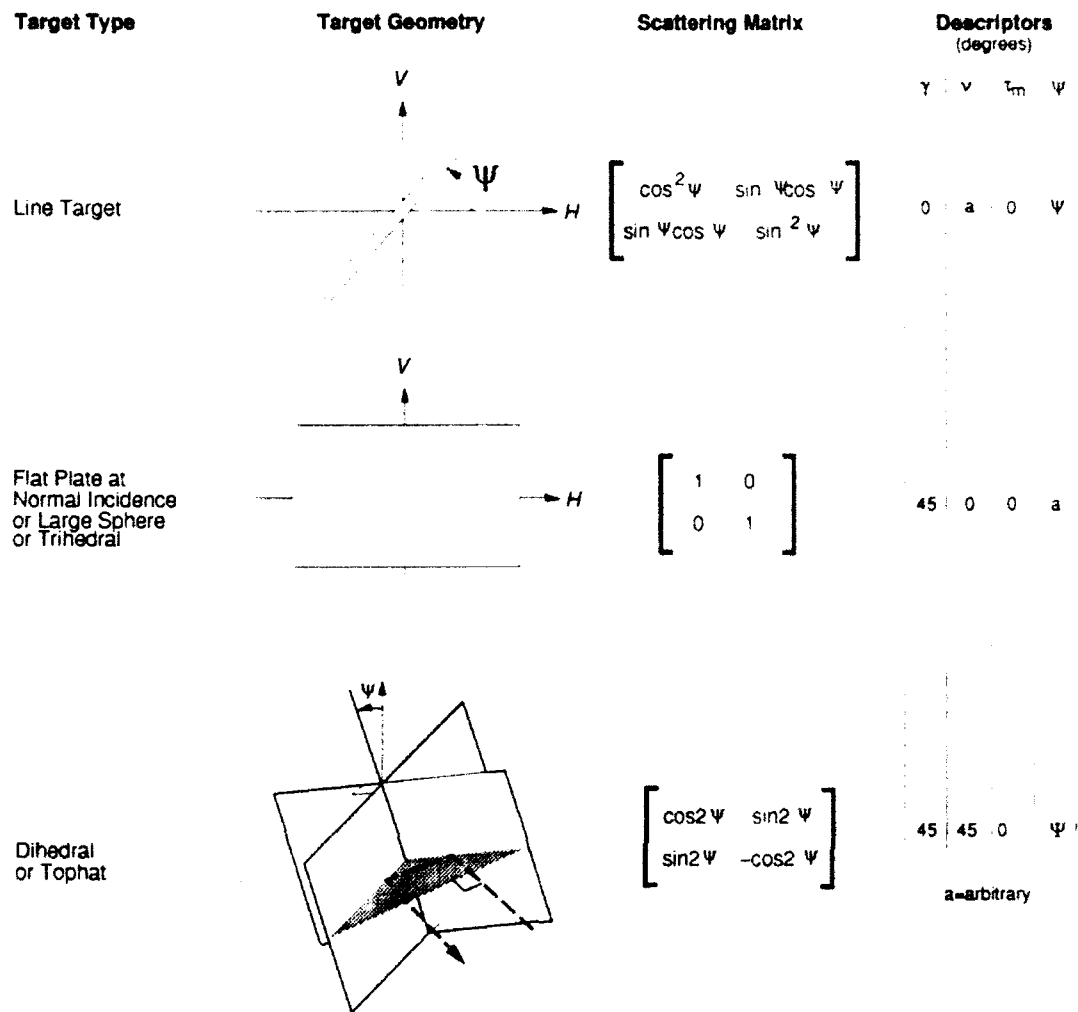


Figure 4.4: Examples of canonical scattering.

scattering matrices of several simple canonical scattering mechanisms. Figure 4.4 presents a few canonical scatterers along with their scattering matrices and some of their geometrical descriptors. In Figure 4.4, the top hat scatterer gives a double-bounce scattering mechanism provided by a vertical cylinder sitting on a ground plane. The top hat scatterer represents, for example, a tree-ground interaction in the forest. These scattering mechanisms and others are expected to comprise target and clutter scattering.

Figure 4.5 shows histograms of the Huynen parameters of the clutter images. Let p_m denote the incident polarization state which causes the scattered signal from a clutter pixel to have maximum amplitude. Then the target amplitude m is that maximum amplitude. The histogram of m is similar to a Rayleigh distribution as would be expected for the maximum eigenvalue of a nearly Gaussian distributed scattering matrix S . If the polarization state $p_m^\perp$ that is orthogonal to p_m is incident on the clutter pixel, then the scattered signal

would have amplitude $m \tan^2 \gamma$ where γ is the characteristic angle. In other words, $\tan^2 \gamma$ gives the ratio of the minimum to maximum eigenvalues of S . The histogram shows that this ratio is $\overline{\tan^2 \gamma} = 0.39$ on the average with a standard deviation $std(\tan^2 \gamma) = 0.20$. The orientation ψ about the line-of-sight (ie. the tilt of p_m) is concentrated around 0° , or horizontal. The histogram for the ellipticity τ_m of polarization state p_m peaks at 0° (ie. linear) and tapers linearly down to zero at $\pm 45^\circ$ (ie. right and left circular). The skip angle ν has a nearly uniform distribution over its range with a slight preference towards -15° . Hence, the clutter has a blending of odd and even bounce mechanisms. The absolute phase ρ is uniformly distributed over its range, as expected.

V. Approaches to Polarimetric Target Detection

In this section, we briefly review the state-of-the-art in polarimetric target detection. We then propose a new approach which treats the target signal as deterministic with a small set of unknown parameters. This approach motivates the design of a polarimetric GLRT (PGLRT).

There are a variety of polarimetric target detectors available in the literature. Among these are the span detector [16, 21], polarization whitening filter (PWF) [13, 16, 21, 24], polarization matched filter (PMF) [16, 25], optimal polarimetric detector for Gaussian clutter (OPD) [16, 25], and the optimal polarimetric detector for K-distributed clutter (OPDK) [16]. Table 5.1 lists these detectors and the rationale that support their use.

From the table we see that the span and PWF do not explicitly distinguish between target and clutter scattering characteristics. The span detector, PWF, and PMF are unable to take advantage of the non-Gaussian nature of forest clutter and the OPDG is specifically designed for Gaussian clutter. The OPDG and OPDK treat the target like clutter resulting in a test that chooses between hypotheses with common means and uncommon covariances. The reason for this approach is that the received signal is typically modeled as

$$\begin{aligned} H_0 : X_C &= C \\ H_1 : X_C &= C + e^{j\phi} T \end{aligned}$$

where we have used the subscript C to denote the complex vector representation of the data as in 3.1. The unknown phase ϕ of the target component T is typically modeled as a uniform random variable $\phi : U[-\pi, \pi]$ resulting in a measurement vector X which is zero-mean under either hypothesis. The covariance of X , however, changes from Σ_c under H_0 to $\Sigma_c + \Sigma_t$ under H_1 .

span	$ HH ^2 + 2 HV ^2 + VV ^2 \underset{H_0}{\overset{H_1}{>}} \beta$	incoherent sum, reduce speckle
PWF	$X^\dagger \Sigma_x^{-1} X \underset{H_0}{\overset{H_1}{>}} \beta$	choose $A \geq 0$, hermitian $\ni$ speckle $\triangleq \frac{s.d(X^\dagger A X)}{E[X^\dagger A X]}$ is minimized
PMF	$ V^{* \dagger} X ^2 \underset{H_0}{\overset{H_1}{>}} \beta$ where $V^* =$ max eigenvector of $\Sigma_c^{-1} \Sigma_t$	choose V to maximize $t/c \triangleq \frac{E(V^\dagger T ^2)}{E(V^\dagger C ^2)}$
OPDG	$X^\dagger (\Sigma_c^{-1} - \Sigma_{t+c}^{-1}) X + \ln \frac{ \Sigma_c }{ \Sigma_{t+c} } \underset{H_0}{\overset{H_1}{>}} \ln \beta$	lrt; common means, uncommon covar. $H_1 : N(0, \Sigma_{t+c}), H_0 : N(0, C_c)$
OPDK	see [16]	lrt; common means, uncommon covar. $H_1 : K(\alpha_{t+c}, 0, \Sigma_{t+c}), H_0 : K(\alpha_c, 0, \Sigma_c)$
where $X = [HH \ HV \ VV]^T = T + C =$ measurement vector. T = target component of X , C = clutter component of X , β = threshold, and $\Sigma_x = E[XX^\dagger]$ = covariance of measurement X , similarly for Σ_c and $\Sigma_{t+c} = \Sigma_t + \Sigma_c$.		

Table 5.1: Polarimetric target detectors.

Our approach is motivated by the fact that, in practice, a scattering center target to be detected is completely known except for its amplitude m , absolute phase ρ , and orientation ψ about the line-of-sight. We explicitly denote these unknowns in our model of the received signal

$$\begin{aligned}
 H_0 : X_R &= C \\
 H_1 : X_R &= C + m \cdot R(\rho)P(\psi)T_n
 \end{aligned}$$

where we have switched to the real vector representation of the data as in 3.2. In this model, T_n is the nominal target with amplitude $m \triangleq 1$, phase $\rho \triangleq 0$, and orientation about the line-of-sight $\psi \triangleq 0$. After premultiplying T_n by the matrices $P(\psi)$ and $R(\rho)$ and the scalar m , the resulting target contribution has amplitude m , phase ρ , and orientation ψ . The form of $R(\cdot)$ and $P(\cdot)$ are found easily and are not presented here. In addition, the clutter is taken to be distributed as $K(\alpha, 0, \Sigma)$, where Σ is given by 3.7 and α is given at the end of Section IIIB.

If m , ρ , and ψ were known then this problem is one of testing for a completely known target in noise of known statistics

$$\begin{aligned}
 H_0 : X_R &: K(\alpha, 0, \Sigma) \\
 H_1 : X_R &: K(\alpha, m \cdot R(\rho)P(\psi)T_n, \Sigma).
 \end{aligned}$$

The resulting polarimetric likelihood ratio test for K-distributed clutter (PLRTK)

$$\frac{p_{X|H_1}(X|H_1)}{p_{X|H_0}(X|H_0)} \underset{H_0}{\overset{H_1}{>}} \beta \quad 5.9$$

chooses between hypotheses with uncommon means and common covariances and its performance relative to the OPDK of Table 5.1, a common means - uncommon covariances test, is not known.

Since m , ρ , and ψ , are typically unknown we estimate them in the maximum likelihood (ML) sense and implement a generalized likelihood ratio test [3]

$$\frac{\max_{m,\rho,\psi} p_{X|m,\rho,\psi}(X|m,\rho,\psi)}{p_{X|H_0}(X|H_0)} \underset{H_0}{\overset{H_1}{>}} \beta \quad 5.10$$

which is simply the standard PLRTK with the unknown quantities replaced by their ML estimates. The performance of this polarimetric GLRT for K-distributed clutter (PGLRTK) as a function of the target-to-clutter ratio (t/c) and relative to the PLRTK and OPDK is not known.

VI. Preliminary Results in the GLRT Design

This section presents our preliminary results in the implementation of the PGLRTK of 5.10. A potential closed form expression for the ML estimate of target amplitude m is derived as a function of the ML estimates of ρ and ψ . This result reduces the search for the maximum of the likelihood function from three dimensions to two bounded dimensions. Next we add a known target to our forest clutter image data, obtain ML estimates of its amplitude, phase and orientation, and view the two-dimensional likelihood function.

We begin by looking for closed form expressions for the ML estimates of m , ρ , and ψ . If we let $X_w = \Sigma^{-1/2}X$ be the whitened version of the measurement vector X then

$$X_w : K(\alpha, m\Sigma^{-1/2}R(\rho)P(\psi)T_n, I)$$

and the likelihood function under H_1 is of the form

$$\begin{aligned} p_{X_w|m_w,\rho,\psi}(X|\rho,\psi) &= Q_1 \|X - m_w D\|^{Q_2} K_{Q_3}(Q_4 \|X - m_w D\|) \\ D &\triangleq \Sigma^{-1/2}R(\rho)P(\psi)T_n \end{aligned}$$

where m_w is the amplitude of the "whitened target" D and the Q_i are non-zero constants independent of m_w , ρ , and ψ . To find the ML estimate of m_w we solve

$$\begin{aligned} 0 &= \frac{\partial p_{X_w|m_w,\rho,\psi}(X|m_w,\rho,\psi)}{\partial m_w} \\ &= \|X - m_w D\|^{Q_2-2} (m_w D^T D - X^T D) f(m_w) \end{aligned} \quad 6.11$$

where $f(m_w)$ is a complicated function of m_w involving the modified Bessel function K . The first term of 6.11 is zero only with probability zero. The second term is zero for

$$m_w = \frac{X^T D}{D^T D}. \quad 6.12$$

This expression for m_w is intuitively pleasing since it is the projection of the measurement vector X onto the "whitened target" target D . Given values for ρ and ψ , we see that D is determined and m_w may be computed from 6.12. Thus, we need to search only over ρ and ψ for a local extremum of the likelihood function. In fact, the two-dimensional search is bounded since $\rho \in [-\pi/2, \pi/2]$, $\psi \in [0, \pi]$. Further investigation is needed to determine whether $f(m_w) = 0$ for some m_w and hence to discover whether the local extremum found from the two-dimensional search is also a global extremum. In the simulations that follow we use the two-dimensional search procedure to generate the "at least local" ML estimates of m_w , ρ , and ψ . Once the ML estimate of m_w has been found then the effects of whitening and of omitting the VH measurement from X may be removed to yield the ML estimate of m

$$m_{ml} = \frac{\|D\|}{\|R(\rho)P(\psi)T_n\|} m_w$$

Next we test the ability of the ML search procedure described above to produce accurate estimates of the true target amplitude, phase, and orientation. To each of 50 clutter pixels $\{C_i, i = 1, \dots, 50\}$ taken from the polarimetric clutter images of Figure 2.2 we added a linear target of amplitude $m = 0.0104$, phase $\rho = 90^\circ$, and orientation $\psi = 0^\circ$ to produce $\{Y_i, i = 1, \dots, 50\}$ given by

$$\begin{aligned} Y_i &= C_i + 0.0104 \cdot R(90^\circ)P(0^\circ)T_n \\ T_n &= [1 \ 0 \ 0 \ 0 \ 0 \ 0]^T. \end{aligned}$$

For each of the Y_i , we computed the ML estimates of m , ρ , and ψ using the two-dimensional search procedure described above. The search was carried out in 5° increments of ρ and ψ . We then repeated the same experiment except that this time we set the linear target orientation to $\psi = 90^\circ$.

Figure 6.6 shows the ML estimates of m , ρ , and ψ for the two experiments plotted against their true values. In addition, the figure shows $t/c = 20 \log(\|0.0104 \cdot R(90^\circ)P(0^\circ)T_i\|/\|C_i\|)$ in dB. When ψ was estimated to be 175° in (c), we plotted it as -5° . Figure 6.6 shows that the experiment with a target orientation of $\psi = 90^\circ$ produced qualitatively better ML estimates of target amplitude m and phase ρ than did the experiment with $\psi = 0^\circ$. This is explained by noting that the clutter histogram of ψ (see Figure 4.5(e)) is concentrated around 0° and is relatively small at $\psi = 90^\circ$. As expected, the ML estimates improve when t/c is large. Figure 6.7 shows contour plots of the likelihood functions $p_{X_w|m,\rho,\psi}(X|m,\rho,\psi)$ versus ρ and ψ for pixel Y_{13} which has $t/c = 20 \log_{10}(m/\|C_{13}\|) = 7.96$ dB. The likelihood functions display

a strong single peak characteristic which is expected for linear target detection. This would suggest the use of a gradient-based search technique for linear scatterers.

VII. Conclusions

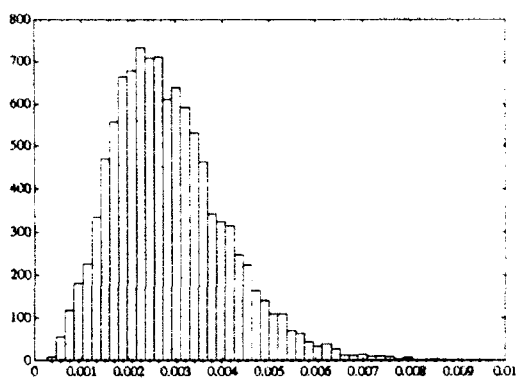
This report presents the results of work accomplished during the 8-week AFOSR summer research program at the AARA lab of Wright Patterson Air Force Base. The goal of this work is a scheme for detecting scattering center targets in clutter using ultra-wideband fully-polarimetric SAR data. First, we analyzed the statistics of forest clutter and found that the K-distribution provides a better model of the clutter than does the Gaussian distribution. In addition, histograms of the Huynen parameters of the forest clutter are provided to characterize the types scattering mechanisms present in the forest and hence to lend insight into the types of man-made scattering centers that are more easily distinguished from clutter with an appropriate detector. Next, we presented a new approach to polarimetric scattering center detection that leads to a generalized likelihood ratio test for scattering center detection in K-distributed clutter (PGLRTK). This test provides detection capability in the presence of unknown target amplitude, absolute phase, and orientation about the line-of-sight. This polarimetric GLRT approach is equally applicable to other clutter types. Further work is needed to determine whether or not the ML estimates resulting from the bounded two-dimensional search do indeed lead to the global maximum of the likelihood function. In addition, the performance of the PGLRTK needs to be compared with currently available detectors. This comparison should be made by first assuming perfect estimation of the unknown parameters and then by studying the loss of performance due to miss-estimation of the parameters in lower target-to-clutter-ratio scenarios.

References

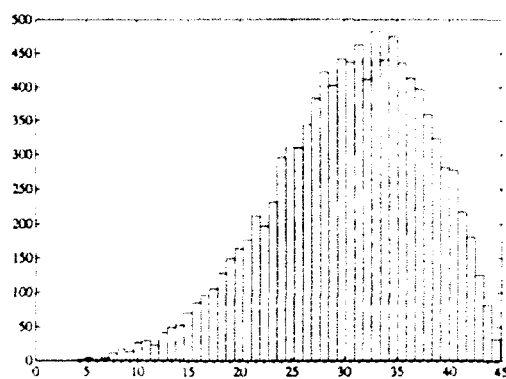
- [1] R. M. Lewitt, "Reconstruction algorithms: Transform methods," *Proceedings of the IEEE*, vol. 71, pp. 390-408, March 1983.
- [2] D. J. Rossi and A. S. Willsky, "Reconstruction from projections based on detection and estimation of objects—parts I and II: Performance analysis and robustness analysis," *IEEE Transactions on Acoustics, Speech, and Signal Processing*, vol. SP-32, August 1984.
- [3] H. L. V. Trees, *Detection, Estimation, and Modulation Theory: Part I*. New York: Wiley, 1968.
- [4] L. L. Scharf, *Statistical Signal Processing*. Reading, Massachusetts: Addison-Wesley, 1991.
- [5] J. R. Huynen, *Phenomenological Theory of Radar Targets*. PhD thesis. Technical University, Delft, The Netherlands, 1970.

- [6] A. P. Agrawal and W.-M. Boerner, "On the concept of optimal polarizations in radar meteorology," in *Proceedings of the Polarimetric Technology Workshop*, vol. I.2. (Redstone Arsenal, Alabama), pp. 179-243, August 16-18 1988.
- [7] G. T. Ruck, D. E. Barrick, W. D. Stuart, and C. K. Krichbaum, *Radar Cross Section Handbook: Volume 2*. New York: Plenum Press, 1970.
- [8] F. T. Ulaby and M. C. Dobson, *Handbook of Radar Scattering Statistics for Terrain*. Norwood, Massachusetts: Artech House, 1989.
- [9] A. D. Poularikas and S. Seely, *Signals and Systems*. Boston, MA: PWS Engineering, 1985.
- [10] J. L. Bauck, *Tomographic Processing of Synthetic Aperture Radar Signals for Enhanced Resolution*. PhD thesis, University of Illinois at Urbana-Champaign, November 1989.
- [11] H.-J. Li and F.-L. Lin, "Near-field imaging for conducting objects," *IEEE Transactions on Antennas and Propagation*, vol. AP-39, pp. 600-682, May 1991.
- [12] S. H. Yueh, J. A. Kong, J. K. Jao, R. T. Shin, H. A. Zebker, T. L. Toan, and H. Ottl, "Chapter 4: K-distribution and polarimetric terrain radar clutter," in *PIER-3: Progress in Electromagnetics Research: Polarimetric Remote Sensing* (J. A. Kong, ed.), New York: Elsevier, 1990.
- [13] L. M. Novak, M. C. Burl, R. D. Chaney, and G. J. Owirka, "Optimal processing of polarimetric synthetic-aperture radar imagery," *The Lincoln Laboratory Journal*, vol. 3, pp. 273-290, summer 1990.
- [14] H. Lim, A. A. Swartz, H. A. Yueh, J. A. Kong, R. T. Shin, and J. J. van Zyl, "Classification of earth terrain using synthetic aperture radar images," *J. Geophys. Res.*, vol. 93, pp. 15252-15260, December 1988.
- [15] H. A. Yueh, A. A. Swartz, J. A. Kong, R. T. Shin, and L. M. Novak, "Bayes classification of terrain cover using normalized polarimetric data," *J. Geophys. Res.*, vol. 93, pp. 15261-15267, December 1988.
- [16] L. M. Novak, M. B. Sechtin, and M. J. Cardullo, "Studies of target detection algorithms that use polarimetric radar data," *IEEE Transactions on Aerospace and Electronic Systems*, vol. AES-25, pp. 150-165, March 1989.
- [17] "Detection in non-gaussian clutter," in *Automatic Detection and Radar Data Processing* (D. C. Schleher, ed.), Dedham, MA: Artech House, Inc., 1980.
- [18] "Noise and clutter rejection in radars and imaging sensors: Proceedings of the second international symposium on noise and clutter rejection in radars and imaging sensors," Kyoto, Japan, November 14-16 1989.

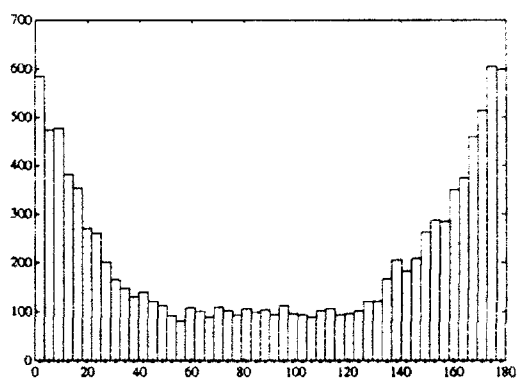
- [19] E. Jakeman, "On the statistics of K-distributed noise," *J. Phys. A: Math. Gen.*, vol. 13, pp. 31-48, 1980.
- [20] J. K. Jao, "Amplitude distribution of composite terrain radar clutter and the K-distribution," *IEEE Transactions on Antennas and Propagation*, vol. AP-32, pp. 1049-1062, October 1984.
- [21] L. M. Novak and M. C. Burl, "Optimal speckle reduction in polarimetric SAR imagery," *IEEE Transactions on Aerospace and Electronic Systems*, vol. AES-26, pp. 293-305, March 1990.
- [22] S. H. Yueh, J. A. Kong, J. K. Jao, R. T. Shin, and L. M. Novak, "K-distribution and polarimetric terrain radar clutter," *J. Electro. Waves Applicat.*, vol. 3, no. 8, pp. 747-768, 1989.
- [23] W. M. Boerner, W. L. Yan, A. Q. Xi, and Y. Yamaguchi, "On the basic principles of radar polarimetry: The target characteristic polarization state theory of Kennaugh, Huynen's polarization fork concept, and its extension to the partially polarized case," *Proceedings of the IEEE*, vol. 79, pp. 1538-1550, October 1991.
- [24] W. W. Irving, G. J. Owirka, and L. M. Novak, "Adaptive processing of polarimetric sar imagery," in *24th Asilomar Conference on Signals, Systems, and Computers*, (Pacific Grove, CA), pp. 388-398, November 1990.
- [25] L. M. Novak and M. B. Sechtin, "Target detection algorithms that use polarimetric radar data: New results," in *Proceedings of the Polarimetric Technology Workshop*, vol. I.1, (Redstone Arsenal, Alabama), pp. 537-555, August 16-18 1988.



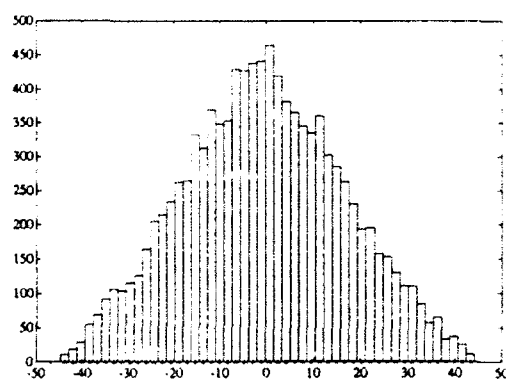
(a) amplitude m



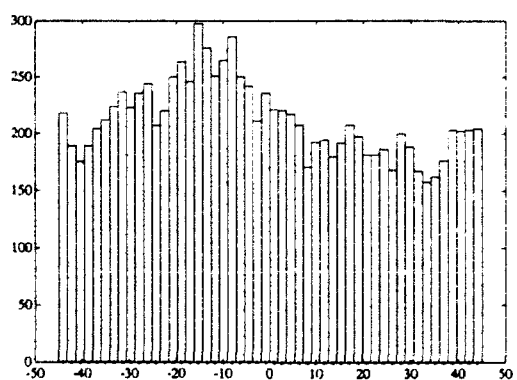
(b) characteristic angle γ



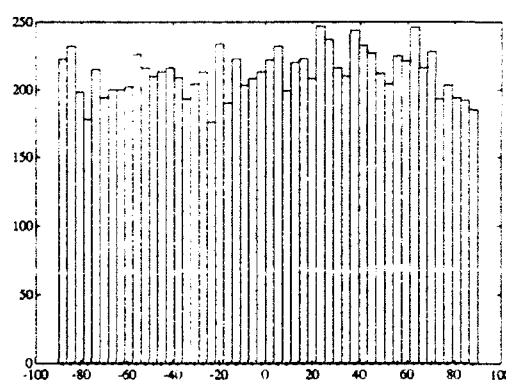
(c) orientation ψ



(d) ellipticity τ

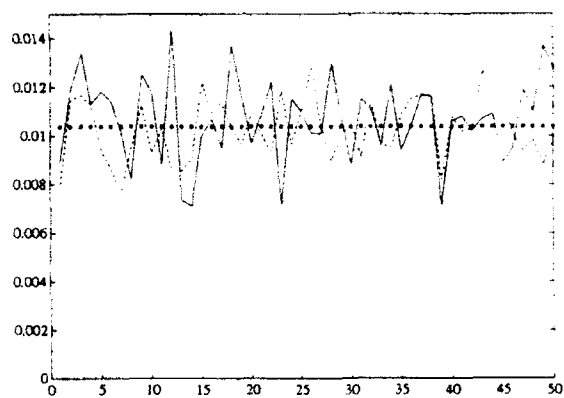


(e) skip angle ν

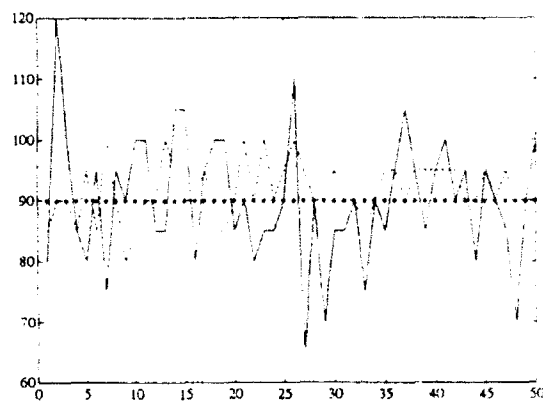


(f) absolute phase ρ

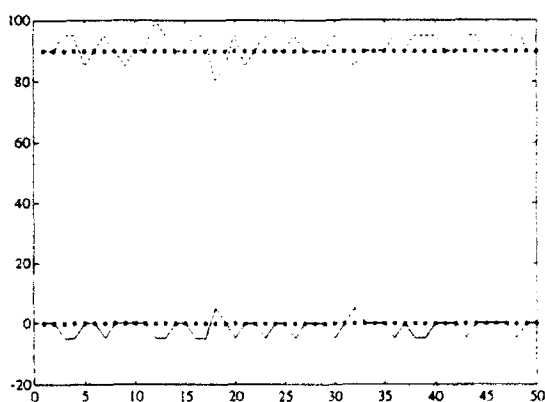
Figure 4.5: Histograms of the Huynen parameters of the clutter data: (a) magnitude m , (b) characteristic angle γ , (c) orientation ψ , (d) ellipticity τ_m , (e) skip angle ν , (f) absolute phase ρ . Y-axes are number of pixels, X-axes of (b) - (f) are in degrees.



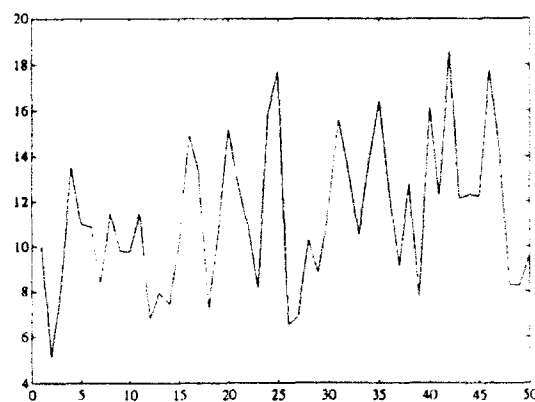
(a) m



(b) ρ

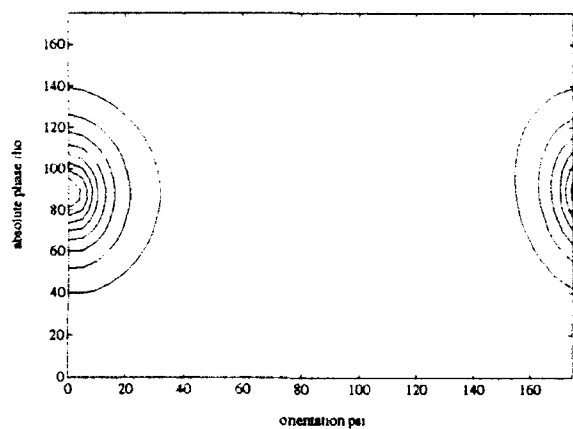


(c) ψ

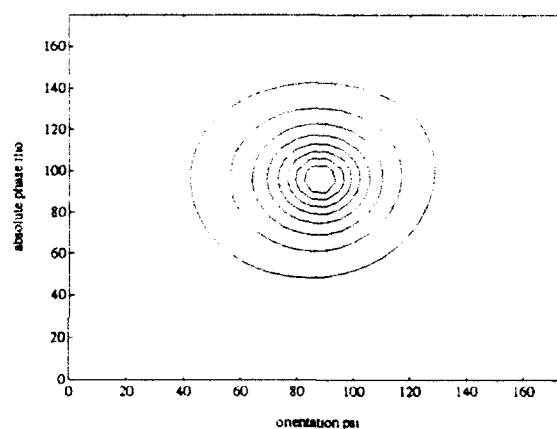


(d) t/c

Figure 6.6: ML estimates and true values of (a) m , (b) ρ , and (c) ψ for the 50 trials. (d) shows the t/c in dB. *:true value, solid line: experiment with $\psi = 0^\circ$, dashed line: experiment with $\psi = 90^\circ$.



(a) likelihood function, true $\psi = 0^\circ$



(b) likelihood function, true $\psi = 90^\circ$

Figure 6.7: Contour plots of likelihood functions for pixel Y_{13} (a) target with $\psi = 0^\circ$ (b) target with $\psi = 90^\circ$. $t/c = 7.96$ dB.

**EFFECTS OF INTENSITY THRESHOLDING
ON THE
POWER SPECTRUM OF LASER SPECKLE**

**Alfred D. Ducharme
Graduate Research Assistant
University of Central Florida**

**Center for Research in Electro-Optics and Lasers
12424 Research Pkwy.
Orlando, FL 32826**

**Final Report for:
Summer Research Program
Wright Laboratory**

**Air Force Office of Scientific Research
Bolling Air Force Base, Washington, D.C.**

September 1992

**EFFECTS OF INTENSITY THRESHOLDING
ON THE
POWER SPECTRUM OF LASER SPECKLE**

**Alfred D. Ducharme
Graduate Research Assistant
Center for Research in Electro-Optics and Lasers
University of Central Florida**

Abstract

Spatial-frequency filtering of laser speckle patterns has proven to be a useful tool in the measurement of MTF for focal plane arrays. Intensity thresholding of the laser speckle patterns offers nearly an order of magnitude savings in digital storage space. The effect of this thresholding on the spatial-frequency power spectral density of the speckle pattern is investigated. An optimum threshold level is found that minimizes distortion of the power spectrum for the classes of speckle data used for MTF testing. Effects of Intensity Thresholding on the Power Spectrum of Laser Speckle

EFFECTS OF INTENSITY THRESHOLDING ON THE POWER SPECTRUM OF LASER SPECKLE

Alfred D. Ducharme

INTRODUCTION

The statistics of speckle phenomena have been the focus of many research efforts since the introduction of the laser. As our knowledge of speckle behavior increases, new ways to use and control it become evident. One useful application of laser speckle has been the measurement of modulation transfer function (MTF) of CCD arrays.¹ In this application a considerable amount of laser speckle data is collected for each MTF calculation which must be stored digitally for future use. Since, digital storage space is always a problem when dealing with computers, it is desirable to reduce the amount space required by this application.

Since the size of the data arrays used in the calculation are fixed by the size of the CCD array being tested, there are only two ways which the data can be reduced. First, the number of points used to plot the MTF could be minimized. This would reduce the accuracy of the calculation since the spacing of data points dictates the size fluctuation in MTF that can be detected. Second, the speckle intensity could be thresholded eliminating all but 2 of the 256 gray levels used to represent the digitized speckle field (Fig. 1). This binarization would transform the 8 bit data into 1 bit data reducing the total amount of data by nearly an order of magnitude.

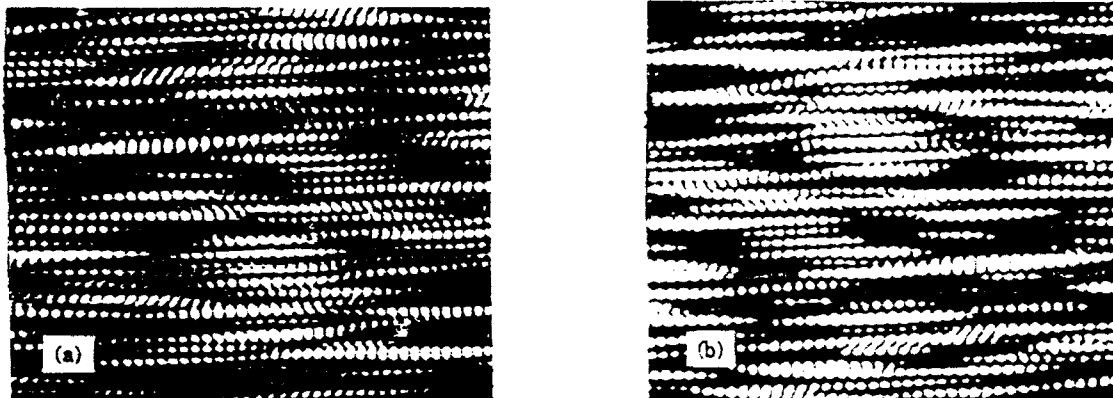


Fig 1 Examples of narrowband laser speckle. (a) Continuous laser speckle with 256 gray levels. (b) Laser speckle reduced to 2 gray levels by thresholding at the mean intensity.

It is the goal of this paper to show that the fidelity of the spatial-frequency power spectral density of the laser speckle intensity is maintained after a thresholding operation is performed. The effects of the intensity thresholding will be investigated and an optimal threshold value will be determined for the narrowband laser speckle used in the MTF calculation.

The second section of this paper will discuss the derivation of the spatial autocorrelation function for the continuous and thresholded intensity of the laser speckle. The derivation will be exemplified using two different cases. The first case will be a simple square aperture and the second will be a double slit aperture used in the calculation of MTF.

The relationship between the autocorrelation and power spectral density of the thresholded laser speckle will be discussed in the third section. The same two examples will be used, and an optimal threshold level will be determined for each case based on a minimum-mean-squared-error analysis in their power spectral densities. A computer simulation of these two examples was performed to check

the validity of this analysis and its results are given in section 4 in terms of power spectrum.

AUTOCORRELATION OF THRESHOLDED LASER SPECKLE

In this section we will derive the spatial autocorrelation function for intensity thresholded laser speckle. To simplify the reading, this function will be referred to as the *thresholded autocorrelation*. The derivation will begin with the spatial autocorrelation function for intensity distribution of non-thresholded laser speckle which will be referred to as the *continuous autocorrelation*.

The continuous autocorrelation is measured using two observations of the speckle intensity. Ideally, these observations are made using two point detectors which sample the speckle field simultaneously at different points in a common plane (x,y). The autocorrelation is formed by finding the expected value of the product of the two observations for each unique separation distance. This is expressed as²,

$$R_I(x_1, y_1; x_2, y_2) = \langle I(x_1, y_1) I(x_2, y_2) \rangle. \quad (2.1)$$

In reality the detectors are of finite size which changes the overall shape of the autocorrelation by convolving eq. (2.1) by the spatial autocorrelation of a single detector. This was not included in this paper so that the singular effect of thresholding could be investigated.

The only factor contributing to the form of the spatial autocorrelation is the amplitude of the field emanating from the generating aperture, $P(\xi, \eta)$. This relationship is described by Goodman² as,

$$R_I(\Delta x, \Delta y) = \langle I \rangle^2 \left[1 + \frac{\left| \iint_{-\infty}^{\infty} |P(\xi, \eta)|^2 \exp \left[i \frac{2\pi}{\lambda z} (\xi \Delta x + \eta \Delta y) \right] d\xi d\eta \right|^2}{\iint_{-\infty}^{\infty} |P(\xi, \eta)|^2 d\xi d\eta} \right]. \quad (2.2)$$

From this we can see that the continuous autocorrelation is basically the magnitude squared of the Fourier transform of $P(\xi, \eta)$. As a result, using eq. (2.2) we can calculate an analytical expression for the continuous autocorrelation for any generating aperture which physically exists.

Our goal is to determine the thresholded autocorrelation, $R_I^{(b)}$. Where we consider the thresholding to occur at both observations at the same value (superscript (b)). We need to determine a relationship between the continuous and thresholded autocorrelations. The speckle field is thresholded using the following operator,

$$I^{(b)}(x, y) = \begin{cases} 1, & \text{if } I(x, y) \geq b \langle I(x, y) \rangle \\ 0, & \text{if } I(x, y) < b \langle I(x, y) \rangle. \end{cases} \quad (2.3)$$

The thresholded autocorrelation can be determined from the continuous autocorrelation using a relationship calculated by Barakat³,

$$R_I^{(b)}(\Delta x, \Delta y) = \frac{\Gamma^2(\alpha, b\alpha)}{\Gamma^2(\alpha)} + \frac{(b\alpha)^{2\alpha}}{\Gamma^2(\alpha)} e^{-2b\alpha} \sum_{n=1}^{\infty} \frac{(R_I(\Delta x, \Delta y))^n}{\binom{n+\alpha-1}{n}} (L_{n-1}^{(\alpha)}(b\alpha))^2, \quad (2.4)$$

where G is the complimentary incomplete gamma function and $L_{n-1}^{(\alpha)}$ is the associated Laguerre polynomial⁵. The parameter α is equivalent to the mean of the intensity squared divided by the variance. In this paper we are assuming

point detector meaning that the value of a is equal to 1. Using $a=1$, eq. (2.4) can be simplified to,

$$R_I^{(b)}(\Delta x, \Delta y) = e^{-b} \left[1 + b^2 \sum_{n=1}^{\infty} \frac{(R_I(\Delta x, \Delta y))^n}{n^2} (L_{n-1}^{(1)}(b))^2 \right]. \quad (2.5)$$

This equation was computed utilizing the recursion formula for Laguerre polynomials⁵ over the range of possible correlation values (0.0 to 1.0). The result is given in Fig. 2 in the form of a mapping. A value of continuous autocorrelation, R_I , can be mapped to the corresponding value of thresholded autocorrelation, $R_I^{(b)}$, using Fig. 2. As an example, for a threshold value of $b = 3$ a value of continuous autocorrelation, $R_I = 0.4$, would map to a value of thresholded autocorrelation, $R_I^{(3)} = 0.25$. Repeating this procedure for all values on a continuous autocorrelation function would yield the thresholded autocorrelation for a particular threshold value.

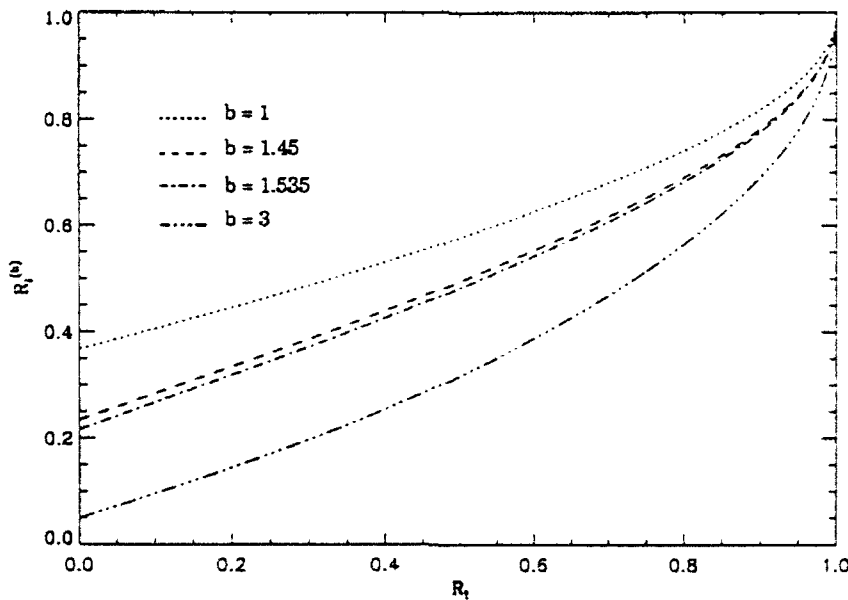


Fig. 2 Mapping for continuous autocorrelation (R_I) to thresholded autocorrelation ($R_I^{(b)}$) for threshold levels: $b = 1, 1.45, 1.535$, and 3 .

Two cases will be used to exemplify this mapping technique. One, a simple

square aperture whose continuous autocorrelation can be found in Goodman³. Two, a double slit aperture which was used in Ref. 1 to produce narrowband filtered laser speckle (Fig. 1).

The amplitude of the field emanating from a square aperture is expressed as,

$$|P(\xi, \eta)|^2 = \text{rect}\frac{\xi}{L} \cdot \text{rect}\frac{\eta}{L}. \quad (2.6)$$

where L is the measure of each side of the square aperture. Substituting eq. (2.6) into eq. (2.2) we find the continuous autocorrelation to be,

$$R_I(\Delta x, \Delta y) = \langle I \rangle^2 \left[1 + \text{sinc}^2 \frac{2L \Delta x}{\lambda z} \cdot \text{sinc}^2 \frac{2L \Delta y}{\lambda z} \right]. \quad (2.7)$$

This is illustrated in Fig. 3 (solid line) for a single dimension with the bias level, $\langle I \rangle^2$, subtracted out. The autocorrelation was normalized so the correlation values would range from 0.0 to 1.0. Each value on the solid line in Fig. 3 is transformed by finding the corresponding value in Fig. 2. This was done to form the three new curves (dotted lines) which represent the thresholded autocorrelations for different threshold values.

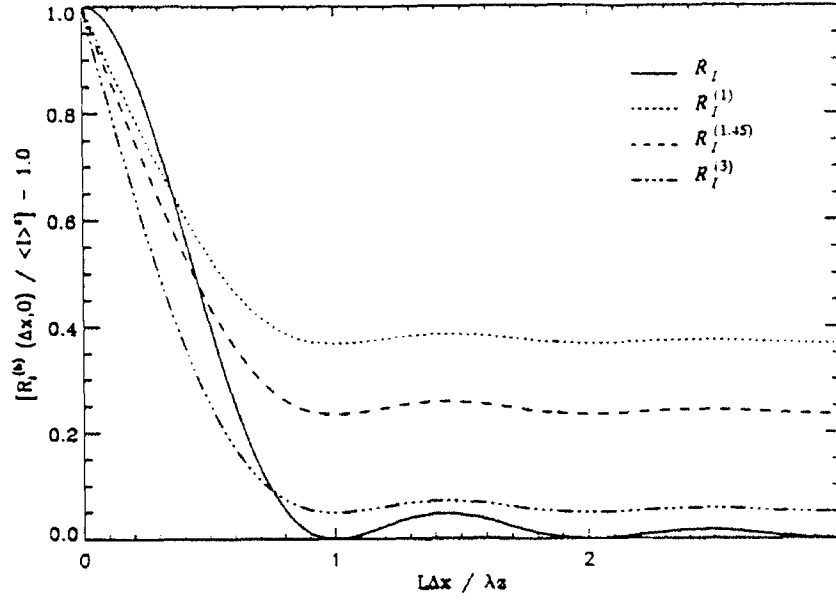


Fig. 3 Autocorrelation function of the speckle intensity resulting from a square aperture.

For the second case the amplitude of the field emanating from the double slit aperture is expressed as,

$$|P(\xi, \eta)|^2 = \text{rect} \frac{\xi}{l_1} \cdot \text{rect} \frac{\eta}{l_2} * \left[\frac{2}{L} \delta \left(\frac{2}{L} \xi \right) \right], \quad (2.8)$$

where l_1 and l_2 are the width and height, respectively, of each rectangle and L is their separation distance in x . The continuous autocorrelation is given by eq. (2.2),

$$R_I(\Delta x, \Delta y) = \langle I \rangle^2 \left[1 + \text{sinc}^2 \left(\frac{l_1 \Delta x}{\lambda z} \right) \cdot \text{sinc}^2 \left(\frac{l_2 \Delta y}{\lambda z} \right) \cdot \cos^2(\pi L \Delta x) \right]. \quad (2.9)$$

Once again, the thresholded autocorrelation can be determined using eq. (2.9) and the correlation mappings in Fig. 2. The resulting thresholded autocorrelations for three different threshold values are shown in Fig. 4. Only the Dx -axis is

plotted since the cosine term, in eq. 2.9, is in 'Dx'.

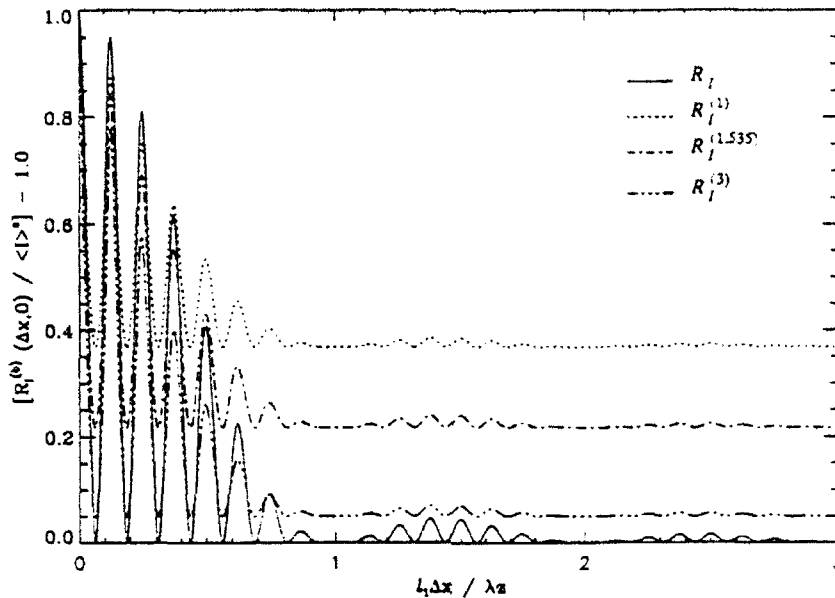


Fig. 4 Autocorrelation function of the laser speckle intensity resulting from a double slit aperture.

The predominant change seen in the transformation of continuous to thresholded autocorrelation is the addition of a correlation bias. This is due to a reduction in the uniqueness of the individual speckles. Consider the speckle intensity on a single dimension. The continuous speckles differ from each other in maximum intensity, length and the rate that the intensity rises to and falls from the maximum intensity. Since, each speckle is unique the correlation tappers off to zero for large separation distances. The thresholded speckle intensity appears like a 'boxcar' signal with each speckle represented by a single box. The speckles all have the same intensity value and only differ in length. Since, each box 'looks' like every other box, except for scaling, there will always be some bias correlation. This effect can be seen in Figs. 3 and 4. The amount of bias is equal to e^{-b} which can be determined by evaluating eq. 2.5 for $R_I(\Delta x, \Delta y) = 0$

POWER SPECTRUM OF THRESHOLDED LASER SPECKLE

The power spectral density is the frequency domain counterpart of the autocorrelation function. This relationship is described by the Wiener-Kinchine theorem⁴ and is expressed as,

$$S_I(v_x, v_y) = \mathcal{F} \{ R_I(\Delta x, \Delta y) \}. \quad (3.1)$$

The power spectral density describes the amount of power that exists at each spatial frequency in the laser speckle intensity distribution. Thresholding the speckle intensity may distort this distribution considerably if the correct threshold level is not chosen. An optimal threshold level occurs when the mean-squared error between the continuous and thresholded power spectrum is minimized. Where the continuous and thresholded power spectrums follow the terminology convention defined in the second section.

The continuous power spectrum for the square scattering area can be calculated by Fourier transforming eq. (2.7) which yields,

$$S_I(v_x, v_y) = \langle I \rangle^2 \left[\delta(v_x, v_y) + \left(\frac{\lambda z}{L} \right)^2 \Lambda \left(\frac{\lambda z}{L} v_x \right) \cdot \Lambda \left(\frac{\lambda z}{L} v_y \right) \right]. \quad (3.2)$$

This expression is the reference from which the mean-squared error will be determined. The reference curve is shown as a solid line in Fig. 5. To find the thresholded power spectrum the thresholded autocorrelation is Fourier transformed numerically. This was done for three different threshold values (Fig. 5, dotted lines). The impulse function at the origin in Fig. 5 comes from the correlation bias described at the end of the section 2.

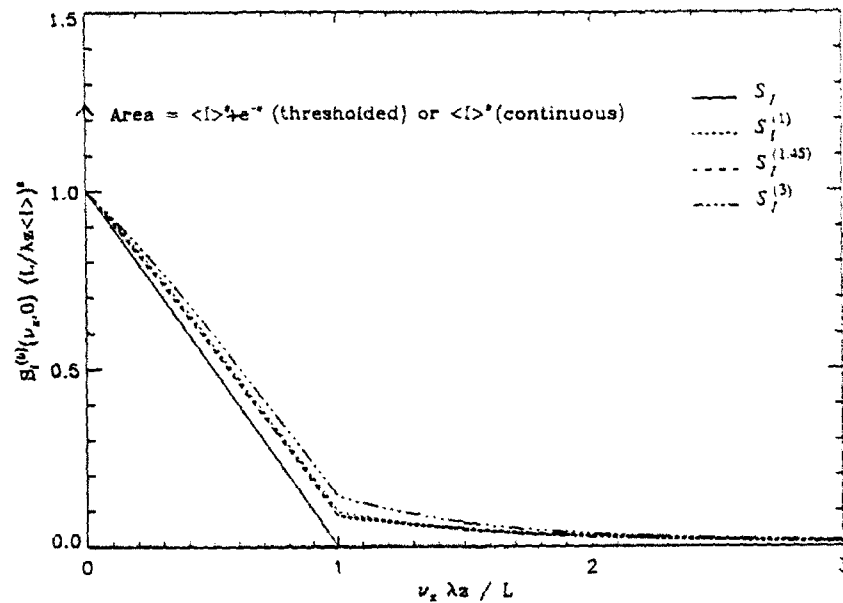


Fig. 5 Power spectral density of the laser speckle intensity resulting from a square aperture.

An analysis of Fig. 5 shows that the thresholded curves approach the shape of the continuous reference curve as the threshold level is increased from $b=1.0$ to $b=1.45$. As the threshold level is increased further, the thresholded curve moves away from the reference. To determine it exactly, the mean-squared error in the power spectral density from $\nu_x = 0$ to $\nu_x = 200L/\lambda z$ was calculated for threshold values between $b=1.0$ and $b=3.0$ (Fig. 6).

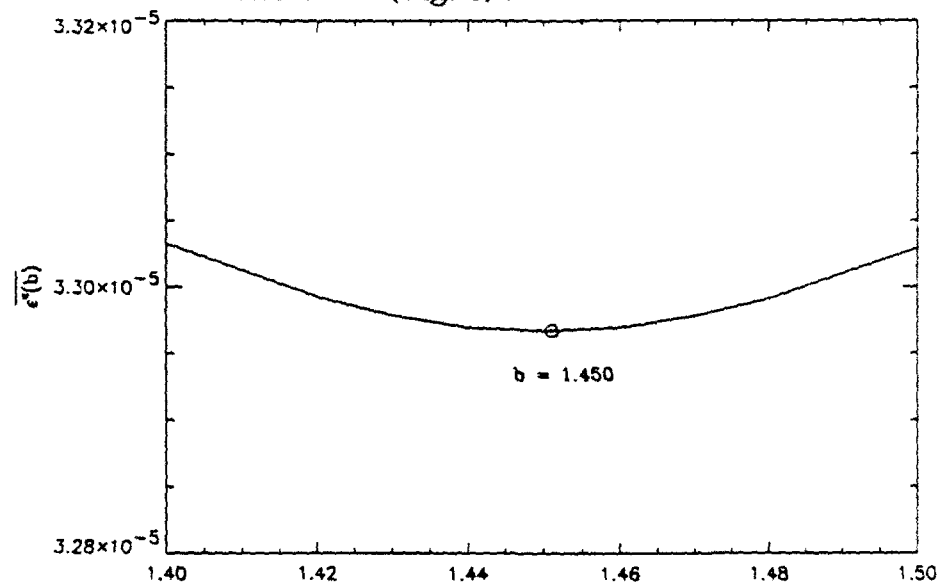


Fig. 6 Mean-squared error found in power spectral density of square aperture as a function of threshold value, b .

A value of $b=1.45$ was determined as the absolute minima and subsequent optimal threshold for the square aperture.

The same technique was applied to the double slit aperture. From eq. (2.9) the reference power spectrum is,

$$S_I(v_x, v_y) = \langle I \rangle^2 \left\{ \delta(v_x, v_y) + \frac{1}{2} \frac{(\lambda z)^2}{l_1 l_2} \Lambda \left[\frac{\lambda z}{l_1} v_x \right] \Lambda \left[\frac{\lambda z}{l_2} v_y \right] \right. \\ \left. + \frac{1}{4} \frac{(\lambda z)^2}{l_1 l_2} \Lambda \left[\frac{\lambda z}{l_1} \left(v_x - \frac{L}{\lambda z} \right) \right] \Lambda \left[\frac{\lambda z}{l_2} v_y \right] \right. \\ \left. + \frac{1}{4} \frac{(\lambda z)^2}{l_1 l_2} \Lambda \left[\frac{\lambda z}{l_1} \left(v_x + \frac{L}{\lambda z} \right) \right] \Lambda \left[\frac{\lambda z}{l_2} v_y \right] \right\}. \quad (3.3)$$

The thresholded power spectrums for three threshold values are given in Fig. 7. Although the lines are closely spaced, close inspection shows that the thresholded power spectrum exhibits the same behavior seen in the first case. The mean-squared error calculation in the power spectral density was calculated from $n_x = L/2lz$ to $n_x = 3L/2lz$. This limited range was used so that the optimal threshold value would be associated with the least amount of distortion in the outer triangle. The result was an optimal threshold value of $b=1.535$ (Fig. 8).

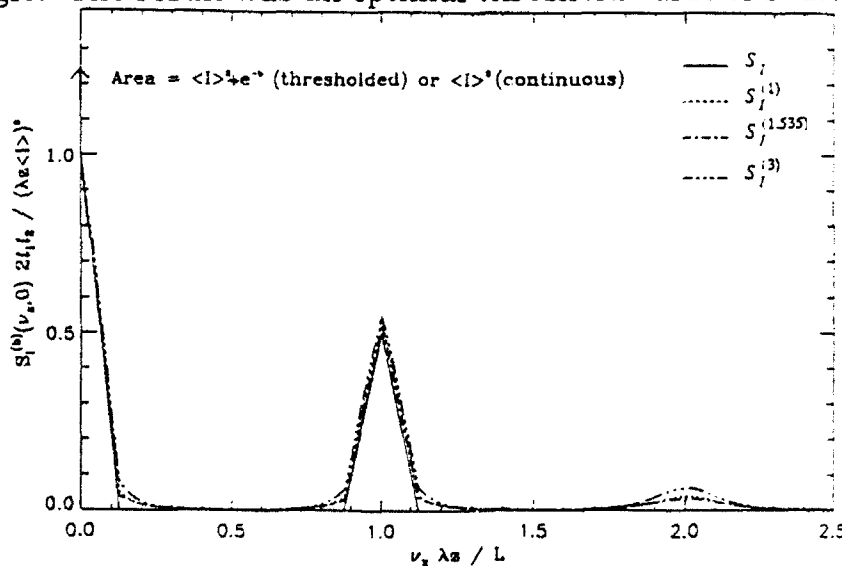


Fig. 7 Power spectral density of the laser speckle intensity resulting from a double slit aperture.

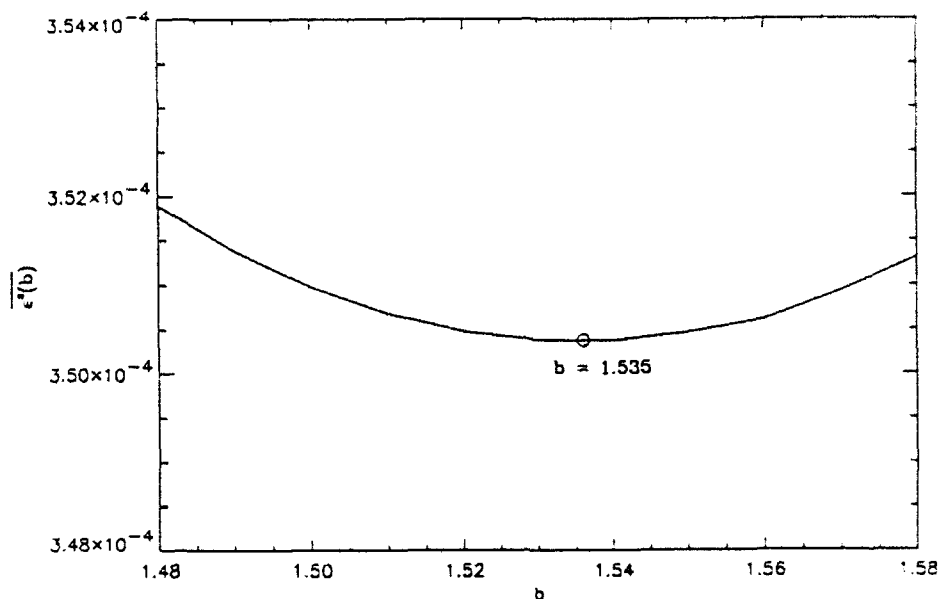


Fig. 8 Mean-squared error found in power spectral density of double slit aperture as a function of threshold value, b .

THRESHOLDED LASER SPECKLE SIMULATION

To verify the shape of the power spectral density for the optimal threshold values obtained in section 3, a Monte-Carlo simulation was performed. The simulation began by numerically propagating light, with unit magnitude and uniformly distributed $(-\pi, \pi)$ phase, from each generating aperture. This created a large continuous one-dimensional data record containing simulated laser speckle. The record was then separated into several segments of equal length. The power spectrum was calculated for each segment and averaged together to form an estimate of the continuous power spectrum for the entire record.

To calculate estimates for the thresholded power spectrums the original data record was first thresholded at the optimal threshold value. The resulting binarized record was segmented and the estimates were calculated as before.

The results of the simulation are given in Figs. 9 and 10. The raw data was

plotted as single points in both figures. The solid lines are the results determined in section 3 for the optimal threshold values. The data shows an increase of deviation with a decrease in frequency. This effect was expected since the record is of finite length and there are less low frequency speckles which can be used to form an average. A decrease in deviation is often seen in simulations of this type as the number of segments averaged is increased.

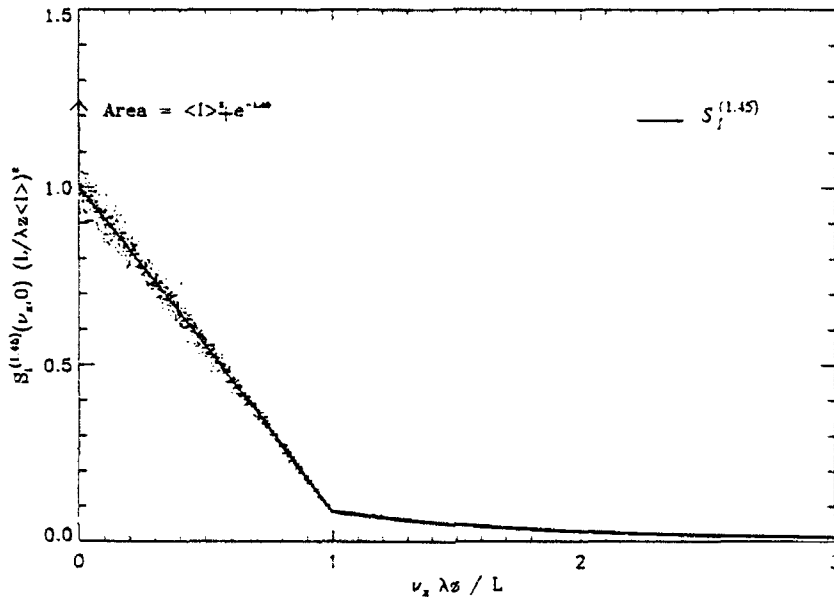


Fig. 9 Power spectral density of simulated laser speckle generated with square aperture and thresholded at $b = 1.45$.

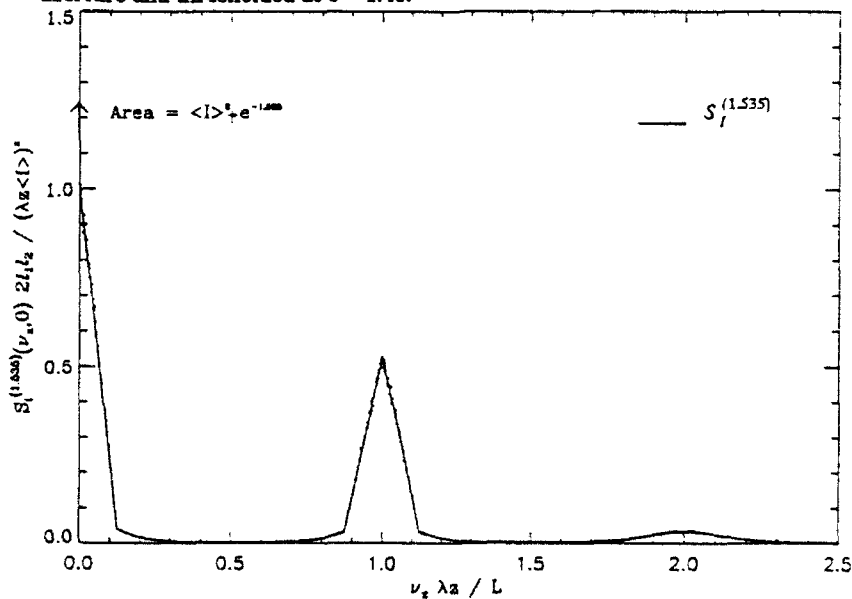


Fig. 10 Power spectral density of simulated laser speckle generated with double slit aperture and thresholded at $b = 1.535$.

CONCLUSIONS

The results presented here show that thresholding performed at the optimal level is a viable method for data reduction. We have shown that, for the MTF application, the important information is contained in the placement and size of the speckles and not in the intensity fluctuation between them. The thresholding operation preserves the important information and discards the rest which is the methodology behind all data reduction techniques. Implementation of this technique would result in an 8 to 1 savings in digital storage space.

REFERENCES

1. M. Sensiper, G. Boreman, A. Ducharme, "MTF Testing of Detector Arrays Using Narrowband Laser Speckle", accepted for publication in Opt. Eng.
2. R. Barakat, "Clipped Correlation functions of Aperture Integrated Laser Speckle", Appl. Opt. **25**, 3885 (1986).
3. J. Goodman, *Laser Speckle and Related Phenomena*, J.C. Dainty, ed., pp. 35-40, Springer-Verlag, Berlin (1975).
4. G.R. Cooper, C.D. McGillem, *Probabilistic Methods of Signal and System Analysis*, (2nd ed., Holt, Rinehart and Winston, Inc., New York, 1971), p.253.
5. L.C. Andrews, *Special Functions for Engineers and Applied Mathematicians*, (Macmillan Pub. Co., New York, 1985), p. 179.

USING X WINDOWS TO DISPLAY EXPERIMENTAL DATA

David E. Frink
Graduate Student
Department of Aerospace Engineering and Engineering Mechanics

University of Cincinnati
Cincinnati, Ohio

Final Report for:
Summer Research Program
Wright Laboratories, MLLP
Wright Patterson Air Force Base

Sponsored by:
Air Force Office of Scientific Research
Bolling Air Force Base, Washington, D.C.

September 1992

USING X WINDOWS TO DISPLAY EXPERIMENTAL DATA

David E. Frink
Graduate Student
Department of Aerospace Engineering and Engineering Mechanics
University of Cincinnati

Abstract

The following report describes the interactive colormap program written in C for X windows. The program is intended to be an alternative to a program already in existence which was not able to be executed with X windows. The purchase of X made a new program necessary.

USING X WINDOWS TO DISPLAY EXPERIMENTAL DATA

David E. Frink

INTRODUCTION

The raw data collected during an experiment often requires an easy method to view it. One method is by assigning colors to the data and representing the data as an image.

Which colors are best depend upon the image itself. Often the set of colors, the colormap, used for one image is not the best set for the next image. To simplify the color selection process, the image is first created using standard colors. An interactive program then allows the colors to be changed until the image best represents the data.

DISCUSSION OF PROBLEM

X windows were just purchased for Vax computers with a VMS operating system. The software previously used to view the Non-destructive evaluations is not compatible with the X windows, which either required the data to be transferred to a computer compatible with the old software or required a new program.

The original program displays an image based upon the experimental data. The colormap of the image can then be altered interactively. The colormap on this system is a set of 256 colors

where each color has a red value, a green value, and a blue value (rgb values.) The system and X windows requires some of the 256 colors leaving 235 colors for the program to use.

The previous software had several desirable features:

- Division of the colormap into two regions
- Derive a colormap if the two limiting colors are given
- Assign particular colors to particular values of the data
- Rotate the colors of the colormap

DIVISION OF THE COLORMAP INTO TWO REGIONS

The original program permits the division of the colormap into two separate colormaps. This allows a lower or upper range of data to be made distinct from the rest of the data. If noise exists it is often convenient to define a threshold the magnitude to must reach. In Figure 1, the data ranges from A to C but all values less than B are considered noise. The colormap has been divided into two regions. The first region is from A to B and is white. The second region is from B to C and is greyscale. Because all the noise is white, attention is placed upon the data range of interest.

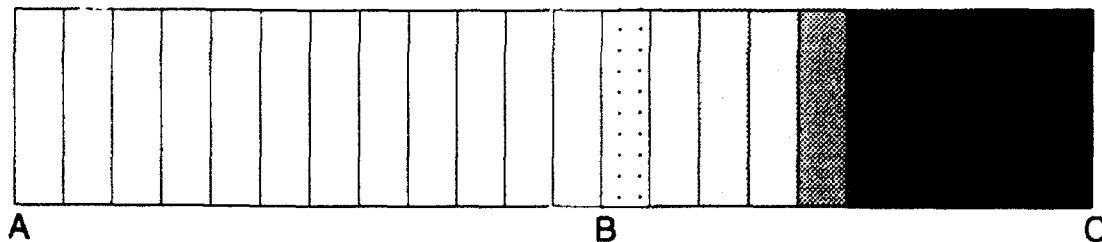


Figure 1 Division of the Colormap into Two Regions

DERIVE A COLORMAP IF THE TWO LIMITING COLORS ARE GIVEN

The colormap consists of 235 usable colors. If each of the 235

colors needed to be defined individually by the user, the program would become tedious. Instead, the first and last color of the colormap, or region of the colormap, is chosen and the program interpolates the remaining colors.

ASSIGN COLORS TO PARTICULAR VALUES OF THE DATA

It is convenient to be able to have the option of being able to define a single color. If a single value is of importance then a greyscale colormap can be chosen with the value of interest set to red. The red color is dominant within the grey colors and will show a particular value of the data.

ROTATE THE COLORS OF THE COLORMAP

The colormap can be rotated in the original program. Figure 2 shows a typical greyscale colormap. When the colormap is rotated by 4 colors to the left, the colormap is transformed to the colormap shown in Figure 3.

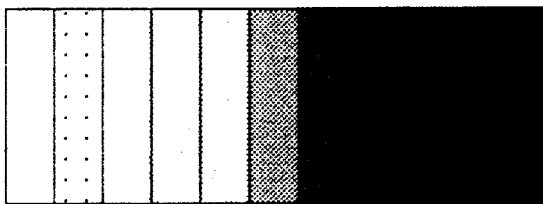


Figure 2 Standard Colormap

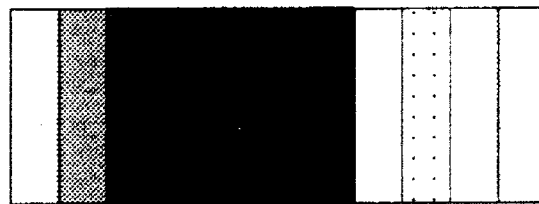


Figure 3 Rotated Colormap

METHODOLOGY

The goal, when creating the new program was to have as many of the features described above as possible, plus some new features. The first step was to create a program to create an image based

upon the experimental data.

The magnitude range of the experimental data is experiment dependent; it is not a fixed value. The first step is to scan the data and find the minimum and maximum values. After knowing the limits of the data, the data is scaled to fit integer values between 0 and 234, inclusive. Each value has been assigned a color on the colormap; 0 is the first color and 234 is the last color.

The image is now displayed using a program adapted by C. J. Fiedler.

The original program - before X windows was available - combined the viewing program and the colormap program. The programs have now been separated in order to easily place several images (sets of data) on the same colormap scale. The X window version reassigns the the default colormap of the workstation so as to be able to change all the images simultaneously.

Features of the new program:

- Multiple regions of the colormap (3 regions)
- Choice of predefined colormaps
- Saves set-up data for use later
- Invert the Colormap
- Noncontinuous colors
- Using stripes
- Rotate the Colormaps
- Print the Colormap

The program has six types of windows/Widgets. (A Widget is a higher-level window used in Motiff.) Two of the widgets are not accessible to the user, they include the parent widget and the form

widget. The parent widget is the interface between the program and the window manager. The form widget is a child of the parent widget and contains all the user accessible widgets.

The four user accessible widgets used, are listed in Table 1. The push button is activated when the mouse button is pressed while the mouse is on the button. The toggle button behaves like a push button except a toggle button changes status from on to off or off to on only when the mouse button is pressed. The push button changes status from off to on and back to off every time the mouse button is pressed. The slider is used for entering numerical data when the data has known upper and lower limits. The slider is position between the limits of the data. The final widget used is the Option Menu. When the mouse button is pressed on the category, a list of choices appears. The list is a collection of push button. This allows many push button to be available without requiring very much space.

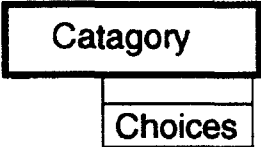
Key	
	Push Button
	Toggle Button
	Slider
	Option Menu ("Pull Down")

Table 1 Symbols for the Widgets Used

Windows or widgets activate a callback function (event handler.) The program is capable of passing predefined and user-defined variables. For example, when a scale widget is moved, the scale widget will execute its callback function. The callback function will be told whether the slider is being moved or has been moved depending upon the parameters of the callback. It will also pass the new location of the slider to the callback function.

When the callback function is finished, the main program waits for the next widget to be activated. The program only can do one widget at a time so the execution time of each callback must be short.

The program has the Widgets set up as shown in Figure 4. The push buttons: quit, print, save, and load are along the top of the window. The quit button calls an event handler which exits the

program. The print button calls an event handler which outputs four files: a list of all the red components of the colormap, a list of all the green components colormap, a list of all the blue components of the colormap, and a file containing the names of the other 3 files. The save and load widgets save, or load, the status of all the widgets so a particular set up can be stored for use again later.

The option menus titled "Colormap 1," "Colormap 2," and "Colormap 3," are a list of the colormaps the program has predefined. The reason for three categories of the same type is the program allows for the colormap to be divided into three regions. (Region 2 is the main region with a supplement colormap on either side of the main region.)

The option menus titled "# of Colors 1," "# of Colors 2," and "# of Colors 3" allow the predefined colormaps to be either smooth, distinct, or stripes. Smooth is when the colormap gradually changes colors. Distinct is when there are only 8 or 16 colors, each far enough apart on the color spectrum to distinguishable. Stripes is a special case of Distinct. Stripes defines every other color to be the same, such as black.

The invert buttons are of the toggle button type. Each of the three will reverse the order of the colors within each region.

The two vertical sliders define the dividing points between the three zones. The left slider gives the upper limit of region 1 and the lower limit of region 2. The right slider gives the upper limit of region 2 and the lower limit of region 3. Using the value x for the

left slider and the value y for the right slider the limits can be expressed as: $0 \leq x < y \leq 234$. Since region two is the main region, it must exist but the other two regions can be turned off by placing the sliders in the extreme positions. The program does not let the left slider have a value greater than the right slider.

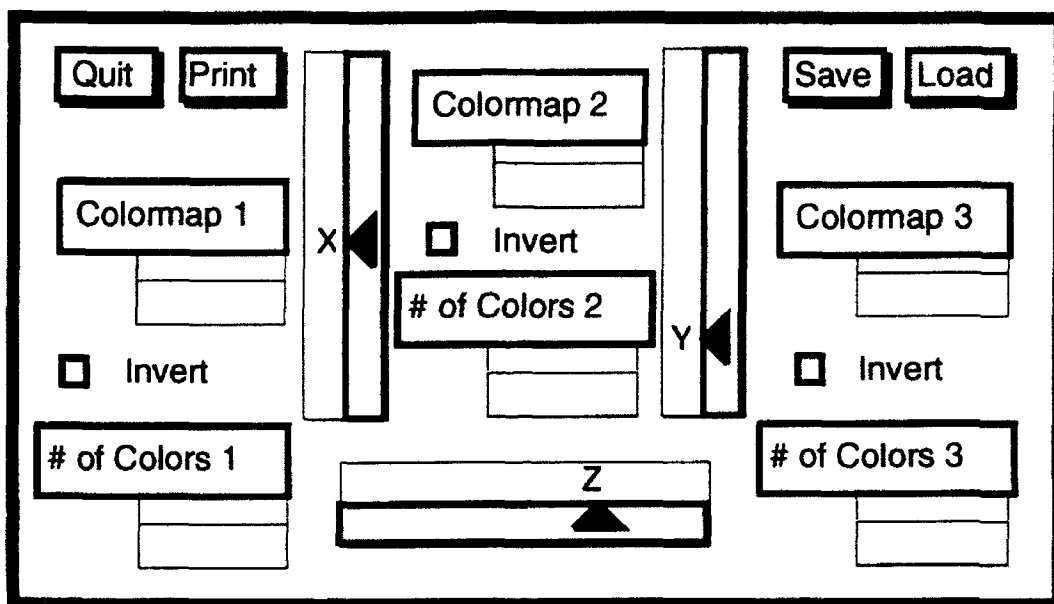


Figure 4 Set-up of Widgets

The horizontal slider along the bottom allows the colormap to be rotated. When the slider has a value of zero the colormap is not rotated. When the value (z) is non-zero, the colormap is rotated to the left by the amount z . This slider acts on all three regions as a group. If the slider is moved an amount less the size of colormap 1 ($z < x$) the colormap of Figure 5a will look like Figure 5b. If $z = x$, the first colormap has entirely rotated to the right side of the colormap, see Figure 5c. The colors can rotated completely - until the order has returned to normal.

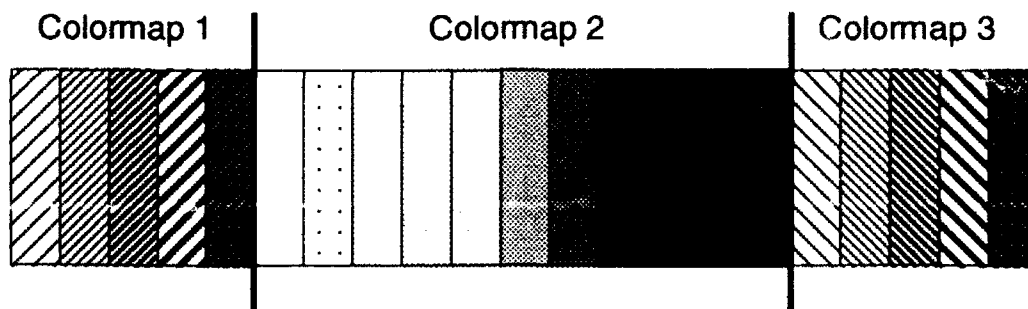


Figure 5a, Sample Colormap

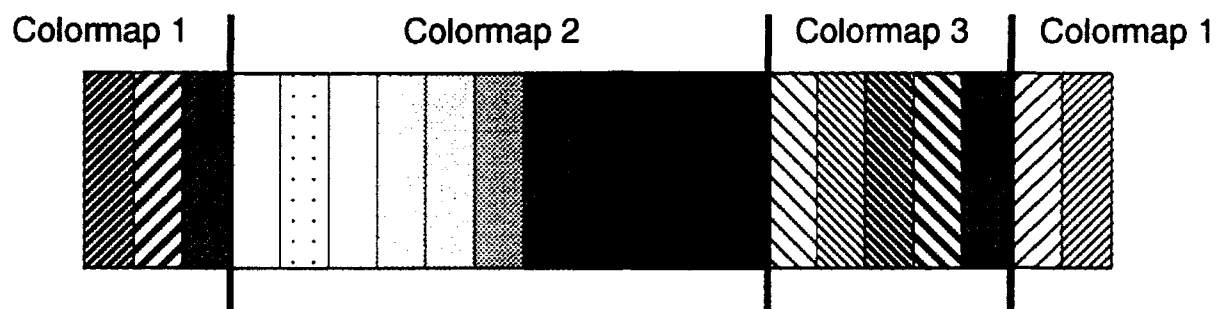


Figure 5b, Slight Rotation of the Colormap

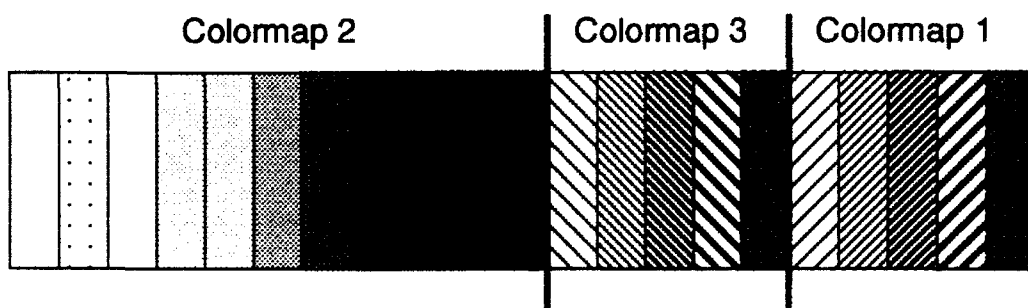


Figure 5c, Larger Rotation of the Colormap

This program also displays the current colormap by using the viewing program modified by C. J. Fiedler. The input of the experimental data has been replaced by a specific set of data which will show the colormap. The colormap program described above, simultaneously executes the specialized viewing case to show the current colormap.

RESULTS

The imaging and interactive colormap programs do have most of the same features as the original program, plus some additional features. These features are summarized in Table 2.

Features	Old Program	New Program
MULTIPLE REGIONS	✓	✓
INTERPOLATE COLORS	✓	
PARTICULAR COLORS	✓	
ROTATE COLORMAP	✓	✓
PREDEFINED COLORMAPS	≈ 5	20
INVERT THE COLORMAP		✓
NON-CONTINUOUS COLORS		✓
STRIPES		✓
PRINT IMAGE & COLORMAP	✓	✓

Table 2 Summary of Features

CONCLUSION

The new program is able to perform most of the capabilities of the old program. The other capabilities which were not done because of the decision to split the program into a viewing programming and an independent program for changing the colormap.

Because the program is written using X windows, the program

can be easily expanded. The only work needed is to write a function to perform the new feature and to assign this function to a Widget. (The new function will become the callback function of the Widget.)

**VLSI SYNTHESIS GUIDING TECHNIQUES
USING
THE
SOAR ARTIFICIAL INTELLIGENCE ARCHITECTURE**

**Joanne E. DeGroat
Assistant Professor
Department of Electrical Engineering**

**Lindy Fung
Graduate Student (Master's Degree)
Department of Electrical Engineer**

**The Ohio State University
205 Drees Laboratory
2015 Neil Avenue
Columbus, OH 43210**

**Final Report for:
Summer Research Program
Wright Laboratories**

**Sponsored by:
Air Force Office of Scientific Research
Bolling Air Force Base, Washington, D.C.**

September 1992

**VLSI SYNTHESIS GUIDING TECHNIQUES
USING
THE
SOAR ARTIFICIAL INTELLIGENCE ARCHITECTURE**

Joanne E. DeGroat
Assistant Professor
Department of Electrical Engineering
The Ohio State University

Lindy Fung
Graduate Student (Master's Degree)
Department of Electrical Engineer
The Ohio State University

The design of VLSI circuits is a very complex problem. As such, there are numerous computer aided design (CAD) tools to assist the designer in the generation of VLSI circuits. These tools range from layout editors, physical design tools, to high level synthesis tools. In all of these tools the design style and architectural structure are very constrained to limit the complexity to a manageable level. These constraints result in circuits that require more area and are slower than circuits produced manually by a skilled VLSI designer. The application of SOAR to the problem of VLSI creates the possibility of relaxing some of these constraints. The design of VLSI circuits is a design process. Existing tools and methodologies are incapable of capturing any of the essence of this design process and therefore must severely limit the design methodology. This research is a first step in the application of the SOAR Artificial Intelligence Knowledge Based Architecture developed at Carnegie-Mellon University to the synthesis of VLSI circuits. This first step is the application of SOAR to the VLSI placement problem. Findings, limitations, and recommendations for future research are presented.

A CHEMICAL INVESTIGATION OF THE OXIDATIVE BEHAVIOR
OF AVIATION FUELS

Ann Phillips Gillman
Research Associate
Department of Chemistry

Eastern Kentucky University
337 Moore Building
Richmond, KY 40475

Final Report for:
Summer Research Program
Research and Development Laboratories
5800 Uplander Way
Culver City, CA 90230

Sponsored by:
Air Force Office of Scientific Research
Wright Patterson Air Force Base, Dayton OH

September 1992

CHEMICAL INVESTIGATION OF THE OXIDATIVE BEHAVIOR OF AVIATION FUELS

Ann Phillips Gillman
Research Associate
Department of Chemistry
Eastern Kentucky University

Abstract

Methods to simulate time/temperature history of aircraft fuel systems have been studied. Jet fuel surrogates have been thermally stressed using a flask test. The resulting fuels were analyzed by gas chromatography-mass spectrometry (GC-MS) and gas chromatography-atomic emission detection spectroscopy (GC-AED). Deposits formed were tested for solvent characteristics, quantitatively measured and analyzed by thermal desorption and pyrolysis/GC-MS.

Through spectrometric analysis, two generalizations were found. Rapidly oxidizing fuels formed products high in low molecular weight aromatics, no higher than naphthalene. Their solids are primarily gummy and acetone soluble. Other fuels behave quite differently when stressed. These fuels oxidize slowly, forming flocculant, mainly acetone insoluble deposits. GC-MS analysis indicate these fuels to be high in phenols. GC-AED analysis of these different fuel types show the presence of sulfur in fuels that form insoluble deposits while "oxidizing", soluble deposit forming fuels do not indicate sulfur.

In general, the flask test produces results quickly and with repeatability. However, to adequately assess fuel stability, the availability of oxygen must be both limited and controlled. The general theory of oxidation of hydrocarbons is based on observed oxygen dependences. This theory has been slightly modified based on the presence of naturally occurring antioxidant molecules which are proposed to play an important role in both inhibiting the oxidation of fuels as well as being precursors to deposit formation.

A CHEMICAL INVESTIGATION OF THE OXIDATIVE BEHAVIOR OF AVIATION FUELS

Ann Phillips Gillman

INTRODUCTION

Fuel deposits are the limiting factor in the temperature to which a fuel can be heated. Deposits formed by thermal decomposition of fuel can clog filters, fuel feed arms, nozzle augments feeds and heat exchange tubes. Heat generated by advanced aircraft components such as avionics, environmental controls, hydraulics and engines must be exchanged to fuel. To accomplish this with no penalty to aircraft performance requires fuels with higher allowable operating temperature to increase the total heat capacity. An immediate goal of the U.S. Air Force is to increase the allowable operating temperature of fuel from 325F to 425F at little or no increase in fuel cost. It is assumed that at these temperatures, although pyrolytic mechanisms can become important at higher temperature, thermal oxidation is responsible for most fuel degradation and deposit formation.

The fuel/metal interface of the engine fuel nozzle is a significant area for deposit formation. Augmenter spray bars and rings can become completely blocked due to deposits forming in these areas. Figure 1a is a total ion chromatogram of a methylene chloride extract of actual deposits taken from a fuel tube of a fighter engine augmentor. This chromatogram indicates fuel component normal hydrocarbons ranging in composition from C_8 to C_{16} . This composition is typical of a JP-8 type fuel and can be considered to be fuel adsorbed onto the deposit matrix. This sample was taken from the spray ring of an augmentor. Figure 1b indicates the chemical nature of the deposits matrix. The fuel oxidation products of the solid sample are defined by pyrolysis-gas chromatography/mass spectrometry (PY/GC-MS) at 750C and shown in table 1. The solids are primarily substituted one and two ring aromatic compounds with a substantial amount of phenols.

Consumption of oxygen is the primary step in deposit formation. Fuels typically contain dissolved oxygen of 50 to 60 ppm. It is probably not feasible to completely eliminate all contact of oxygen and fuel. Coupled with heat, the presence of oxygen causes thermo-oxidative degradation in fuels, leading to deposit formation.

Elemental analysis of both laboratory and jet engine surface deposits have similar high oxygen levels. Nitrogen and sulfur content in sediments and surface deposits vary, but are also at higher levels than are present in bulk fuel. A number of measurements have shown, particularly by G.F. Bolshakov, that sediments typically contain low, e.g. 400 to 600, molecular weight compounds, indicating that oxidized dimers and trimers are involved. Studies have also shown that the quantity of sediment formed is less than the total quantity of oxidized fuel compounds. Numerous studies have been made of the low temperature liquid phase auto-oxidation of pure hydrocarbons of the type found in distillate fuels. Such reactions are typically free radical chain reactions where the initial products are hydroperoxides. Such knowledge would seem to indicate that the deposition tendency of a fuel would be directly proportional to the ease of oxidation of that fuel, but some fuels known to be "thermally stable" (as judged by deposition), such as JP-7 and JP-TS, oxidize more readily and extensively than other fuels which are oxidation resistant but easily form deposits.

The goal of this study is to better understand the chemistry that causes fuel deposits so that the knowledge can be applied to production of more stable fuels. This is accomplished by gas chromatography-mass spectrometry (GC-MS) for analysis of soluble and insoluble oxidation products of two reference fuels and a 12 component surrogate fuel. Thermal desorption and pyrolysis are used in conjunction with GC-MS to determine the composition of insoluble solid deposits.

The surrogate fuel was used to develop analytical methods and as a simple model for oxidation and deposition processes and thermal oxidative stress. The two reference fuels behave quite differently when stressed. One reference fuel, Propulsion Directorate Branch Fuel (POSF) 2747 (Sun-super K-1), is a highly hydro-treated fuel with a narrow boiling range having marker hydrocarbons ranging from C_{10} to C_{14} . POSF 2747 oxidizes readily and quickly, forming lacquer like, mainly acetone soluble deposits when stressed. The other, POSF 2827 (Shell-Jet A), is a non-hydro-treated fuel with a broader boiling range containing hydrocarbons from C_8 to C_{16} . POSF 2827 does not oxidize easily and forms large amounts of acetone insoluble, flocculent deposits. Figures 2a and 2b contrast the chemical nature of these two fuels. The larger chromatographic peaks in these total ion chromatograms of 2747 and 2827 are normal hydrocarbons (paraffins), with smaller peaks being polar fuel components. Analysis and observation of the properties of these fuels suggested doping studies that we believe should shed substantial light on the mechanism(s) of deposit formation and the ability to predict fuel deposition behavior.

EXPERIMENTAL

Materials

Fuels: Surrogate fuel (JP-8S); blended from Aldrich 99 + % compounds, its composition is listed in table 2. Table 3 is a comparison of the surrogate with a real petroleum derived JP-8 fuel. The reference fuels were two commercial grade Air Force aviation fuels (Shell Jet-A) and a thermally stable, hydrotreated fuel (Sun-Super K-1). Their chemical and physical properties as analyzed by SA-ALC/SFTLA, Wright Patterson AFB, OH, are listed in table 4a.

Doping chemicals: Phenol, 2-hexanethiol, phenylethylmercaptan, 3,4-dimethylphenol, 3-methylthiophene and all Aldrich reagent grade.

Solvents: acetone, heptane, methanol and petroleum ether of Aldrich reagent or HPLC grade.

Organic-sulfur standard: A mixture of organo-sulfur compounds was used to prepare a standard in which the limits of detection with mass spectrometry detection were investigated using a Hewlett-Packard gas chromatograph. Figure 3 includes chromatographic data, and table 5, Aldrich reference data for the following compounds: thiophene, 3-methylthiophene, tetrahydrothiophene, 1-hexanethiol, ethyldisulfide, butylsulfide, benzylmethylsulfide, phenylethylmercaptan, 3,4-dimethylthiophenol, phenylsulfide and benzyldisulfide. Four of the above sulfur compounds were used in doping experiments with JP-8S surrogate which are described later in this report.

Flask Stress Test Apparatus

In order to quickly induce thermal oxidation in our working fuels, a simulative flask test designed by Dr. William Schulz of Eastern KY University was employed. It's primary purpose was to simplify and accelerate the study of oxidation, inexpensively and with repeatability. Figure 4 is a schematic representation of the flask test apparatus. A; Approximately 200mL of fuel is contained in 250mL three-neck 24/40 round bottom boiling flasks. B; Refrigerated coolant is supplied via Friedrich condensers attached to the boiling flask by Claissen adaptors. C; The center opening of the flask and adaptor is used to position the thermometer and nitrogen/oxygen gas feed lines. Gas feed lines consist of 0.53mm Megabore^R deactivated, polyimide coated fused quartz capillaries inserted through Vitron-B^R thermometer seals and positioned to within 0.6mm of the bottom of the flask and parallel to the thermometer. D; Either inert nitrogen or oxidizing oxygen can be delivered to the stress vessel by way of a valved system. E; Boiling flasks are heated by fabric heating mantles, controlled by a 120 volt Therm-o-watch sensor slipped directly to the thermometer to maintain a temperature of 175C. G; Access to the fuel throughout the stress period is maintained by way of the third flask neck which is sealed with a 24/40 ground glass stopper. 10mL fuel samples were withdrawn directly from the stress flask by pasteur

pipette at set time intervals during the stress test. H; Fuel distillation products can be retained when the flask apparatus is modified with a 10mL Dean-Staerk trap placed between the Friedrich condenser and boiling flask. The Dean-Staerk trap was used to investigate the formation of a fuel insoluble phase produced only by POSF 2747 upon thermal stress.

Deposit Isolation

Thermal oxidation products of hydrocarbon fuels are categorized into three groups based upon the solvent characteristics of the oxidized molecules. 1: Soluble gums are those compounds that are soluble or perhaps suspended in fuel and do not precipitate. 2; Insoluble gums are compounds that are insoluble in fuel but soluble in polar solvents such as acetone. 3; Insoluble solids, which are fuel precipitates, are insoluble in acetone. Isolation of solid deposits was achieved using Antrop^R-25 plus 25mm, 0.2mm glass membrane filtration apparatus. Vacuum was produced through water aspiration.

Stressed fuel samples were filtered with Whatman^R 47mm GMF 150 grade binder free glass filters. Filtered fuel was saved for extraction. Filteres were washed with heptane and partially dried by maintaining air flow for 15-20 minutes. After difficulty arose in obtaining useful information from acetone soluble gum samples, "gums and solids" were dried *in vacuo* (20-30 torr, 80C) 36-48 hours and analyzed by thermal desorbtion and pyrolysis/GC-MS. Oxidation products were extracted from filtered fuel by solid phase and liquid-liquid extraction. Solid phase (S.P.E.) was by 1.0g silica gel cartridges, conditioned from methanol to heptane as per supplier recommendation. Then, 10mL of filtered fuel was forced dropwise through the S.P.E. cartridge. The cartridge was washed with 3 X 2mL portions of heptane, purged with 100mL air (syringe) and eluted with 2mL of acetone. Alternate extraction was 10mL of fuel, extracted with 3 X 2mL methanol with extracts pooled and back extracted with 3 X 2mL of heptane. Extraction was done in 12 X 150mL culture tubes, centrifuged and last traces of heptane was removed by pastuer pipette after centrifuging.

Chemical Analysis

Liquid fuel samples, containing soluble gums, were analyzed on the Hewlett-Packard 5890 Series II gas chromatograph with mass detection scanned from 35 to 550m/z. Samples were injected via HP 7673 auto-injector and electronic pressure programming set to maintain 30cm/sec carrier gas flow. Samples were carried through a 50M x 0.5mm DB-5 column. Dilute samples were injected splitless, up to 4mL, with the flow rate adjusted to 30cm/sec at 150C.

Acetone soluble and insoluble deposits were analyzed by thermal desorption and pyrolysis/GC-MS using a CDS model 1000 "Pyroprobe". It contains a resistively heated platinum filament pyrolyzer which heats samples held in a quartz tube. The pyroprobe is interfaced to the GC by means of a heated chamber which houses the filament rod during pyrolysis. This method produces highly resolved chromatographic peaks of distinct deposit components, which are not otherwise chromatographable. Figure 5 represents the optimum peak resolution that can be obtained from pyrolysis/GC-MS at 950C. This total ion chromatogram is of acetone soluble deposits from, JP-8S surrogate fuel, stressed for 46 hours at 175C.

RESULTS AND DISCUSSION

Chemistry of the Flask Test

Figure 6 shows the development of insoluble solids and gums formed by the surrogate fuel and the two working fuels at constant time (46 hours) and temperature (175C), with an oxygen flow rate of 100mL/min. These fuels, under identical stress conditions, exhibit very different oxidative characteristics. POSF 2747 forms minimal amounts of insoluble solids. POSF 2827 behaves much differently, forming extensive amounts of insoluble solids and minimal insoluble gums. Figure 7 represents the amount of total deposits formed by these fuels during various times throughout a 46 hour stress period. Although POSF 2747 forms more deposits overall than POSF 2827, the majority of these deposits are acetone soluble. 99.26% of the total deposits formed by 2747 are soluble in acetone, leaving only 0.73% being insoluble, while an alarming 57.33% of the total deposits formed by POSF 2827 are acetone insoluble. Table 4b indicates that the amount of deposits formed by POSF 2747 are primarily acetone soluble. POSF 2747 oxidizes more easily than POSF 2827. When total deposits formed by these fuels after 46 hours of stress were compared, 2747 produced 2.59g (1.65%) and 2827 formed 0.75g (0.50%) of deposits. This perhaps suggests that fuels which oxidize easily are invariably stable when measured by filterable deposits.

Upon thermal stressing, fuels produce highly polar oxidation products in the form of organic acids, alcohols, aldehydes, esters and ketones. Solid phase and liquid extraction techniques were implemented to concentrate and separate these polar oxidation products from non-polar fuel components. Both extraction methods effectively separated and concentrated oxidation products from the fuels. GC-MS chromatograms of extracted fuel samples display an array of highly resolved peaks of oxidized fuel components. Figures 8a and 8b, total ion chromatograms of

stressed POSF 2747 and 2827 fuels, illustrating this optimum peak resolution. These particular samples were stressed for 46 hours, extracted with methanol, then back extracted with heptane.

The noticeable difference between the two fuels are shown by the chromatograms in figures 9a and 9b with peak identifications listed in tables 6 & 7. The soluble oxidation extract of 2827 in figure 9a is high in phenols and furanone derivatives. Figure 9b, methanol extract of 2747, is high in alkenes and furanone derivatives. This may suggest that the presence of phenols, as in 2827, is indicative of the chemical initiator for the formation of measurable deposits. GC-Atomic Emission Detection (AED) analysis (Figures 10 & 11) verify the presence of sulfur in 2827 while 2747 contains no detectable sulfur. Sulfur atoms appear to be significantly involved in the chemical reactions which take place under stressing. The sulfur atoms in stressed 2827 decreased by more than a factor of 2. Sulfur atoms tend to concentrate in the deposits as oxidation progresses.

Figure 12 is a total ion chromatogram of insoluble gum produced from POSF-2827 and figure 13 is the chromatogram of insoluble gum produced from POSF-2747 with peak identifications in tables 8 and 9. There is probably some degree of thermal rearrangement of deposit compounds involved in this method of analysis, but it is the first method that has succeeded in identification of deposit composition.

Thermal desorption chromatograms of the two fuel deposits are quite consistent with the chromatograms of the soluble oxidation products. The 2827 deposit contains a great many phenols and few alkenes or alcohols. The 2747 deposit contains alcohols and alkenes, with few phenols. Both deposits contain (or yield) furanone derivatives. These must be oxidation products of alkanes (paraffins) but it is unclear as to what role they play in deposition if it is not that they are simply polar, fuel insoluble compounds that will aggregate and precipitate.

The result of the flask test as an accurate and precise measurement of thermal stability depended either upon the amount of oxygen supplied or the method of solid collection. Determination of solid formation with time was tedious due to the limited amounts of solids available to work with. Oxygen saturation was necessary to obtain measurable amounts during an accelerated test period.

Doping Experiments

The contribution of phenolic compounds to deposit formation is significant. Results from doping experiments of phenolic compounds in JP-8S surrogate and POSF-2747, stressed under previously described conditions, produced deposits similar to those of sediment forming fuels like

POSF-2827. Figure 14 illustrates the stability changes in JP-8S surrogate as a result of addition of 2.0% phenol.

A mixture of phenolic compounds, similar to those found in POSF-2827 were added to JP-8S surrogate. The results were similar to those of the previous experiment; doping with phenol. The phenolic mixture containing 2,4-dimethylphenol, 2-sec-butylphenol, 2-tert-butylphenol and 2-propylphenol (1.0% total volume) was added to ~200mL of JP-8S surrogate. Figure 15, as well as observations that phenol alone, illustrates results in increased deposit formation, suggesting that quantity and type of phenolic compound present in fuel may be relevant, and that their presence does play a significant role in deposit formation.

In an effort to determine the effect of sulfur-containing compounds on the stability of fuels, we doped JP-8S surrogate with a mixture of organo-sulfur compounds containing sulfides and thiols, in addition to phenols, and subjected the fuel to previous flask stress conditions. To ~200mL of surrogate fuel, the following were added: 1-hexanethiol, phenylethylmercaptan, 3,4-dimethylthiophenol and 3-methylthiophene (1.94% total volume) along with 1.0% phenolic mix. The average percentage of deposits increased significantly from the previous stress test with phenol alone, particularly acetone insoluble deposits. These insoluble solids increased by 48.66% as shown in table 10. These results indicate that doping stable fuels with sulfides and thiols significantly reduce fuel stability and add to deposit formation as shown in figure 16.. Sulfur compounds play a significant role in jet fuel stability. That exact role is still uncertain. They may play a special role on interaction with other compounds. Phenols are strong oxidizing activators for the benzene ring. This characteristic perhaps allows sulfur to sulfate the ring, making it insoluble to polar solvents.

CONCLUSION

1. The flask test does not model any real aircraft fuel environment, but as an oxidative system, it, provides fuel oxidation products very similar to real engine deposits quickly and reproducibility.
2. GC/MS and pyrolysis GC/MS proves to be effective chromatographic techniques for the separation and identification of component fuel matrices.

3. GC/MS analysis of deposits from non-hydrotreated (insoluble deposit forming) fuels, indicates high levels of phenols. These fuels contain naturally occurring antioxidants in the form of phenols and anilines.
4. Sulfur containing compounds are found to be present in fuels that form large amounts of insoluble deposits like POSF-2827.
5. Doping experiments using phenolic and organo-sulfur compounds with oxidizable fuels modeled by JP-8S surrogate, change their oxidative behavior, causing them to develop characteristics of unoxidizable (non-hydrotreated) fuels like POSF-2827. This doping results in the formation of large amounts of acetone insoluble deposits by fuels that form mainly acetone soluble deposits when stressed clean.
6. Oxygen consumption is a fundamental variable in deposit formation. Hydrotreated fuels like POSF-2747 consume oxygen at a rapid rate, therefore oxidizing and forming products quickly and then leveling-off. Non-hydrotreated fuels like POSF-2827 slowly consumes oxygen, hence, these fuels gradually oxidize, forming deposits at a steady rate.

Figure 1a: Total Ion Chromatogram of Methylene Chloride Extract of Fuel Tube Deposits.

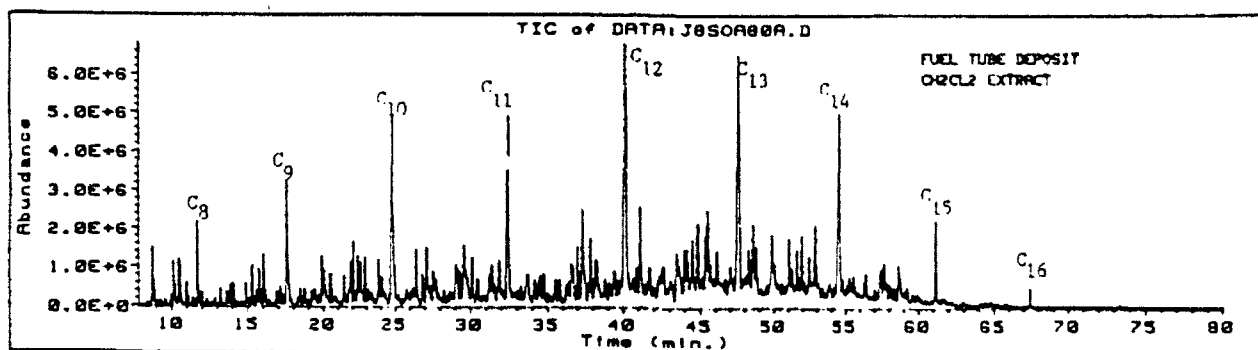


Figure 1b: Total Ion Chromatogram of Pyrolysis @750C of Insoluble Solids from Spray Ring Deposits.

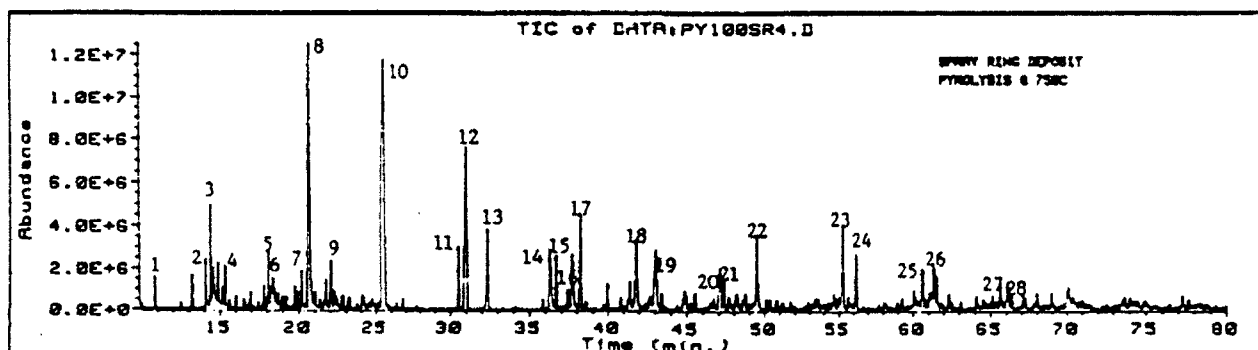
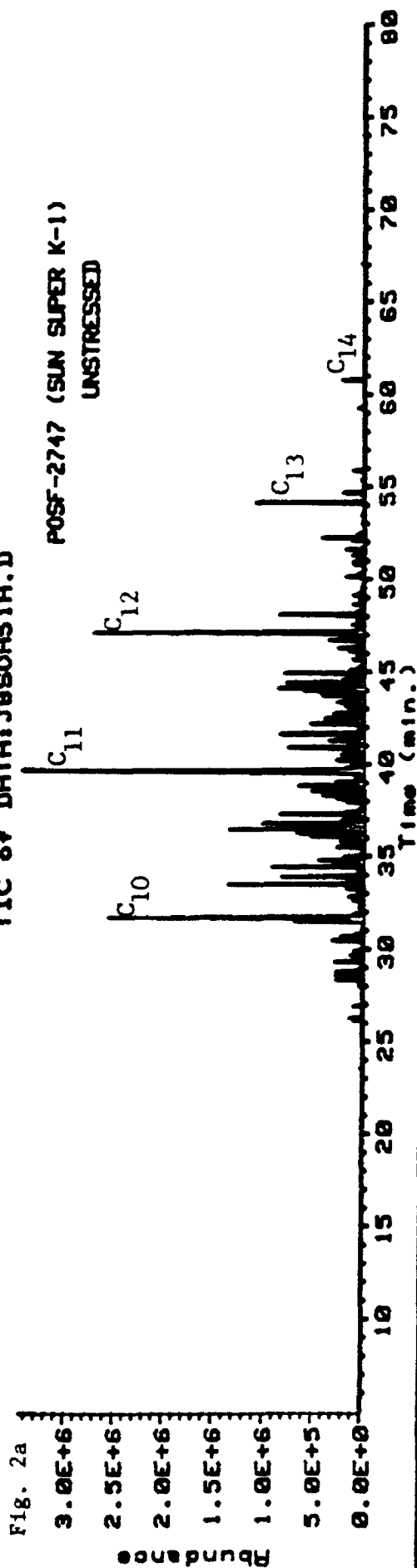


Table 1. Identification of Chromatographic Peaks for Figure 1b (PY100SR4.D). Pyrolysis of Insoluble Solids from Spray Ring Deposits @750C.

Peak #	Rt(min.)	Compound
1.	10.86	2-butene
2.	14.08	C5 alkene
3.	14.41	C5 alkene
4.	14.89	C5 alkene
5.	18.11	hexane
6.	18.38	4-methylpentane
7.	20.16	C6 alkene
8.	20.72	methylcyclopentene
9.	22.06	heptane
10.	25.62	toluene
11.	30.53	ethylbenzene
12.	31.09	m,p-xylene
13.	33.41	o-xylene
14.	36.31	1-ethyl-2-methylbenzene
15.	36.72	1,2,4-trimethylbenzene
16.	37.74	phenol
17.	38.27	1,2,5-trimethylbenzene
18.	41.86	2-methylphenol
19.	43.10	3-methylphenol
20.	47.08	2,4-dimethylphenol
21.	47.36	2-methylindene
22.	45.54	naphthalene
23.	55.32	2-methylnaphthalene
24.	56.23	1-methylnaphthalene
25.	60.67	dimethylnaphthalene
26.	61.39	dimethylnaphthalene
27.	65.85	naphthalenol
28.	66.11	mixed naphthalenol and trimethylazulene

TIC of DATA:J080A51A.D



TIC of DATA:J080A52A.D

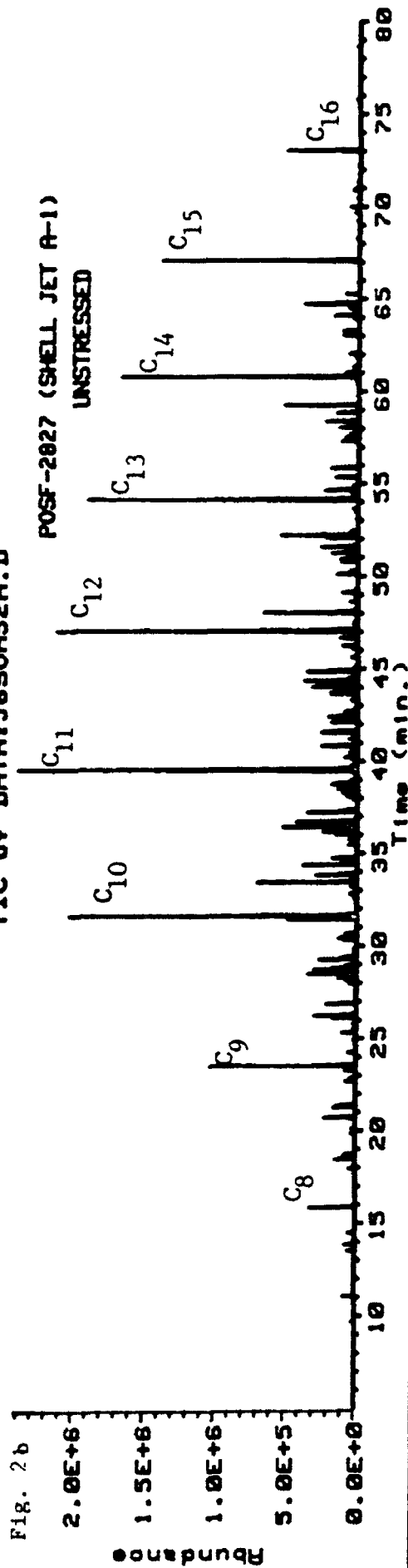


Table 2: Aldrich and Chromatographic Data of JP-8S Surrogate Fuel

Peak#	Ret. Time	Compound	Wt%	D ₁₅ (cc)	B.P. (C)
1	9.07	isooctane	5.0	0.690	98
2	10.27	methylcyclohexane	5.0	0.770	101
3	15.98	m-xylene	5.0	0.868	138
4	18.40	cyclooctane	5.0	0.934	151
5	21.71	decane	15.0	0.730	174
6	23.38	butylbenzene	5.0	0.860	183
7	25.58	1,2,3,4-tetramethylbenzene	5.0	0.838	197
8	27.06	tetralin	5.0	0.973	207
9	28.74	dodecane	20.0	0.749	215
10	31.87	1-methylnaphthalene	5.0	1.001	240
11	34.74	tetradecane	15.0	0.763	252
12	40.01	hexadecane	10.0	0.773	287

Table 3: Comparison of Surrogate and Authentic JP-8 Fuels

ASTM#	Test	Surrogate		JP-8	
		WPAFB	Tank 3-15	JP-8	JP-8
D1319	Aromatics, Vol %	22.0	22.0	22.0	22.0
D1319	Olefins, Vol %	3.7	3.7	3.7	3.7
D2887	Distillation Initial B.P. C	92	124	92	124
D2887	Distillation 10% Recovered C	135.0	160	135.0	160
D2887	Distillation 20% Recovered C	169.0	173	169.0	173
D2887	Distillation 50% Recovered C	205.0	212	205.0	212
D2887	Distillation 90% Recovered C	255.0	259	255.0	259
D2887	Distillation End Point, C	286.0	296	286.0	296
D1298	Density, Kg/L	0.8	0.811	0.8	0.811
D93	Flash Point, Deg C	26.0	59	26.0	59
D2386	Freezing Point, Deg C	-14.0	-54	-14.0	-54
D445	Viscosity @ 20C, cS	3.9	3.9	3.9	3.9
D3338	Net heat of Combustion, MJ/Kg	43.1	43.1	43.1	43.1
D3343	Hydrogen Content, Wt%	13.7	13.6	13.7	13.6

Table 4a: Comparison of Analysis for POSF 2747 and POSF 2827 Fuels.
(Data from SA-ALC/STTLA, WPAFB, OH 45433-6503)

ASTM#	Test	Result's	
		POSF 2747	POSF 2827
3242	Total Acid Number, mg KOH/g	0.0	0.001
1319	Aromatics, Vol %	19	19
3227	Mercaptan Sulfur, Wt %	0.000	0.001
4294	Sulfur, Total Wt %	0.0	0.1
93	Flash Point, Deg C	60.0	50.0
1298	Specific Gravity, 15.6/15.6 Deg C	0.8076	0.8072
2386	Freezing Point, deg C	-60	-4.3
445	Viscosity @ 20 Deg C, cS	4	5
1322	Smoke Point, mm	22	24
130	Copper Strip Corrosion	1	1
3241	Thermal Stability @ 260, Deg C DELTA P, mm	0	0
381	Existent Gum, mg/100ml	0	1
1094	Water Reaction Interface	1	0
2624	Fuel Electrical Conductivity, pS/m	230	183
5327	Fuel System Icing Inhibitor, Vol %	0.00	0.00

Table 4b: Comparison of Deposits Produced by POSF 2747 and POSF 2827 Fuels in Flask Tests at 175C and 100ml/min Oxygen Flow

Mass of Deposit Wt % of Fuel	Time for Visible Deposit	Result's	
		POSF 2747	POSF 2827
		2.849g	0.8964g
		1.650%	0.0597%
	over 2 hours		
	less than 10 minutes		
Physical Nature	brown, adherent gum	dark brown, flocculant solid	
Gums and Solids	Insoluble Gum	Solid	Insoluble Gum
Mass	2.5855g	0.0191g	0.3197g
Wt % of Fuel	1.8400%	0.0121%	0.3660%
Melting Point C	90 onset	183 onset	doesn't

FIGURE 3 GC/MS Chromatogram of Organo-sulfur Mix 1.

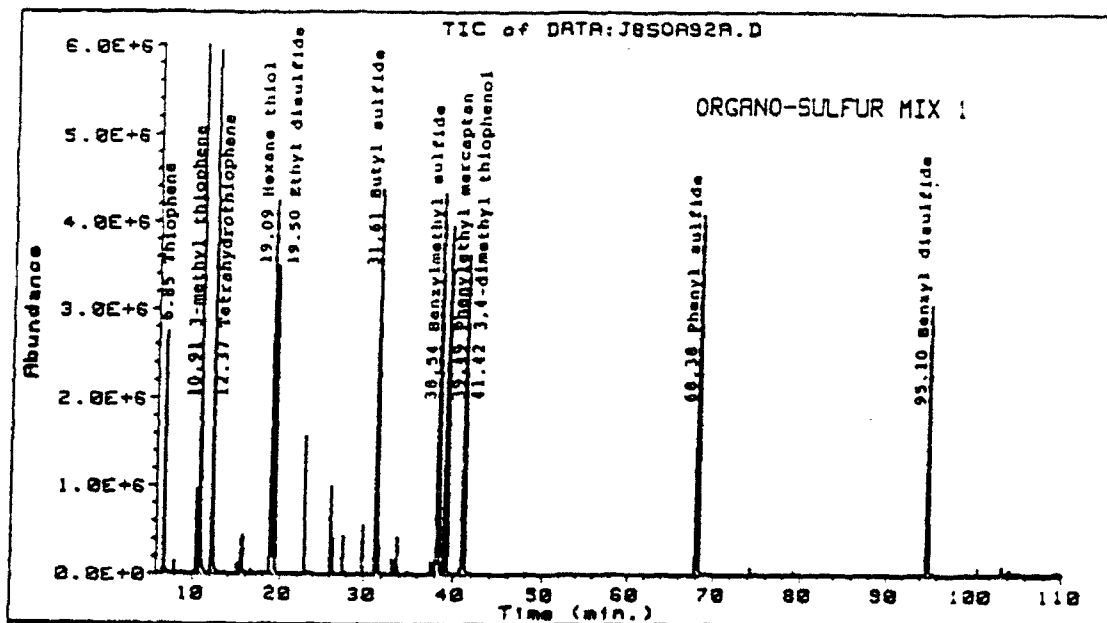
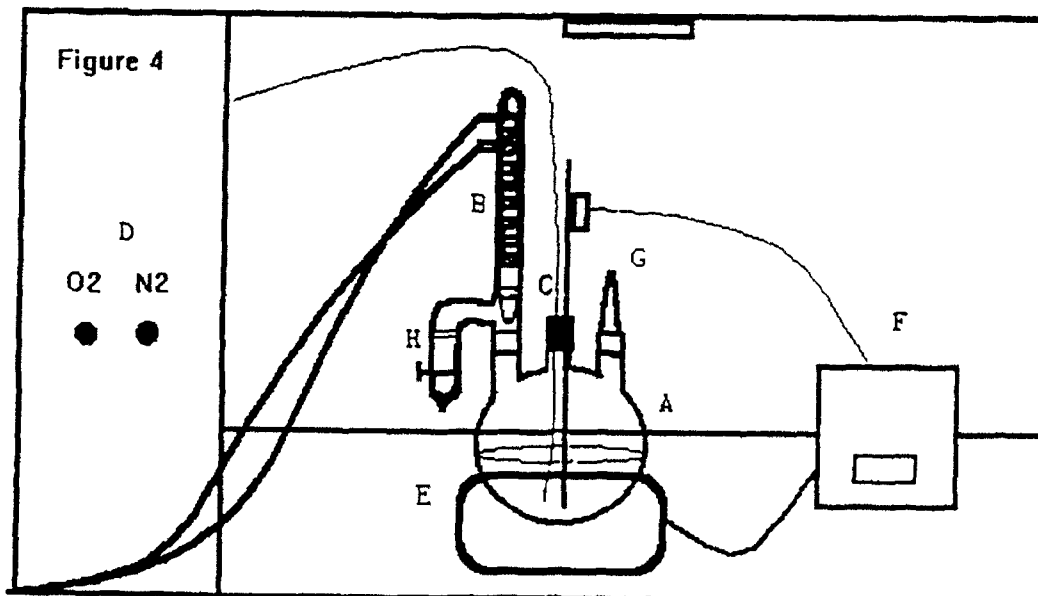


Table 5: Chromatographic and Physical Data of Organo-sulfur Mix 1.

Compound	t_R (min)	B.P. (°C)	M.P. (°C)	F.W. (g)	CAS#
Thiophene	6.85	84	-38	84.14	110-02-1
3-methyl thiophene	10.91	114	-69	98.17	616-44-4
Tetra-hydro thiophene	12.37	119	-96	88.17	110-01-0
1-hexane thiol	19.09	150-154	-81-80	118.24	111-31-9
Ethyl disulfide	19.50	153	-	122.25	110-81-6
Butyl sulfide	31.61	188-189	-	146.30	544-40-1
Benzylmethyl sulfide	38.54	195-198	-	138.23	766-92-7
Phenylethyl mercaptan	39.49	217-218	-	138.32	4410-99-5
3,4-dimethyl thiophenol	41.42	218	-	138.32	18800-53-8
Phenyl sulfide	68.38	296	-	186.28	139-66-2
Benzyl disulfide	95.10	-	66-69	246.39	150-60-7



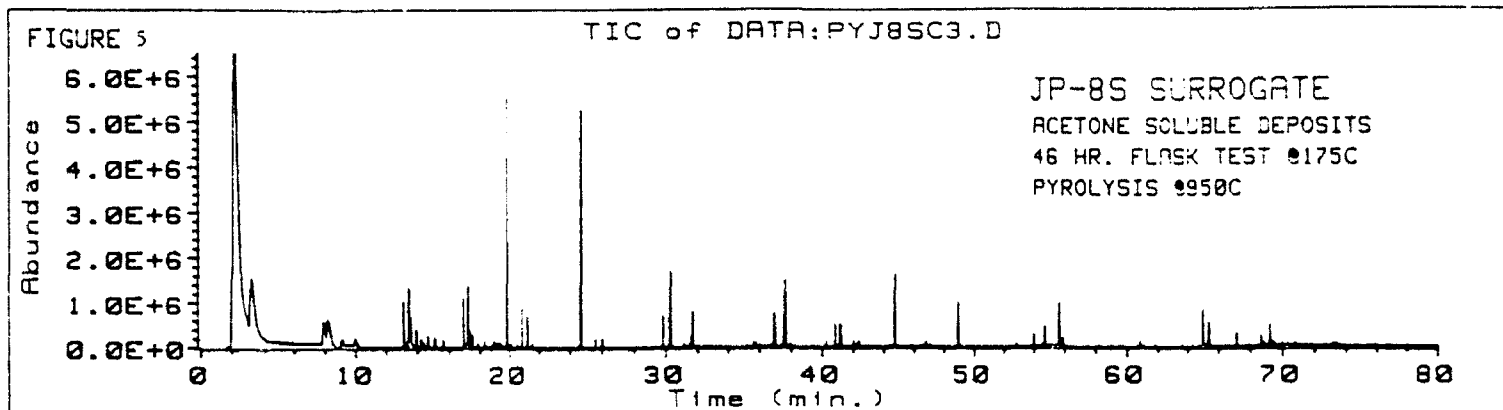


FIGURE 6. GRAVIMETRIC DATA FOR SURROGATE AND WORKING FUELS. 46 HR FLASK TEST @175C & O₂ @100mL/min.

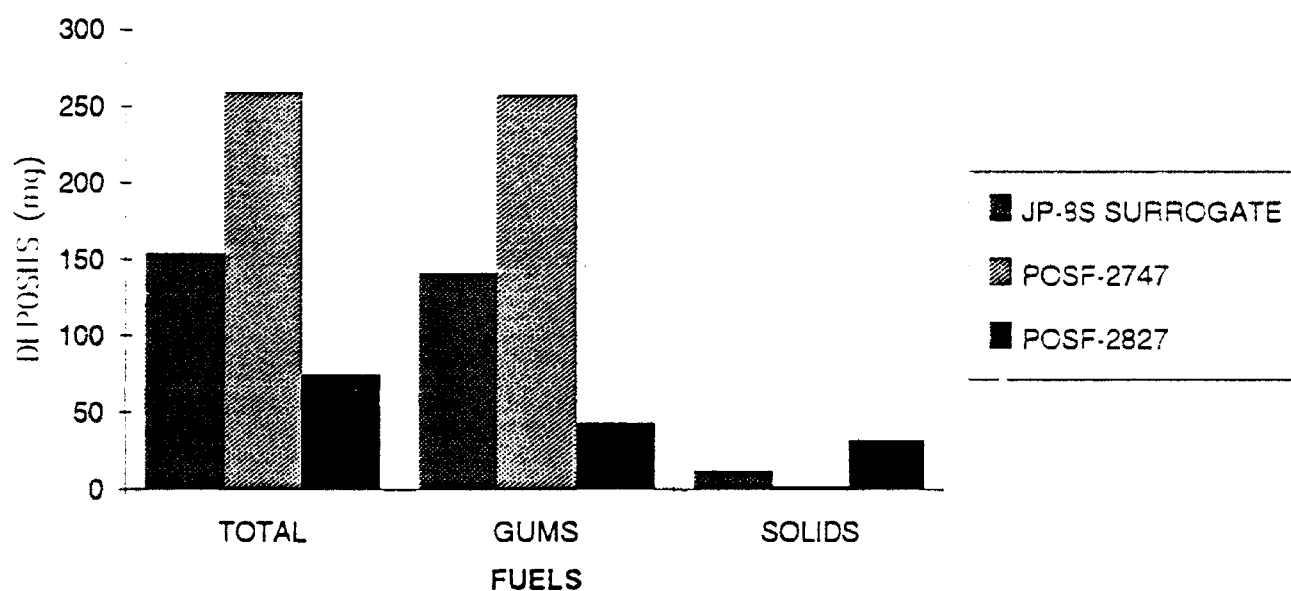
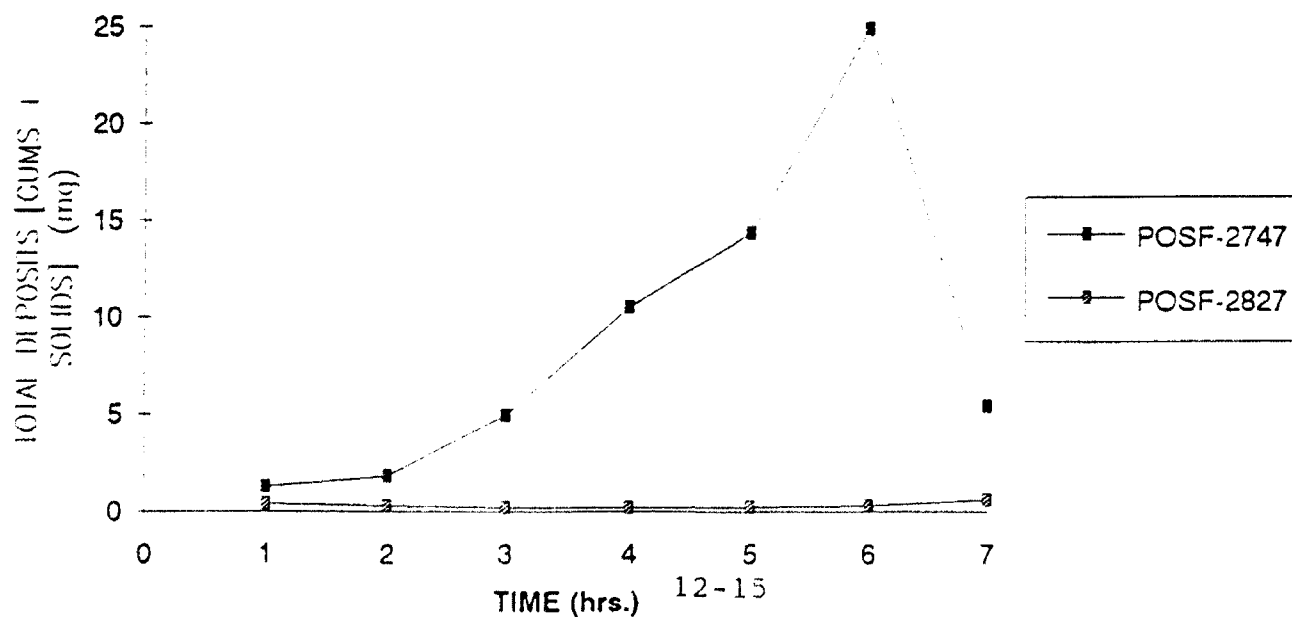


FIGURE 7. TOTAL DEPOSITS (mg) FORMED BY POSF-2747 & POSF-2827 AT VARIOUS STAGES DURING 46 HR. FLASK TEST @175C.



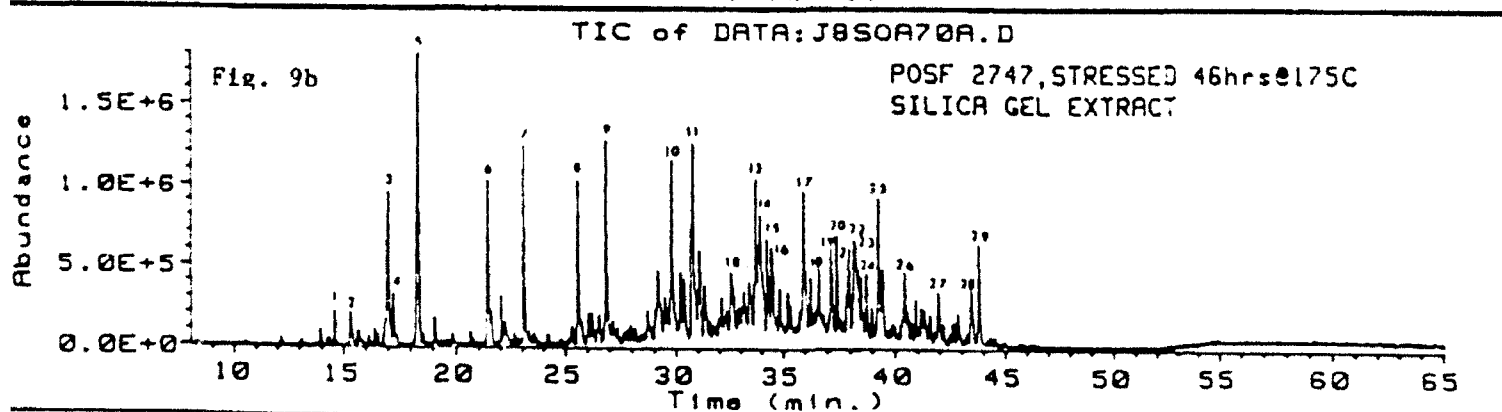
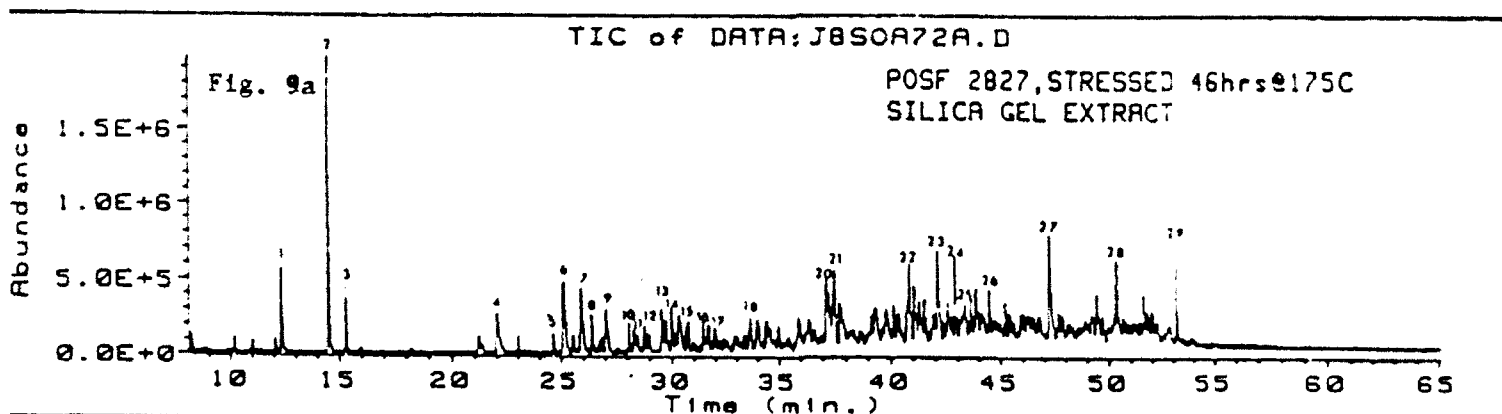
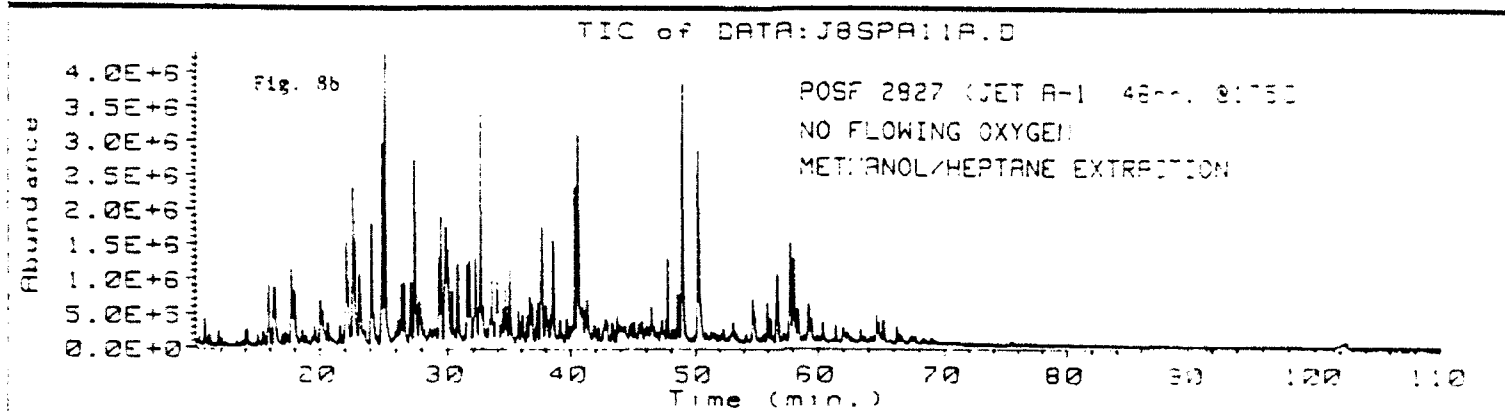
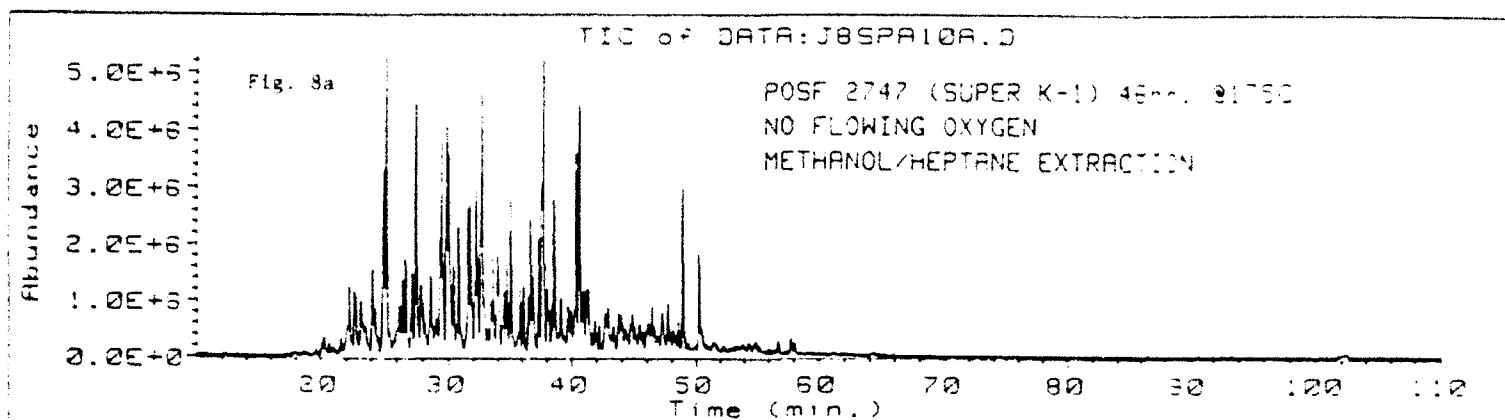


Figure 9a and 9b. Total Ion Chromatograms of Oxidation Products Extracted from Fuels Stressed for 46 hrs. with 100 mL/min. Oxygen Sparge. Figure 9a- is POSF 2827 extract, peak identification is table 6. Figure 9b- is POSF 2747 extract, peak identification in table 7.

Table 6: Peak Identification for Figure 9a. Extracted Oxidation Products of POSF 2827.

1.	12.43	1 propoxypentane
2.	14.59	3,3-dimethyl-2-hexanone
3.	15.34	2,2-dimethylpentanol
4.	22.08	4-methylphenol
5.	24.66	ethylphenol
6.	25.18	dimethylphenol
7.	25.98	dimethylphenol
8.	26.43	dimethylphenol
9.	27.09	dimethylphenol
10.	28.07	2?-propylphenol
11.	28.31	propylphenol
12.	28.78	propylphenol
13.	29.58	ethylmethylphenol
14.	30.03	trimethylphenol
15.	30.26	C ₆ phenol
16.	31.46	C ₆ phenol
17.	31.70	C ₆ phenol
18.	33.70	subst. tetrahydrofuranone
19.	34.45	decanol
20.	37.19	methylisobenzofurandione
21.	37.57	undecanol
22.	40.90	mixed spectra, subst. benzoic acid?
23.	42.17	dimethylbenzopyran-2-one
24.	42.95	alkene
25.	43.69	p-cyclohexenylphenol?
26.	44.52	methylnaphthoquinone?
27.	47.31	subst. phenol
28.	50.38	methoxyphenanthrene?
29.	53.20	subst. furandione?

Table 7: Peak Identification for Figure 9b. Extracted Oxidation Products for POSF 2747.

1.	14.59	3,3-diaethyl-2-hexanone
2.	15.32	dihydro-2(3H)-furanone
3.	17.06	dihydro-5-methyl-2(3H)-furanone
4.	17.24	tetrahydro-2H-pyran-2-one
5.	18.37	4,4-dimethyl-dihydro-2(3H)-furanone
6.	21.43	subst. furanone
7.	23.10	5-ethylidihydro-5-methyl-2(3H)-furanon
8.	25.59	5-propyl-dihydro-2(3H)-furanone
9.	26.86	3-ethyl-2,5-furandione
10.	29.83	5-butylidihydro-2(3H)-furanone
11.	30.80	2-undecanol
12.	32.55	1,3-isobenzofurandione
13.	33.76	5-pentylidihydro-2(3H)-furanone
14.	33.97	sec-butylethylbenzene
15.	34.28	1-ethyl-3-(1-methylethyl)-benzene
16.	34.49	subst. 3,5-furandione
17.	35.96	4-methylisobutylfurandione
18.	36.62	subst. furanone
19.	37.19	4-methylisobenzofurandione
20.	37.45	5-hexylidihydro-2(3H)-furanone
21.	37.97	methyl-1(3H)-isobenzofurandione
22.	38.20	methyl-1(3H)-isobenzofurandione
23.	38.29	methylisobenzofurandione
24.	38.75	2,4,6-trimethylphenyl-1-ethanone
25.	39.29	4-methylphthalic acid
26.	40.42	5,5-diaethyl-3(2H)-benzofuranone
27.	41.95	dimethylbenzofuranone
28.	43.41	ethylmethylbenzofuranone?
29.	43.79	ethylmethylbenzofuranone?

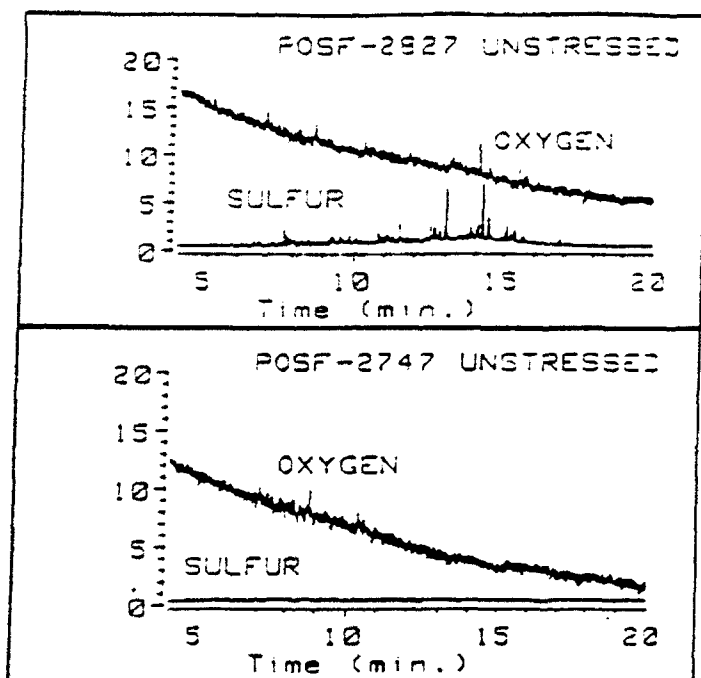


Figure 10: GC/AED of Unstressed POSF-2747 & 2827.

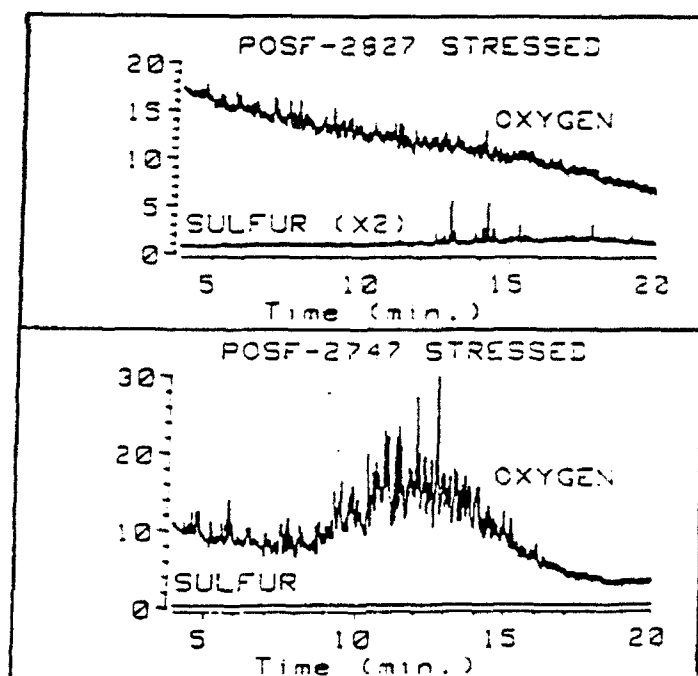


Figure 11: GC/AED of Stressed POSF-2747 and POSF-2827. 46 hr. Flask Test @175C.

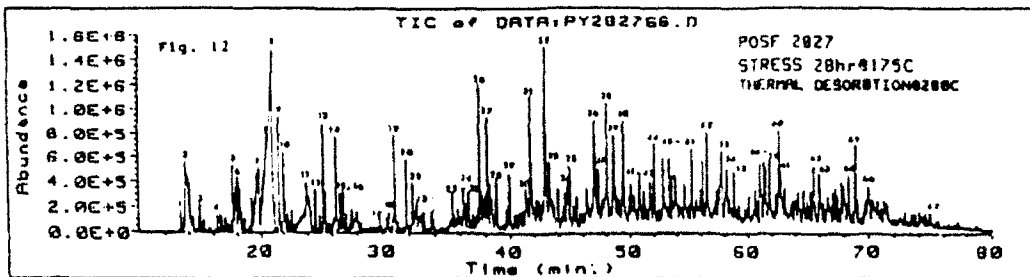


Figure 12: Total Ion Chromatogram of Insoluble Gum Formed by POSF 2027 Stressed 28 hours at 175C with 100 mL/min Oxygen Spurge. Peak identification is table 11.

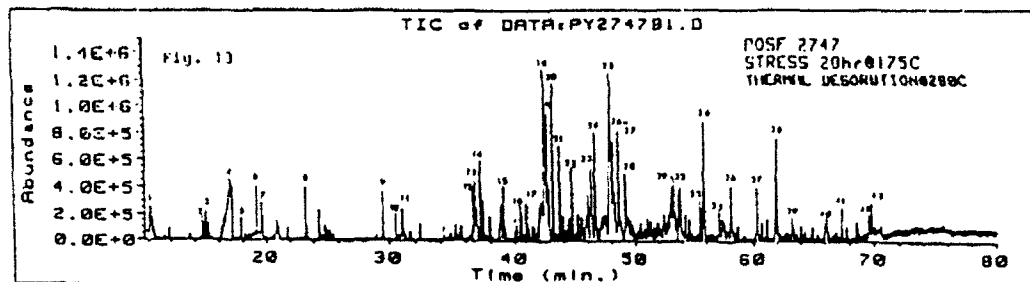


Figure 13: Total Ion Chromatogram of Insoluble Gum Formed by POSF 2747 Stressed 28 hours at 175C with 100 mL/min Oxygen Spurge. Desorption at 280C. Peak identification is table 12.

Table 1: Peak Identification For Figure 12: Thermal Desorption Chromatogram of Stressed POSF 2027 Insoluble Gum

1.	13.34	1-butene
2.	13.74	butane
3.	14.88	diethylcyclopropane
4.	15.34	aethylpentene
5.	17.56	1-hexene
6.	18.06	2-butanone
7.	19.38	mixed MS
8.	21.06	acetic acid
9.	21.63	diethylpentene
10.	22.01	3-methylbutanal
11.	23.99	propanoic acid
12.	24.71	aethylpentanol
13.	25.39	toluene
14.	26.38	2-heptenal
15.	26.69	dihydropyran
16.	26.91	sec-alcohol
17.	29.76	aethylcyclopentanone
18.	30.52	ethylbenzene
19.	30.97	a & 2 - xylene
20.	31.89	1-octanol
21.	32.38	o-xylene
22.	32.69	incomplete MS
23.	35.43	ene-ol
24.	36.33	ethylaethylbenzene
25.	36.71	ethylaethylbenzene
26.	37.63	phenol
27.	38.27	triethylbenzene
28.	38.98	ethylhexanol?
29.	40.02	isopropylbenzene
30.	41.37	2-hydroxybenzaldehyde
31.	41.80	2-aethylphenol
32.	43.03	4-aethylphenol
33.	43.32	sec-alcohol
34.	44.88	dimethylphenol
35.	45.05	2-aethylbenzofuran
36.	47.39	dimethylphenol
37.	47.48	aethylindene
38.	48.15	diethylphenol
39.	48.72	diethylphenol
40.	49.52	diethylphenol
41.	50.17	C ₆ phenol
42.	50.88	C ₆ phenol
43.	51.78	C ₆ phenol
44.	52.14	2-hydroxybenzeneacetic acid
45.	52.84	C ₆ phenol
46.	53.42	C ₆ phenol
47.	53.63	C ₆ phenol
48.	53.89	C ₆ phenol
49.	54.87	1H-indene-1-one
50.	56.05	ethylaethylphenol
51.	56.24	aethylnaphthalene
52.	56.62	isobenzofurandione
53.	57.85	aethylbenzofuranone
54.	58.28	aethylbenzoic acid
55.	58.87	dimethylphenyl-1-ethanone
56.	61.10	dimethylbenzofuranone
57.	61.43	diethylphenylbenzene
58.	61.62	diethylphenylbenzene
59.	61.91	dimethylbenzofuranone
60.	62.74	4-aethylphthalic acid
61.	62.88	2,2-diethyl-3,4-dihydro-1(2H)-1-benzopyran
62.	65.60	aethylnaphthalenol
63.	66.10	1-phenylfuran?
64.	68.47	subst. phenol?
65.	69.04	9H-fluorene
66.	70.06	2-aethyl-1-naphthalenol
67.	74.96	diethylphenylbenzene

Table 2: Peak Identification For Figure 13: Thermal Desorption of Stressed POSF 2747 Insoluble Gums

1.	10.24	butane
2.	14.58	diethylcyclopropane
3.	14.95	2-butanone
4.	16.31	acetic acid
5.	17.33	benzene
6.	19.13	3-butene-2-one
7.	19.56	heptane?
8.	23.28	toluene
9.	29.52	xylene
10.	30.51	octanol?
11.	31.01	xylene
12.	36.56	phenol?
13.	36.72	ene-ol?
14.	37.20	ethylaethylbenzene
15.	39.03	ethylaethylbenzene
16.	40.40	3-methylcyclohexene?
17.	40.94	aethylisopropylbenzene
18.	42.42	2-undecene
19.	42.68	4?-undecene
20.	43.16	7-undecene
21.	43.71	undecene
22.	44.70	C ₆ benzene
23.	46.30	hydroxybenzaldehyde?
24.	46.61	aethyldimethylbenzene
25.	47.92	1-dodecene
26.	48.13	7-dodecene
27.	48.61	7-dodecene
28.	49.15	undecanol?
29.	53.10	dodecanol?
30.	53.25	dimethylphenyl-1-ethanone
31.	53.69	alkene
32.	53.81	2,3-dihydro-1H-indene-1-one
33.	55.45	dimethylphenylethanone
34.	55.83	1,3-isobenzofurandione
35.	57.07	1H-indene-1,3(2H)-dione
36.	58.09	aethylphenylbutanedione
37.	60.29	dimethylbenzofuran
38.	61.96	aethylphthalic acid
39.	62.22	subst. isobenzofurandione
40.	66.14	dimethylbenzofuranone
41.	67.37	dimethylisobenzofurandione
42.	69.53	isobutylisobenzofuranone
43.	69.79	alkane

FIGURE 14. COMPARISON OF DEPOSITS FORMED FROM PHENOL DOPED JP-8S SURROGATE FUEL AND UNDOPED FUEL

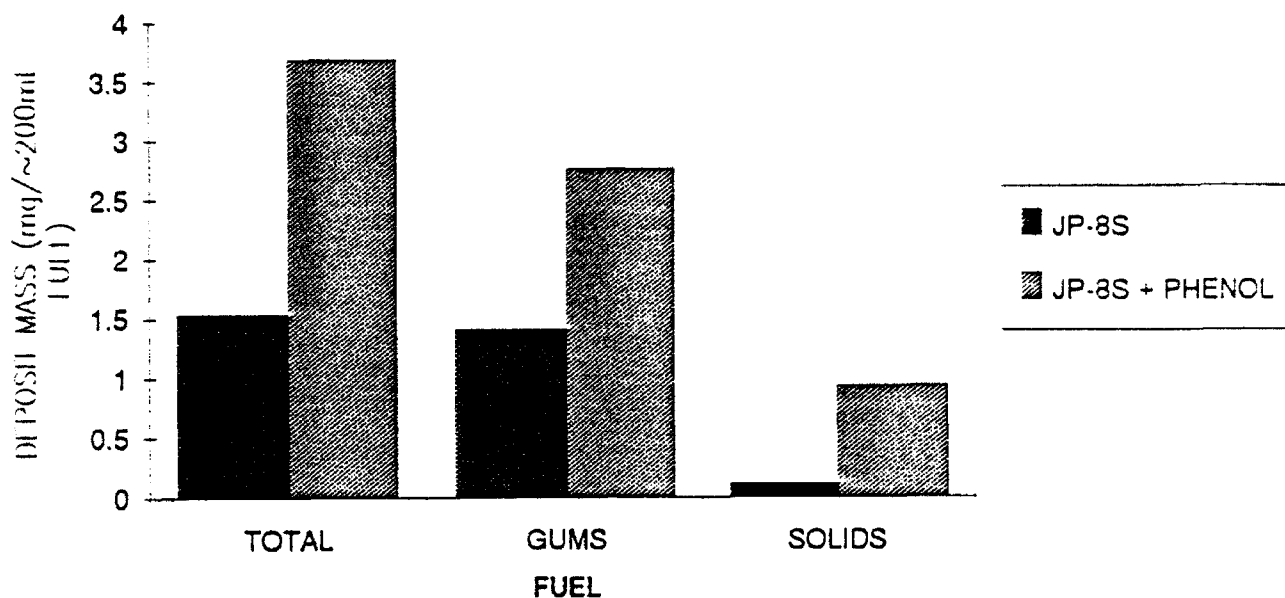


FIGURE 15. COMPARISON OF DEPOSITS FORMED FROM PHENOL AND PHENOL COMBINATION DOPING OF JP-8S SURROGATE FUEL AND UNDOPED FUEL

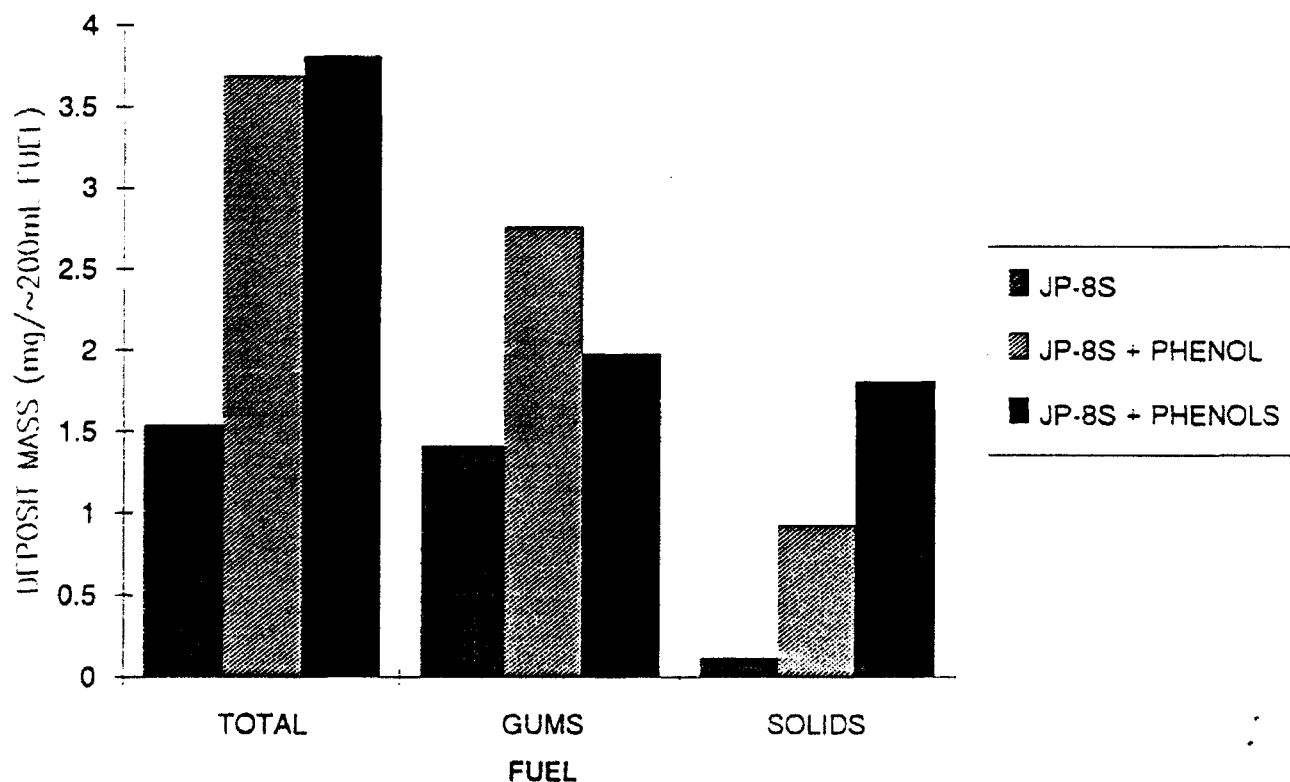


FIGURE 16. COMPARISON OF DEPOSITS FORMED FROM PHENOL AND PHENOL/SULFUR COMPOUND DOPING OF JP-8S SURROGATE FUEL.

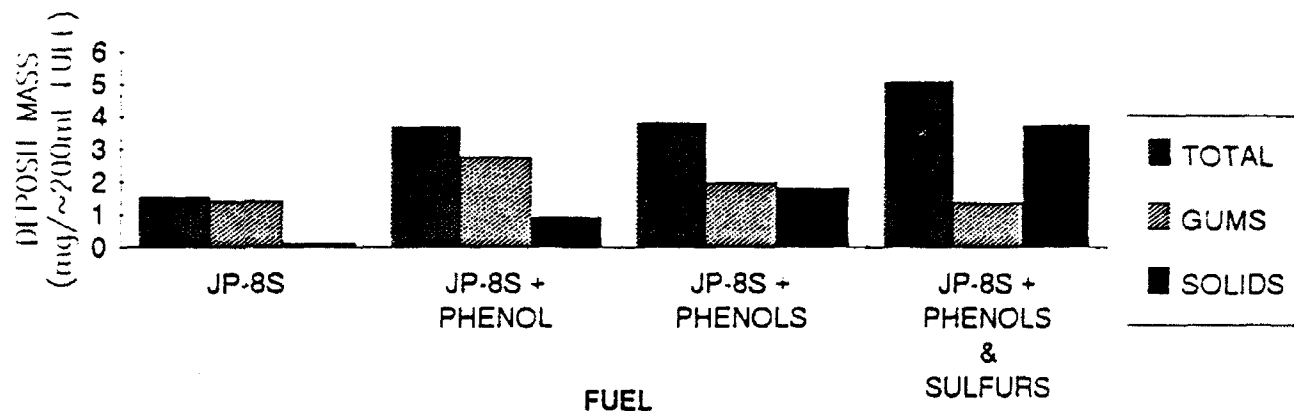


TABLE 10. GRAVIMETRIC ANALYSIS OF JET FUELS. 46 HOUR FLASK TEST @175C WITH O₂
@100mL/min.

FUEL	TOTAL	MASS (mg/~200mL FUEL)	
		GUMS	SOLIDS
JP-8S	15.40	14.10	1.20
JP-8S + 2% PHENOL	36.90	27.60	9.30
JP-8S + 1% PHENOLS	38.10	19.80	18.10
JP-8S + 1% PHENOLS & 2% SULFUR COMPOUNDS	50.70	13.50	37.20

FINITE ELEMENT ANALYSIS OF INTERLAMINAR
TENSILE TEST SPECIMENS

Diane Hageman

Summer Graduate Student
University of Cincinnati

Final Report for:
AFOSR Summer Research Program
Wright Laboratory
Materials Directorate

Sponsored by:
Air Force Office of Scientific Research
Bolling Air Force Base, Washington, D.C.

September 1992

Abstract

To use carbon-carbon composites correctly, the interlaminar strengths and stiffnesses must be measured experimentally. The test procedures, however, can cause stress concentrations within the test specimens that make measurements inaccurate. To help in designing test specimens that produce uniform stress distributions within specimens, the stress distribution of composite test specimens having different geometries was studied through finite element analysis. The interlaminar shear stresses were also analyzed to confirm that test specimen failure occurs due to the normal stresses; shear stresses are insignificant.

The model was first analyzed in three dimensions and later as an axisymmetric problem. The interlaminar shear stresses were found to be insignificant compared to the normal stresses in the thickness direction. The normal stress in the thickness direction was found to be the most constant along the midplane for test specimens with a slightly smaller radius at the mid-plane than at its upper and lower surfaces.

A second problem was analyzed by finite element analysis to understand the dependence of stiffness measurements on test specimen size for rectangular specimens. For this analysis, the specimen was modelled in plane stress and plane strain. The results showed that finite element analysis is a good prediction of Young's modulus for both thick and thin test specimens. It also showed agreement with analytical calculations of interlaminar stiffness for thick specimens, and experimental measurements for both thin and thick specimens.

FINITE ELEMENT ANALYSIS OF INTERLAMINAR TENSILE TEST SPECIMENS

Diane Hageman

1.0 INTRODUCTION

Two interlaminar tensile specimen problems were analyzed by finite elements. The first analysis compared different geometries of carbon-carbon tensile specimens to determine which shape produces the most uniform load distribution within the test specimen. The second analysis was used to calculate the interlaminar stiffness of carbon-epoxy tensile specimens. This report is divided into two parts, one for each analysis.

2.0 PART I : INTERLAMINAR TENSILE TESTING OF 2D CARBON-CARBON

Carbon-carbon composite materials are desirable for many applications because of their ability to withstand extremely high temperatures. Yet, designing with two-dimensional carbon-carbon composites is often limited by the stiffnesses and strength in the direction perpendicular to the lamina-plane (known as the interlaminar direction) which is one order of magnitude lower than the lamina in-plane stiffnesses and strengths. Experimental measurements of the interlaminar, or "through-the-thickness", stiffnesses and strength are required to use carbon-carbon composites properly; however, these tests have shown poor repeatability. To aid in the evaluation and the optimization of experiments conducted at the Material Directorate of Wright-Patterson Laboratory, a finite element analysis of a two-dimensional fabric reinforced carbon-carbon composite loaded in interlaminar tension was performed.

2.1 MODELING OF THE EXPERIMENTAL SETUP

A current problem with experimental measurements is the stress concentration in the test specimen caused by the test procedure itself. To measure interlaminar strengths and stiffnesses, a circular disk of carbon-carbon material is glued between two aluminum or steel pull tabs which are then loaded in tension. Since the composite material itself has a much lower poisson ratio and higher in-plane Young's Modulus than the pull tab, the test specimen and the pull tab generally do not have the same in-plane deformation. This mismatch in material properties causes a stress concentration at the free edge of the interface which is much higher than the stresses in the rest

of the test specimen, and thus failure tends to occur at this location. Unfortunately, there is no way of measuring the stress concentration (only the applied load is known) and thus the measured strength is much lower than the actual strength.

To produce a more uniform load distribution within the test specimen, test specimens are often made with a smaller radius at their midplane than at their upper and lower surfaces, as shown in Figure 6. Experiments have shown that this geometry causes failure to occur at the desired location, the midplane of the specimen. Another design used in past experiments is a constant radius test specimen with pull tabs that have smaller radii than the test specimen⁴. In addition, combinations of these two geometries have also been considered in this analysis. This work uses a finite element analysis to show which geometry produces the most uniform load distribution.

Although the shear strength of carbon-carbon composites is higher than their interlaminar normal strengths, it must be shown that the shear stresses in the experimental measurements of interlaminar stiffness or strength are small compared to the normal stresses. In this way, the experimental failure of the test specimen and measurements of Young's Modulus can be attributed entirely to normal stresses. In isotropic materials, specimens loaded in uniform tension would not have shear stresses in the plane transverse to the loading direction, but in composites, there is coupling between normal and shear loads. The degree of this coupling depends on the composite properties such as ply layup and fiber orientations. For this reason, this finite element analysis also investigates the interlaminar shear stresses.

2.2 FINITE ELEMENT MODELING

Due to material orthotropic symmetry, an eighth of a test specimen (a 45 degree wedge) was first modeled in three dimensions, as shown in Figures 1 and 2. Since composite materials are orthotropic, it was initially believed that axisymmetry would not be satisfactory. However, as shown by Figures 3-4, stresses were found to vary little in the circumferential direction. The specimen was then modeled in two-dimensions as axisymmetric since this simplified the computation and allowed for a geometrically more desirable mesh.

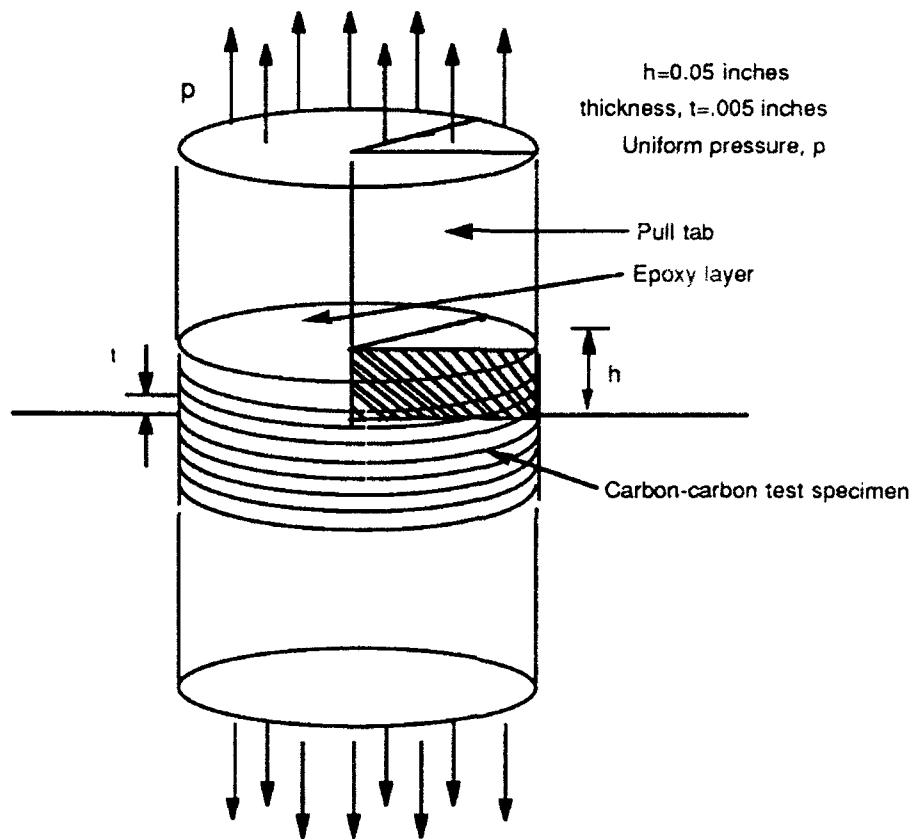
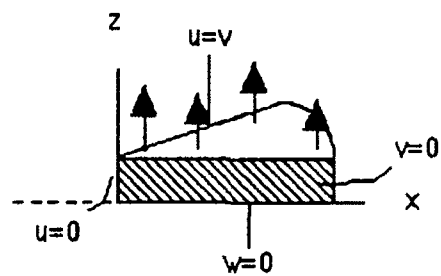


Figure 1



Boundary Conditions:

Only Radial Displacements

- At $x=0$, $u=0$
- At $z=0$, $w=0$
- At $y=0$, $v=0$
- At $x=y$, $u=v$

Figure 2

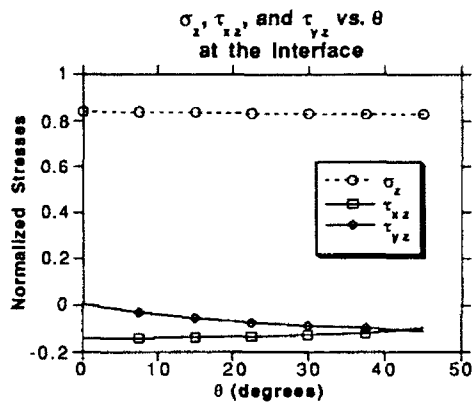


Figure 3

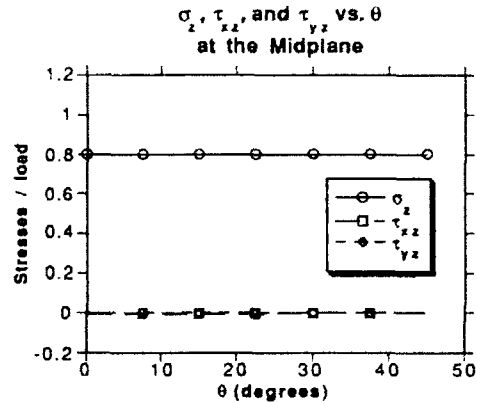


Figure 4

For modeling purposes, the bonds between the lamina were considered perfect. Also, the epoxy bond between the test specimen and the pull tab was assumed to be perfect. One case was run in which a thin isotropic bond layer between the pull tab and the specimen was included. The change in stresses was seen as insignificant; thus the bond was no longer modeled. Properties for the materials and this epoxy bond can be found in Appendix A.

The initial case studies were more concerned with the pattern of stresses than quantitative data. The convergence of mesh designs was only considered for the two-dimensional problems.

Four types of models with slightly different geometries were analyzed. The pull tab and specimen were always loaded in one of two ways. Either a uniform displacement or a uniform pressure force was applied to the upper surface of the pull tab. Figures 5-8 show boundary conditions and loading for each type of model.

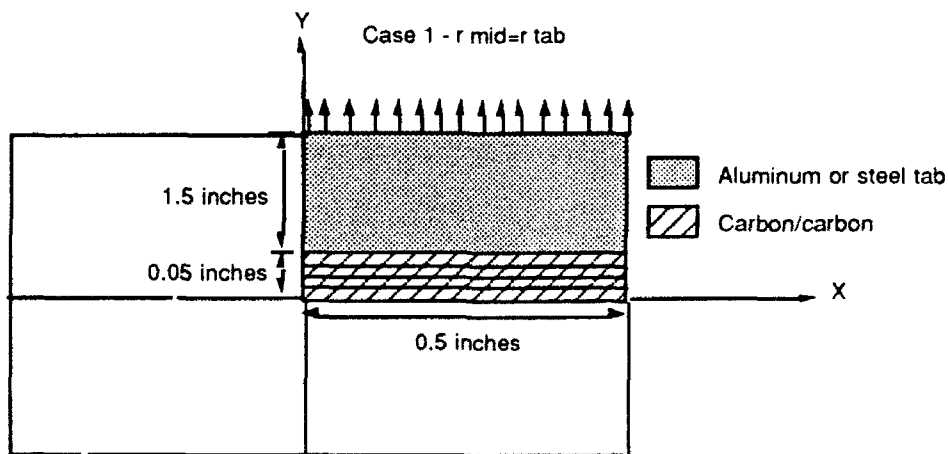


Figure 5

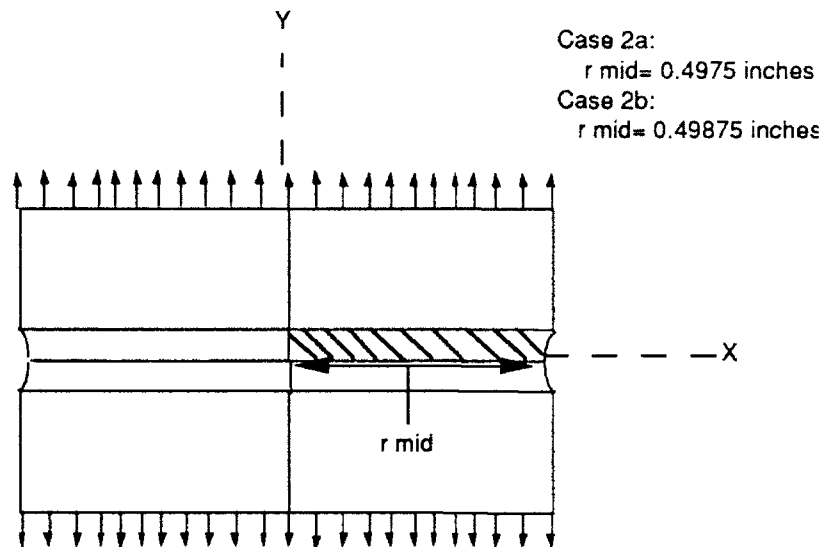


Figure 6

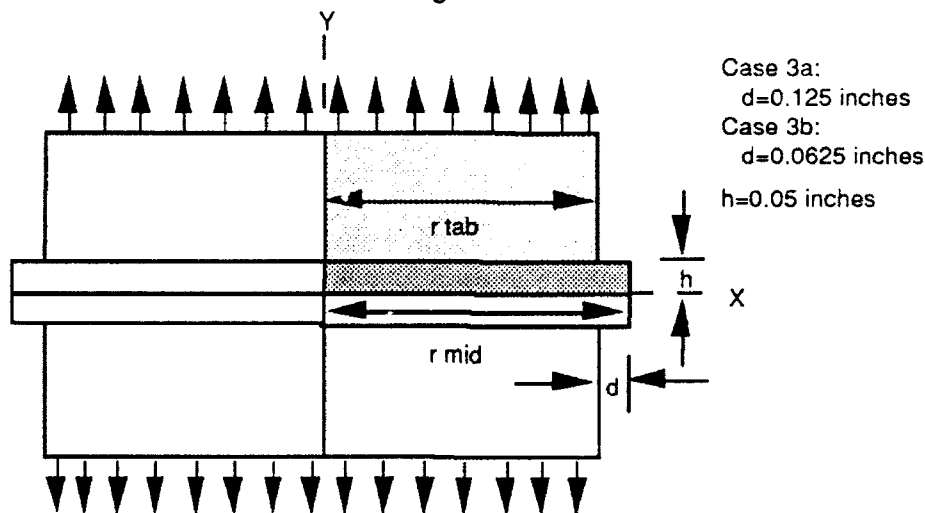


Figure 7

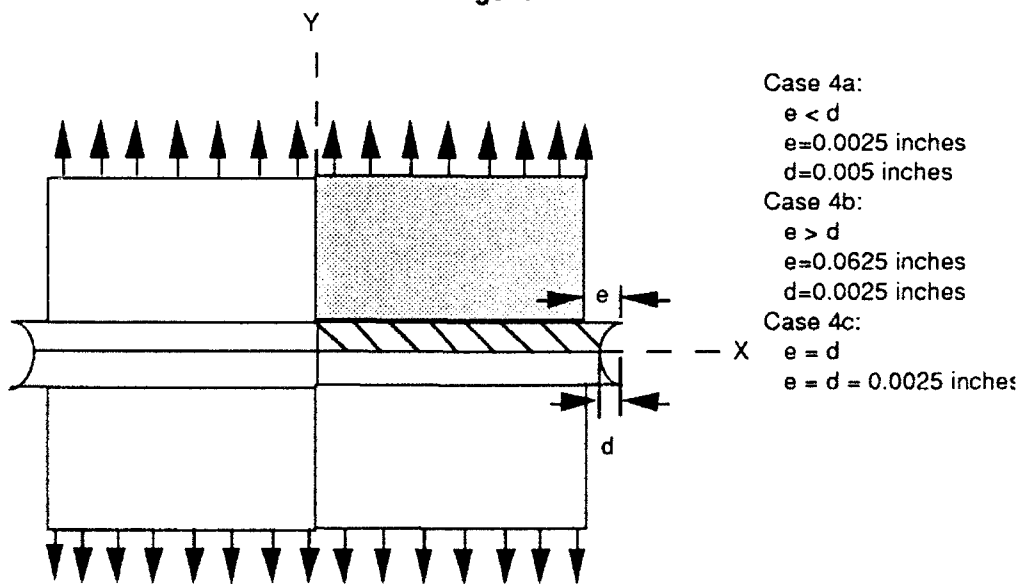
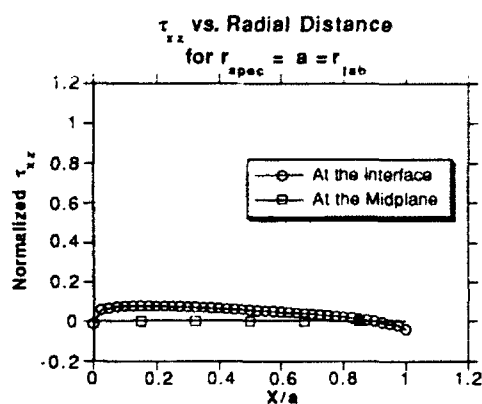
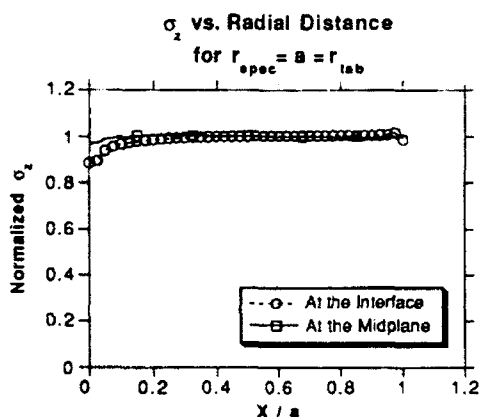


Figure 8

2.3 Results

As shown in Figures 9-17, the normal stress in the z-direction at the midplane for each of the cases is fairly constant, except at the center ($r=0$) and near the free edge ($r=0.5$ inches). For Case 1, where the test specimen has a constant radius which is the same as the radius of the pull tab, there is a stress concentration at the interface free edge ($r=0.5$ inches) which is slightly higher than the concentration at the midplane. These results are as expected because of the mismatch in the poisson ratios of the composite material and that of the aluminum or steel pull tab.

Case 1



The results of Case 2, where the radius at the midplane of the specimen is smaller than at the interface, are slightly different than for Case 1. The normal stress is again relatively constant along the midplane and the interface. For this case, however, the highest stress concentration occurs at the mid-plane free edge, not the interface free edge. The maximum shear stress occurs near the center of the specimen, and is less than 7.5% of the normal stress.

Case 2

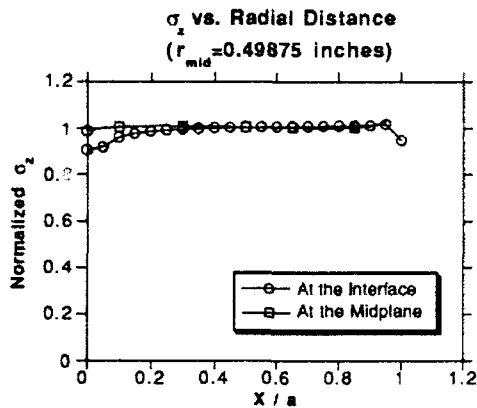


Figure 10-a

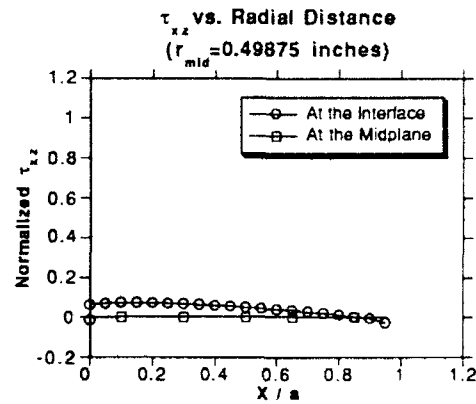


Figure 10-b

For Case 3, the specimen has a constant radius, but it is larger than the pull tab radius. This geometry causes an extremely high stress concentration along the interface at the free edge of the pull tab in both the normal and shear stresses. At this location, the shear stress becomes nearly 90% of the normal stress.

Case 3

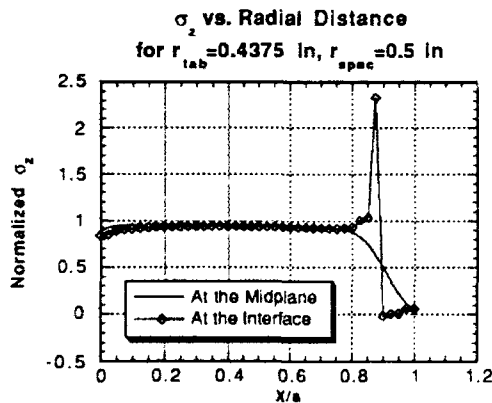


Figure 11-a

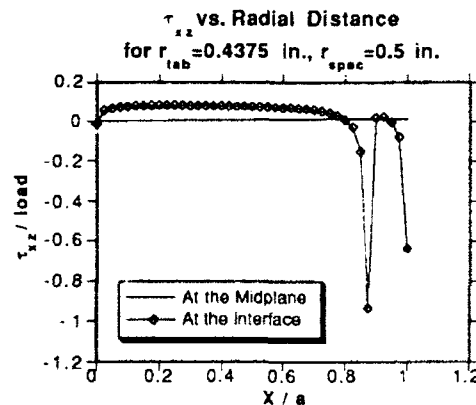


Figure 11-b

Finally, for Case 4, the radius of the test specimen varied parabolically so that its smallest radius was at the midplane. The radius of the pull tab was always less than the radius of the upper surface of the specimen. The geometry for this case was then varied in three ways: $r_{tab} > r_{mid}$, $r_{tab} < r_{mid}$, and $r_{tab} = r_{mid}$.

For Case 4a, σ_z at the midplane and at the interface remain constant until near the free edge. There is a slight increase in σ_z along the midplane, and a sharp decrease in it at the interface free

edge. Therefore, the maximum stress occurs at the mid-plane, as seen in Figure 12-a. Figure 12-b shows that the shear stress is nearly zero at the mid-plane, and less than 30% of the normal stress σ_z at the interface. These results indicate that failure should occur at the midplane, and that it is from the normal stress σ_z which is the same magnitude as the applied tensile load. These results are similar to those of Case 2.

Case 4a

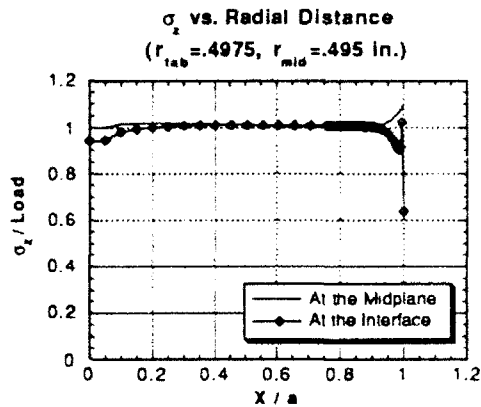


Figure 12-a

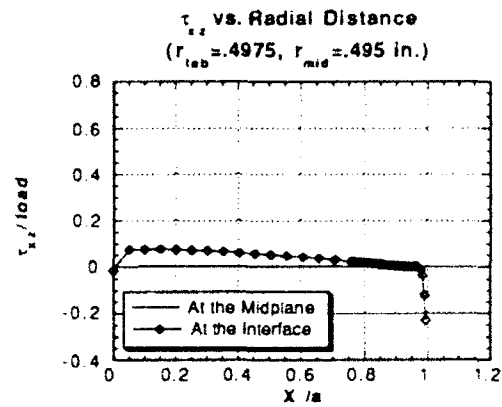


Figure 12-b

Figure 13-a, Case 4b results, shows a different pattern in stress than Case 4a. Here, σ_z decreases near the free edge at the midplane, but increases near the free edge of the interface. Failure would most likely occur at the interface free edge. The shear stress at the interface also increases sharply near the free edge to a value that is about 40% of the normal stress. This value is significant, and thus the shear stress may contribute to failure. These results are similar in nature to Case 1.

Case 4b

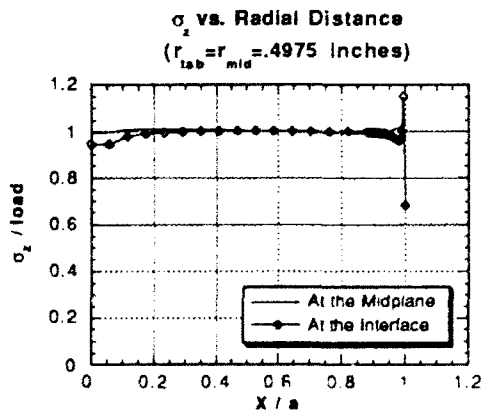


Figure 13-a

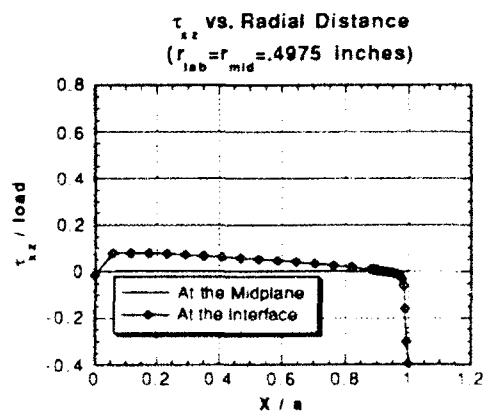


Figure 13-b

The results of Case 4c, where $r_{mid} > r_{tab}$, are more extreme than Case 4b. The interface normal stress increases sharply near the free edge. Its magnitude is much greater than the values at the midplane, as shown in Figure 14-a. As shown in Figure 14-b, the magnitude of the shear stress at the interface free edge also increases sharply. The highest shear stress is about 90% of the value of the normal stress which is significant. These results are comparable to Case 3.

Case 4c

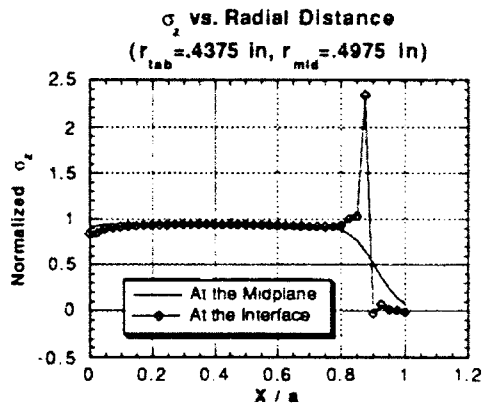


Figure 14-a

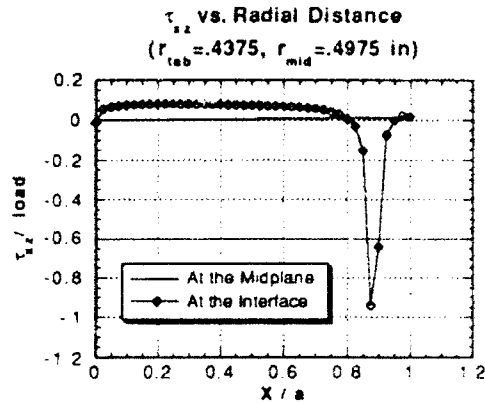


Figure 14-b

2.4 CONCLUSIONS

For measurements of interlaminar strengths and stiffnesses, the best tensile specimen geometry is that of Case 2 where the midplane radius is slightly smaller than its upper and lower surfaces, creating a dog-bone effect. The pulltabs for this geometry have the same radii as the outer surfaces of the specimen. This geometry creates a relatively constant normal stress over the cross section of the test specimen. More importantly, the highest normal stress, and thus failure, occurs at the midplane, not at the interface. The shear stresses are insignificant compared to the normal stresses for this model.

It should be noted, however, that the stress concentrations for the constant radius geometry were found to be insignificant. Yet, for experimental measurements, the dog-bone shape is better because it provides control over where the failure occurs.

Also, the worst geometry for interlaminar strength and stiffness measurements was found to be a constant radius specimen having pulltabs with smaller radii. This geometry caused extreme

stress concentrations near the free edge of the interface. The shear stress is significant in this case, and thus may contribute to failure.

3.0 PART II:

INTERLAMINAR STIFFNESS CALCULATION FOR A $[0/90]_{2s}$ AND A $[0/90/45/-45]_s$ LAMINATE

3.1 INTRODUCTION

In the past, interlaminar stiffnesses of laminates were usually assumed to be equal to the transverse stiffness. While this is true for unidirectional composites, analytical techniques (Roy & Tsai: 1991¹, Sun and Li: 1988², and Pagano: 1974³) have shown that the interlaminar modulus is higher than the transverse moduli for nonunidirectional composites. Experimental measurements of the modulus in the lamina thickness direction are higher than the transverse modulus, but are somewhat lower than the analytical solutions for thin test specimens. This dependence on the test specimen aspect ratio is attributed to the "edge effects" of the composite material which result from property mismatches between the lamina and which were not included in the analytical solutions. The purpose of this *finite element analysis* is to include the composite edge effects in determining the interlaminar stiffness.

3.2 FINITE ELEMENT MODELING

Symmetry of this problem allowed for a quarter of the specimen cross-section to be modeled. The whole test specimen, and the portion modeled are shown in Figure 15. Figure 16 shows the boundary conditions of the finite element model.

Both plane stress and plane strain were analyzed since these conditions represent two extreme boundary conditions for comparison with experimental results: a thin plate (very small thickness) and an infinite slab of material (an infinite thickness). Experimental measurements are expected to fall somewhere in between the plane stress and plane strain values because test specimens have a finite thickness. A three-dimensional analysis was not performed initially since it involves more computation and was not deemed necessary.

Each lamina was created with orthotropic properties in the x-z coordinate system so that the interlaminar stresses, and ultimately, the edge effects, would be modeled. In other words,

properties were not smeared across the laminate. Two different ply layups were analyzed:

$[0/90]_{2s}$ and $[0/90/45/-45]_{2s}$.

To calculate Young's modulus, Hooke's law was used. The average stress along the upper surface of the specimen was divided by the average strain value.

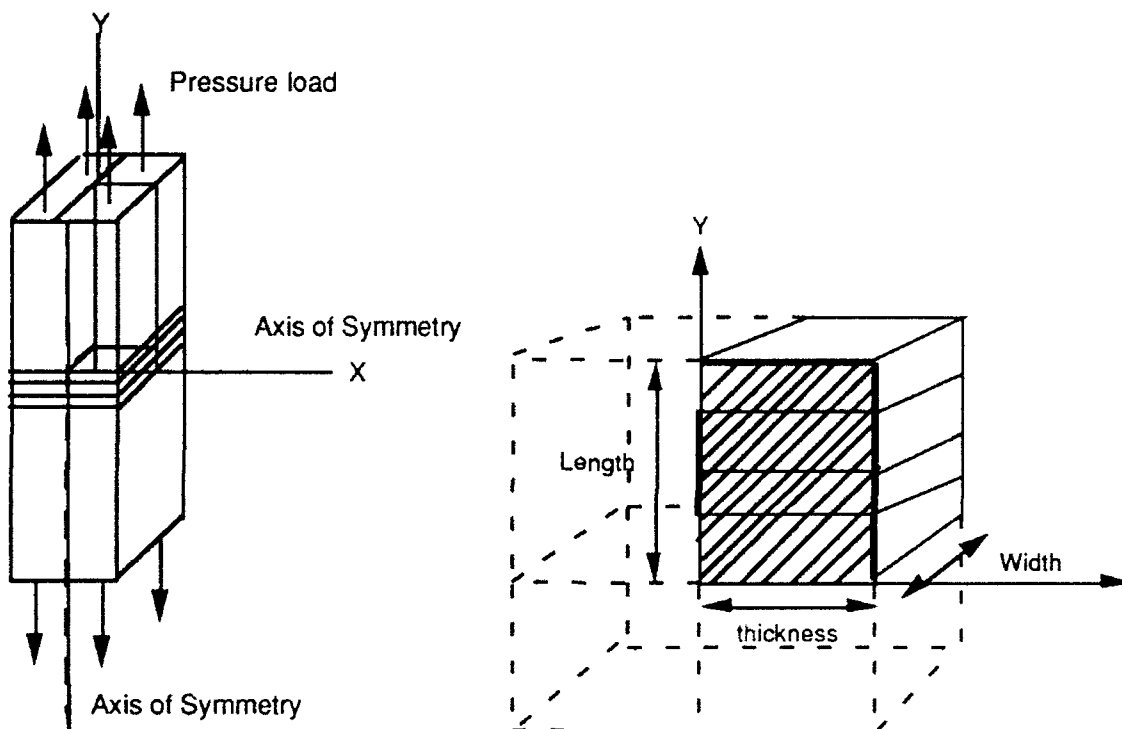


Figure 15

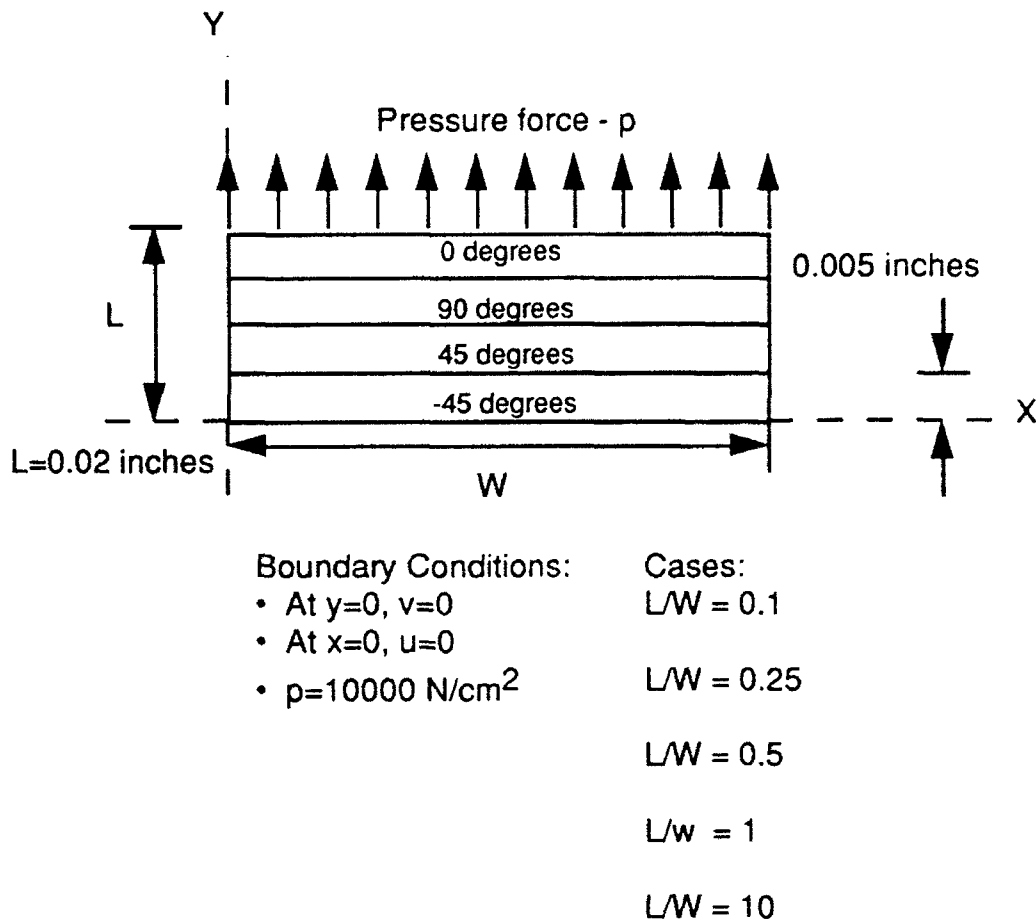


Figure 16

3.3 Results:

These results show consistency with experiments. Thin specimens had smaller interlaminar stiffnesses than thick specimens. As the thickness of the specimen increased, the interlaminar stiffness converged asymptotically. Since test specimens are neither thin plates nor infinite slabs, their stiffness measurements were expected to be between the plane stress and plane strain values. As shown in Figures 17 and 18, with one exception, the experimental values do fall within this range. The plane stress cases were the lower bounds and the plane strain cases were the upper bounds.

The results also show some consistency with analytical results. While the finite element stiffness values vary with test specimen thickness, they are close to the asymptotic value for a

thickness of 2 mm. Once this convergence is achieved, the analytical results of Roy and Tsai (1991) are approximately halfway between the plane stress and plane strain values. The results of Sun & Li (1988) were closer to the plane strain results, but still fell between the plane stress and plane strain values.

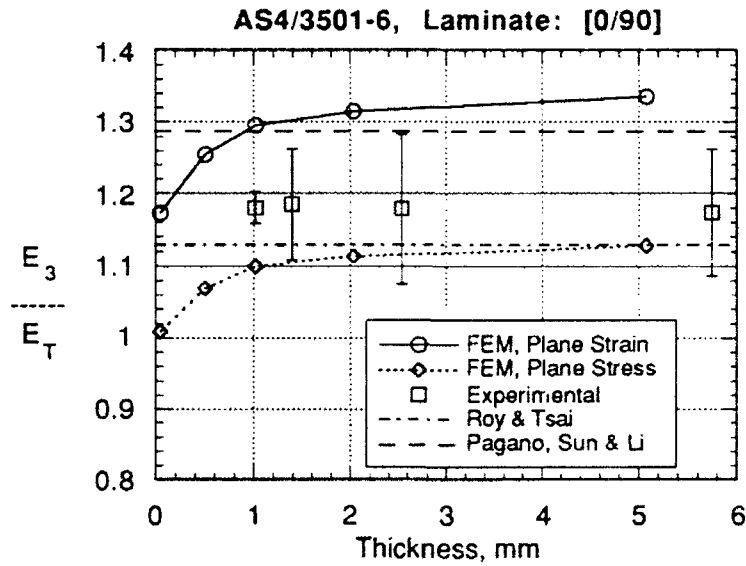


Figure 17

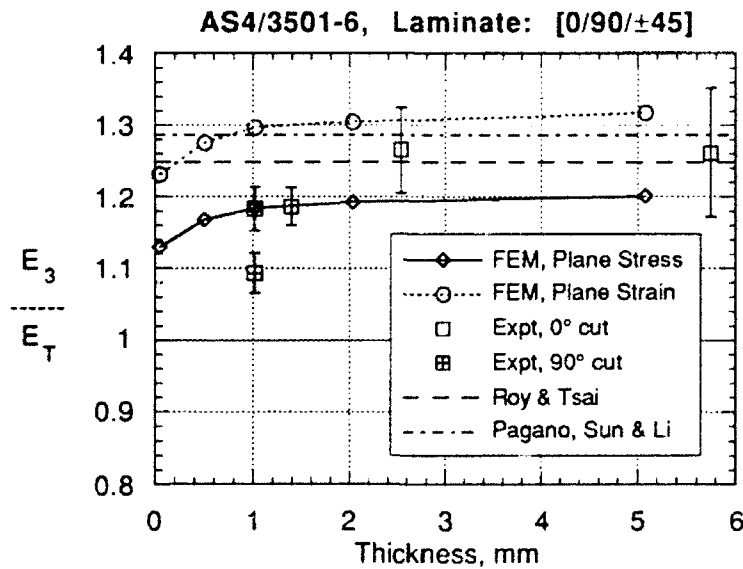


Figure 18

3.4 CONCLUSION

These results show that finite element analysis can predict the interlaminar stiffness for both thin and thick specimens. Also, FEM suggests that the analytical calculations are good predictions of the interlaminar stiffness for thick specimens in which the edge effects are insignificant.

ACKNOWLEDGEMENTS

The author would like to thank Dr. Ajit Roy for his suggestion of the problems and his guidance.

The author also is indebted to AFOSR for its sponsorship.

Appendix A

MATERIAL PROPERTIES

Carbon-carbon reinforced weave (orthotropic):

$E_1 = 17.5e6 \text{ psi}$ $E_2 = 16.5e6 \text{ psi}$ $E_3 = .737e6 \text{ psi}$
 $\nu_{12} = 0.1$ $\nu_{23}=\nu_{13}=0.1$
 $G_{12}=.856e6 \text{ psi}$ $G_{23}=G_{13}=.75e6 \text{ psi}$

Epoxy (isotropic):

$E_1 = 0.640e6 \text{ psi}$
 $\nu_{12} = 0.34$

Aluminum (isotropic):

$E_1 = 10.667e6 \text{ psi}$
 $\nu_{12} = 0.3$

Carbon-epoxy Lamina (orthotropic):

0 degrees

$E_1=138e9 \text{ Pa}$ $E_2 = 10.6e9 \text{ Pa} = E_3$
 $\nu_{12} = 0.3$ $\nu_{23}=0.52$ $\nu_{31}=0.023$
 $G_{12}=5.17e9 \text{ Pa}$ $G_{23}=3.48e9 \text{ Pa}$ $G_{13}=5.17e9 \text{ Pa}$

90 degrees

$E_1=10.6e9 \text{ Pa} = E_2$ $E_3=138e9 \text{ Pa}$
 $\nu_{12}=0.52$, $\nu_{23}=0.023$ $\nu_{31}=0.3$
 $G_{12}=3.48e9 \text{ Pa}$ $G_{23}=G_{13}=5.17e9 \text{ Pa}$

+/- 45 degrees

$E_1=18.24e9 \text{ Pa}$ $E_2=13.588e9 \text{ Pa}$ $E_3 = 18.24e9 \text{ Pa}$
 $\nu_{12}=0.1162$ $\nu_{23}=0.08656$ $\nu_{31}=0.3$

REFERENCES

1. Roy, A. K. & Tsai, S. W., Three-Dimensional Effective Moduli of Orthotropic and Symmetric Laminates. *Journal of Applied Mechanics* (ASME), **59** (1) (1992) 39-47.
2. Sun, C. T., and Li, S., 1988, "Three-Dimensional Effective Elastic Constants for Thick Laminates," *Journal of Composite Materials*, Vol. 22, No. 7, pp.629-639.
3. Pagano, N. J., 1974, *Exact Moduli of Anisotropic Laminates: Mechanics of Composite Materials*, Vol. 2, G. P. Sendeckyj, ed , Academic Press.
4. Pereira, J. M., "Flat Tension Strength of Damaged Composite Laminates", US-Japan Composite Conference, Orlando, Fl., June 1992

**THE DESIGN OF A NO₂ CHEMILUMINESCENCE
TEST CHAMBER**

**Andrew P. Johnston
Graduate Student
Department of Mechanical and Aerospace
Engineering and Engineering Mechanics**

**University of Missouri-Rolla
Rolla, Missouri 65401**

**Final Report for:
AFOSR Summer Research Program
Wright Laboratory
Flight Dynamics Directorate
Aeromechanics Division
Experimental Engineering Branch
Aero Optics Instrumentation Group**

**Sponsored by:
Air Force Office of Scientific Research
Bolling Air Force Base, Washington, D.C.**

August 1992

THE DESIGN OF A NO₂ CHEMILUMINESCENCE TEST CHAMBER

Andrew P. Johnston
Graduate Student
Department of Mechanical and Aerospace
Engineering and Engineering Mechanics
University of Missouri-Rolla

Abstract

The design of a test chamber to aid in the investigation of the use of NO₂ chemiluminescent radiation as a new diagnostic tool for hypersonic flows is presented. The reaction of interest is nitric oxide reacting with atomic oxygen to yield nitrogen dioxide plus a photon. The initial theory and design work has been completed earlier, and a proof of principle experiment is to be performed to answer remaining questions. The conditions within the continuous flowing cryogenic vacuum chamber will approximate typical conditions found in the free stream of a Mach 12 to Mach 14 hypersonic wind tunnel. The flexible design will allow for studies of the chemiluminescent reaction rates, nitric oxide depletion, ultraviolet laser beam penetration, and required nitric oxide concentration. The design includes the necessary equipment needed for safe delivery of the nitric oxide, dry air, and nitrogen to the test cell.

THE DESIGN OF A NO₂ CHEMILUMINESCENCE TEST CHAMBER

Andrew P. Johnston

INTRODUCTION

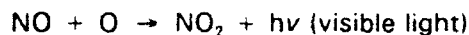
The mission of the Aero Optics Instrumentation Group is to support the Flight Dynamics Directorate through the implementation of flow diagnostics in the wind tunnel facilities as well as through the investigation and development of new flow diagnostic techniques. Throughout the summer, my activities involved working within both of these mission areas. A significant portion of time was designated to assisting in the set-up and operation of a 2-dimensional Laser Doppler Velocimeter in the 20 inch hypersonic wind tunnel (HWT), which operates at Mach numbers between 12 and 14. Assistance was also provided in the Subsonic Aerodynamic Research Laboratory (SARL) during laser light sheet operations for flow visualization purposes. The group is organizing a new laser laboratory, for which I participated in the installation of an argon-fluoride excimer laser and in the construction of a 15 foot long ventilation system. A seminar providing an overview of Particle Image Velocimetry, a relatively new experimental technique and the topic of my thesis, was organized and presented to an audience of 20 people from Flight Dynamics. Throughout the summer, a design for a chemiluminescence test cell to aid in the investigation of using NO₂ chemiluminescent radiation as a new diagnostic tool for hypersonic flows was developed.

The purpose of this report is to present a preliminary design of a cryogenic vacuum chamber to be used in the investigation of NO₂ chemiluminescence reactions under conditions found in the HWT. A discussion on the capabilities and cost is also included.

NO₂ CHEMILUMINESCENCE

The term chemiluminescence refers to the emission of light at low temperatures due to a

chemical reaction. There are a number of chemiluminescent reactions, and this proposed diagnostic tool utilizes the radiation emitted from the reaction



The application of this technique in the wind tunnel is summarized as follows. Nitric oxide (1% of the tunnel mass flow) is injected upstream of the test section and mixes with the tunnel air as it convects downstream. Atomic oxygen is formed in the test section through the dissociation of molecular oxygen using ultraviolet light. An argon-fluoride excimer laser used in conjunction with a Raman cell will deliver a 1mm to 3mm diameter beam of ultraviolet light at 166.5 nm. The chemiluminescent reaction and the resulting emission of light takes place along the "beam" of atomic oxygen as it convects downstream through the test section.

According to Aerodyne Products Corporation, which performed the initial theoretical and design work under contract with the United States Air Force, the possible measurements are velocity, temperature, and pressure as well as flow visualization of flow structures such as shocks and boundary layers.[1,2] These values will be obtained through analysis of radiation data taken with a CCD camera and a photomultiplier tube. This technique is suited for hypersonic conditions because of its applicability to low density - low pressure flows, as opposed to other techniques such as Schlieren photography. However before NO₂ chemiluminescence can be applied to the HWT, several questions concerning the reaction at hypersonic conditions should be answered.

The continuous flowing cryogenic vacuum test cell will be utilized at minimum as a proof of principle and answer the question of whether or not the reaction takes place with a high enough rate. Further experimental investigations concerning the depletion of NO during mixing and the absorption (penetration depth) of the 166.5 nm beam can also be performed. Initial calibration of radiation intensity versus temperature and pressure is possible along with the verification of the 1% NO concentration. These investigations are necessary, either in the test cell or in the HWT, because little

information is available on the reaction rates at cryogenic conditions.

The typical conditions utilized during the design of the test cell are listed in Table 1. The actual free stream conditions in the HWT can not be duplicated in a static or low flow test chamber due to the condensation of air. Instead, a near free stream condition was chosen by Mr. John Schmisser, Project Engineer, which matched the free stream density at a temperature slightly above the 86°R condensation temperature. Another typical condition that would be encountered in the HWT is the flow behind a 10 degree oblique shock. A relatively large range of temperatures and pressures will be obtainable as opposed to a single condition. The system is considered to be in a medium vacuum, as opposed to a high vacuum. Calculations of the molecular mean free path and collision frequency show that a continuum exists. It is also noted that a test cell can not reproduce the rapid expansion that occurs in the HWT, but instead only reproduce the static conditions and reaction times (not rates) during mixing of the NO and air.

Table 1. Design conditions to approximate hypersonic flow.

CONDITION	PRESSURE (psia\mm Hg)	TEMPERATURE (R\ K)	DENSITY (lbm/ft ³)
FREE STREAM (M = 12.5)	0.0056 \ 0.29	56 \ 31	2.7×10^{-4}
APPROXIMATE FREE STREAM	0.011 \ 0.56	109 \ 60	2.7×10^{-4}
BEHIND 10° OBLIQUE SHOCK	0.059 \ 3.0	151 \ 84	1.0×10^{-3}

TEST CHAMBER

The design of the test cell evolved over the 12 week period into a final design that meets the design requirements along with being adaptable and safe. The major requirements of the test cell are

to allow for continuous flow, to obtain the test conditions, allow for laser beam and optical access, and to obtain pressure and temperature measurement. The vacuum chamber must be capable of handling toxic gases at cryogenic conditions. Figures 1 - 3 show the final design of the chamber. Several adaptive features are incorporated into the design which enable further studies beyond just a proof-of-concept experiment. The basic use of the test cell will be as follows. Two gases, dry air and NO will enter the chamber through two supply lines and flow through the center of a cryogenic cell to be cooled during mixing. After exiting the cryogenic vessel the flow passes through a reaction zone through which the ultraviolet laser beam passes. The chemiluminescent reaction takes place, and radiation measurements are taken through an optical access. The reacted gases flow out of the test cell, through a filter, through the vacuum pump and into the ventilation system. Note that the entire test chamber is not kept at cryogenic temperatures but instead uses local cooling with the surrounding vacuum acting as an insulator. The major components of the test cell are the main chamber, the liquid nitrogen (LN_2) system, gas feed-throughs, laser beam entrance and exit ports, the optical access, instrumentation feed-throughs, and the outlet to the vacuum pump.

The main chamber is constructed of a 12 inch length of 8 inch diameter stainless steel tube. Due to the corrosive nature of NO and its byproducts, all components will be either 304 or 316 stainless steel. Removable 10 inch diameter flanges will seal the top and bottom of the main chamber and will be machined for gas and temperature measurement feed-throughs. The cell will utilize either "Conflat" or "Del-Seal" type flanges throughout. ASA type flanges are not recommended.

The design allows for obtaining the low temperatures by means of local cooling of the air and NO mixture by convection as it flows through a one inch diameter tube in the center of a LN_2 vessel. The 4" x 7" thin walled cylindrical vessel will hold approximately a quart of LN_2 and will adequately cool the flow through the center to the temperature of the LN_2 . If a slightly higher temperature is desired the air supply tube can be inserted further into the test cell to allow less flow time along the

LN₂ tube. The temperature of the LN₂ can also be reduced approximately 12 K (down to 63 K) by pulling a vacuum in the vessel. The LN₂ system, shown in Figure 4, consists of a slightly pressurized Dewar flask, supply line and valve, a thin walled stainless steel vessel, outlet line to a relief valve, and a 2-way valve leading to a vacuum pump or to an overflow tube. A vacuum gauge may be added to the system for pressure indication.

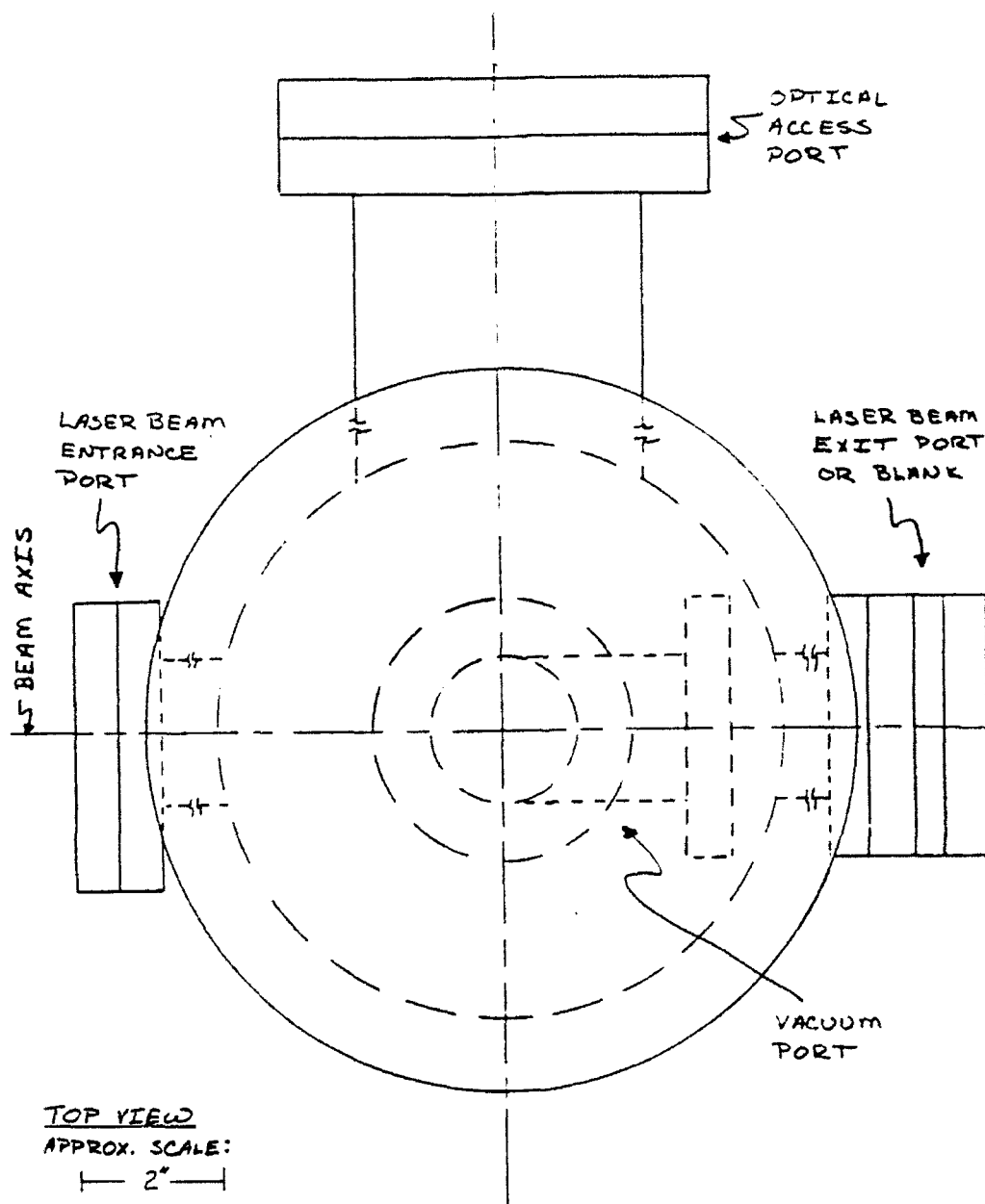


Figure 1. Top view of stainless steel test chamber.

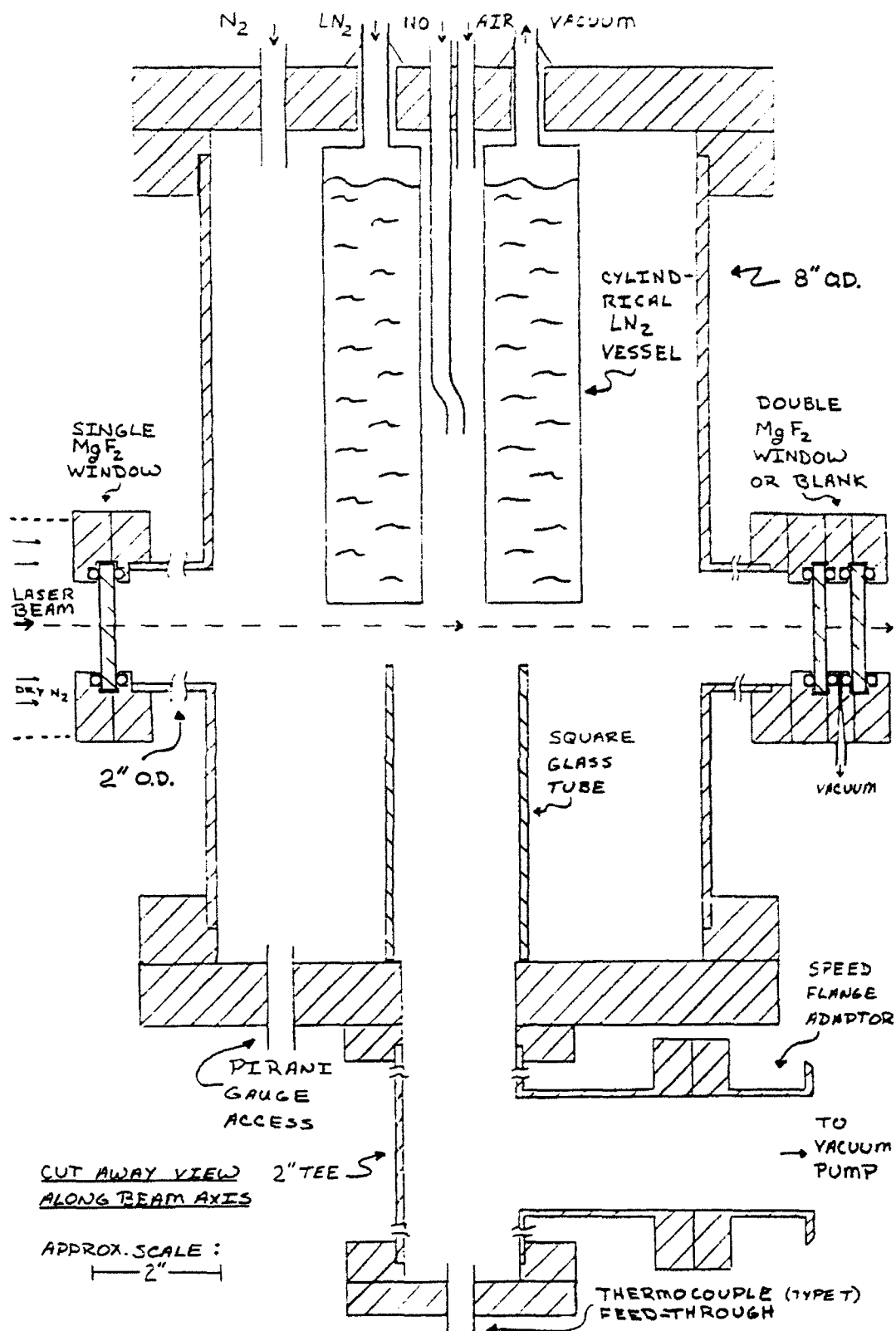


Figure 2. Cut away view of the test chamber along the laser beam axis.

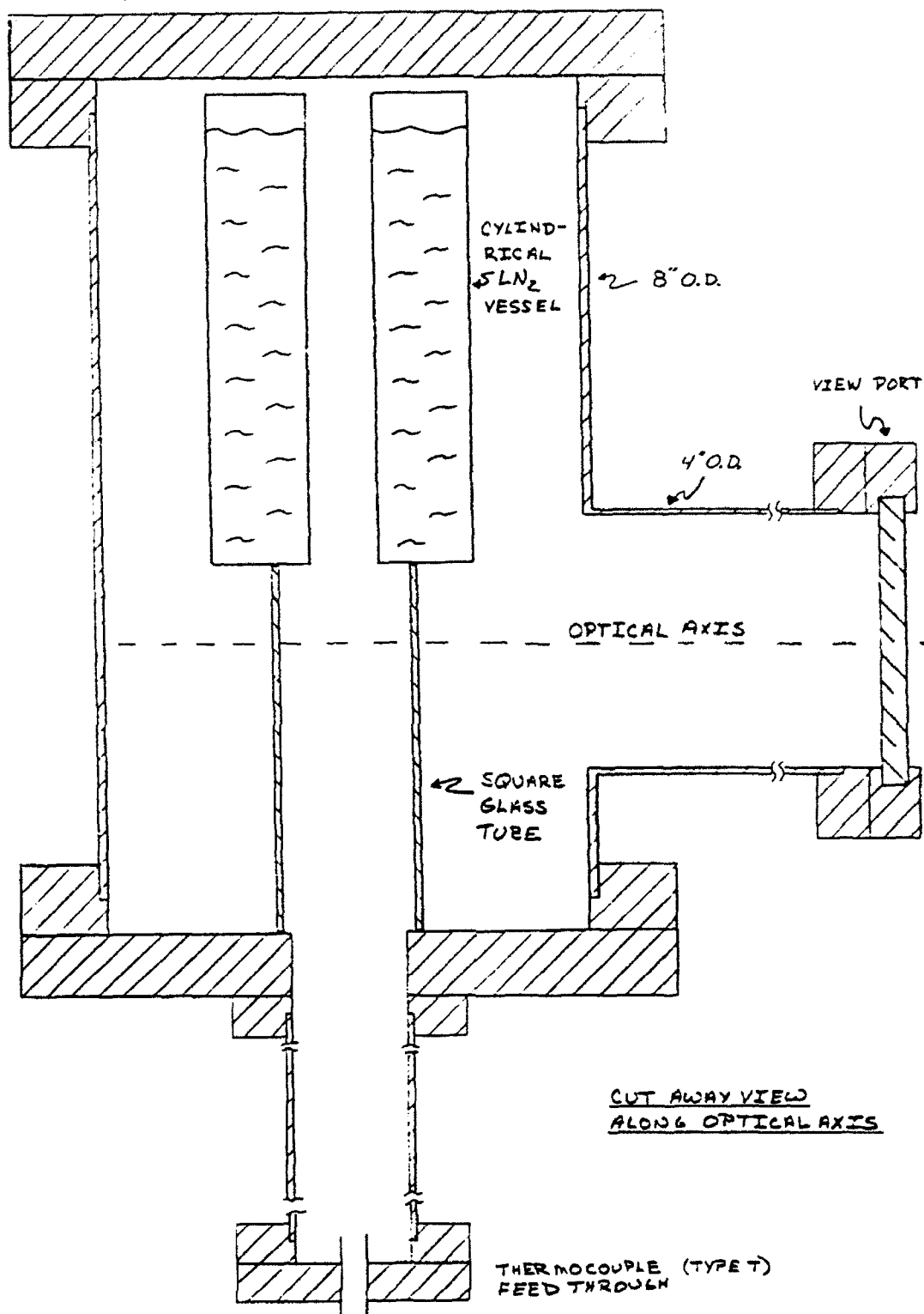


Figure 3. Cut away view of the test chamber along the optical axis.

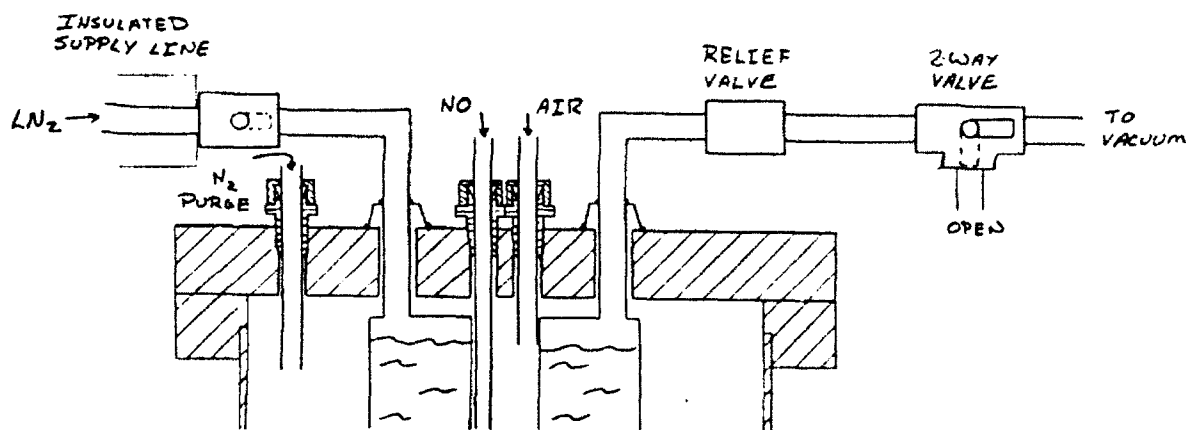


Figure 4. Top flange of the test cell showing gas feed-throughs and the LN₂ system.

The gases utilized in the test cell are nitrogen, dry air, and nitric oxide. The NO and air enter the test cell through two long supply tubes that extend into the test cell. The mixing length and the resulting reaction time between the air and NO is adjustable by sliding the NO tube up or down relative to the air supply tube. "Swage-Loc" or "Crouse-Heinz" type fittings with Teflon ferrules are used to allow for repositioning and resealing of the tubes. Condensation and/or sublimation of the NO on the sides of the supply tube and the LN₂ vessel are of concern. The tip of the NO tube should be centered in the one inch diameter tube in order to eliminate sublimation of NO on the sides of the LN₂ vessel. If condensation of NO occurs inside the NO supply tube, a double walled, vacuum insulated, supply tube for the NO can be designed and easily installed. It is noted that experimentation with krypton gas, which has similar critical and triple points as NO and is not corrosive or toxic, could be used to investigate problems with condensation and sublimation. The nitrogen will be used as a purge gas and will enter the cell through the N₂ port and the NO line during a purging procedure.

Two external ports are utilized for laser beam access to the reaction zone. The laser beam will enter the test chamber through a MgF₂ window after exiting a N₂ purged beam delivery tube. A 2 inch diameter port is arc welded to the side of the main chamber to mount the window. The ultraviolet

beam is either absorbed at the opposite side of the test cell or passes through a set of double MgF_2 windows - allowing for power measurements of the exiting beam for beam absorption studies. The vacuum insulated pair of windows prevents any condensation from forming on the windows. A double window is not necessary for the beam entrance due to the dry N_2 purge in the beam delivery tube. The MgF_2 windows will be custom made using a basic O-ring design for sealing. Extension tubes of various lengths may be added to the entrance port for penetration depth studies.

A 4 inch diameter optical viewport is located perpendicular to the beam ports. The optical access will be used for taking measurements using a CCD camera or a photomultiplier tube. The port should be about 6 inches in length to assure that it is near room temperature so no condensation forms. If condensation should form, an extra view port can be added and a vacuum pulled between the two. The port is positioned along the main chamber such that a majority of the reaction area will be in the center of the window.

The desirable location of the reaction is along the long axis of the chamber, flowing from top to bottom. Because of the local cooling, the large temperature gradients in the cell will give rise to substantial natural convection currents. These currents in turn disperse the NO and air mixture throughout the entire test cell. In order to help channel the flow from top to bottom along the center of the cell, a rectangular glass tube is added. Small openings on the top allow the laser beam to pass through. The tube will not isolate the NO and air mixture in the center of the cell, but it will help channel the flow through the desired reaction zone.

The measurement of temperature and pressure in the reaction area is provided for. A four channel Pirani gauge will be utilized for pressure measurements. The placement of the gauge out of the main flow region is due to possible corrosion of the tungsten filament. Three other gauges can be placed throughout the system if desired. A typical compression fitting like those used for the gas feed-

throughs are used with both the Pirani gauge and the thermocouple. A shielded and grounded type T thermocouple is suggested for the temperature measurement. A type K may work, but the calibration may shift when used in a vacuum. The reference junction is at a convenient 0°C. These instruments are in-hand.

The outlet line from the chamber to the vacuum pump is shown in Figure 5. The required speed of the vacuum pump is 40 cfm. A 75 cfm vacuum pump is in-hand. A 2 inch line should provide enough conductance if the distance is kept relatively short. A throttle valve is utilized to manually control the conductance of the tube and thus the pressure in the test chamber. A multi-turn popit valve is recommended. Speed flanges are added for convenience. A filter must be placed in-line before the vacuum pump in order to protect it from the corrosive gases and acids. This will also help cut down on the contamination of the test cell with hydrocarbons from outgassing from the pump. The actual filter material still needs to be selected. Activated alumina may be used to trap the acids.

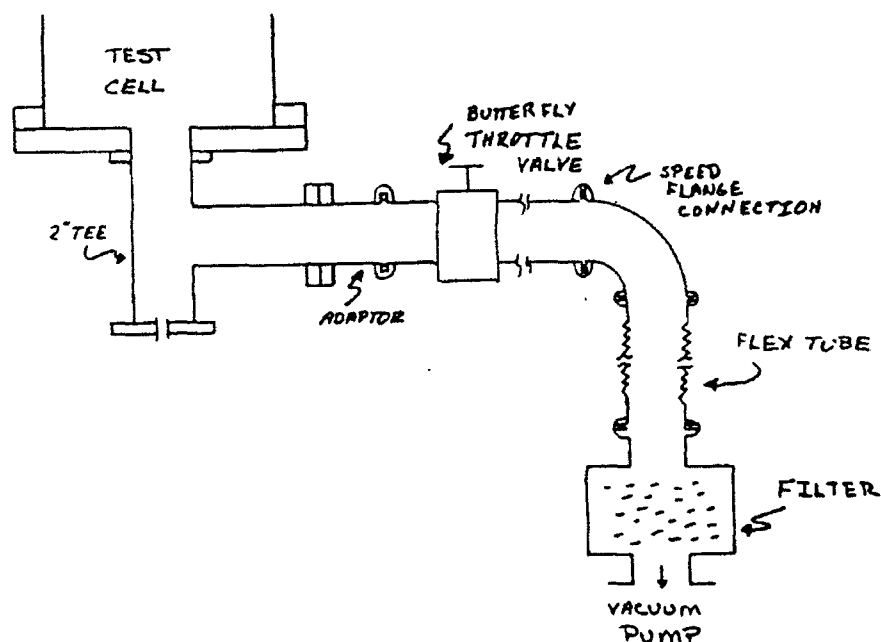


Figure 5. Sketch of vacuum pump side of the test chamber.

GAS DELIVERY SYSTEM

The gas delivery system allows for the supply of NO, dry air, and N₂ into the test chamber. A schematic of the proposed configuration is shown in figure 6. Special considerations include the safe handling of the NO, which poses a severe health hazard and a possible explosion hazard, and the control of the extremely low flow rates into the test chamber. The maximum required delivery pressure should be less than 30 psig. The major components are the NO supply panel and gas cabinet, air and N₂ pressure regulators, and low flow control valves for all three gases.

The basis for the design of the NO delivery panel is a level IV toxic gas panel. Several features are incorporated to prevent leakage, contamination, and corrosion - particularly during a purge procedure. Corrosive acids are formed when H₂O is introduced into the system. All lines are stainless steel with VCR and/or butt weld fittings. The NO panel, bottle and detector are enclosed in a vented gas cabinet. A stainless steel pressure regulator with a tied-diaphragm design is recommended. The purge system incorporates vacuum assisted vents on both the high and low pressure sides of the regulator. Quarter-turn, diaphragm type, stainless steel shut-off valves are utilized. Several check valves are added to prevent contamination by eliminating back-flow.

Added safety options for the NO panel are recommended. An emergency shut off (ESO) valve will stop the supply of NO from the bottle if a safety device is tripped. Two such safety devices are an excess flow switch and a vent switch. The excess flow switch is activated if the mass flow rate through the system exceeds a set limit, which would occur if a line broke. The vent switch will prevent the NO bottle from being opened without the ventilation system turned on and will close the bottle if flow through the vent declines. The NO detectors already in-hand which test the room air for unsafe concentrations of NO may also be tied into the ESO system.

The pressure control of the air and N₂ is less critical. A typical single stage brass regulator with

stainless steel diaphragms is recommended for the N_2 . A double stage, high purity brass regulator with stainless steel diaphragms is recommended for the dry air. This design will keep a constant delivery pressure for decreasing cylinder pressure. The grades of N_2 and dry air used should be such that the amount of H_2O is very low - less than 5 ppm. A fine control needle valve will be adequate for the supply of N_2 to the test cell.

One of the difficulties in designing a continuous flowing vacuum cell is the control of the low flow rates into the test cell. This case is particularly difficult because of the control of the mass flow of the NO, which is 1% of the air mass flow. Table 2 shows the required flow rates and estimated pumping rates for the two design cases. These flow rates will match the reaction time (not rate) of 0.03 seconds, which is the reaction time that the mixture of NO and air will have between the injection location and the test section.[1] These flow rates assume a three inch mixing length within a one inch diameter tube. The trade-off between the extremely low flow rate into the cell and the required pumping speed was optimized. If much larger flow rates through the test cell are desired, the required pump size will not be practical. If a smaller pump is to be used, the flow rate of the NO will become extremely difficult to monitor. The possibility exists that a premixed gas or mixture outside of the test cell will work if the NO depletion studies (theoretical or experimental) are carried out. This would allow for a smaller pump to be used, since only the total flow rate has to be monitored.

Table 2. Flow rates through test cell for the design conditions.

CONDITION (M = 12.5)	MASS FLOW RATE (lbm/ft ³)		VOLUME FLOW RATE (scfm)		VACUUM PUMPING RATE (cfm)
	AIR	NO	AIR	NO	
APPROXIMATE FREE STREAM	7.3×10^{-4}	7.3×10^{-6}	9.5×10^{-3}	9.5×10^{-5}	40
BEHIND 10° OBLIQUE SHOCK	2.8×10^{-3}	2.8×10^{-5}	3.7×10^{-2}	3.7×10^{-4}	30

There are a few possible ways to control the flow, such as variable leak valves, rotameters, or mass flow controllers. The variable leak valves are recommended for several reasons. The large range of flow rates from this single valve will allow for any desired flow rate, even slightly lower than the design conditions. A micrometer type turn counter indicates how far the valve is open. These valves, once calibrated, should be fairly repeatable if they are not used as shut-off valves. Lastly, the large range of flow rates allows for the same type of valve to be utilized for both the NO and air. The rotameter is also an option that may work nicely. The flow tube will provide a visual indication of the amount of flow. However, the range will be limited and perhaps a series of different size flow tubes would have to be purchased to run all of the desired cases. The NO flow rate is so low that a custom tube will have to be made to measure the 9.5×10^{-6} scfm. Yet another option is to use a 4 channel mass flow controller. Verification of accurate control at these low flow rates is needed. This instrument may be rather expensive for a proof of principle experiment.

DISCUSSION

The design presented meets all of the requirements to be used as a proof of principle test cell for the examination of NO₂ chemiluminescent radiation at hypersonic conditions. Adaptability incorporated into the chamber expands its possible use to answering questions about penetration depth, NO depletion, initial temperature and pressure calibrations, and the best concentration of NO to be used. A few questions remain and should be answered before the final design is decided upon. The foremost question is that of the true explosive nature of NO. Conflicting reports indicate that the unstable behavior results when NO is in a condensed state or in a high temperature or pressure state.[2,3] Weak portions in the design itself include precise temperature control and adequate mixing. Enhanced mixing through mechanical means may be necessary. The placement of stainless steel mesh inside the one inch flow tube is an option if condensation does not occur.

A preliminary cost estimate of the major components of the test chamber and gas delivery

system is shown in Table 3. Several companies were contacted and provided budgetary quotes. The minimum cost to put an entire test cell and gas system together is approximately \$17,000. The option to perform some of the assembly in house is included. It is noted that the equipment purchased for the supply of NO will be used at the HWT, and thus the safety options should strongly be considered.

Table 3. Approximate price list for test chamber and gas delivery system.

ITEM	PRICE	MANUFACTURER	CONSTRUCTION
TEST CHAMBER	\$ 9,000	NOR-CAL	NOR-CAL/IN HOUSE
	14,540	COOK VACUUM	COOK VACUUM
NO PANEL	3,225	SPECTRA GASES	IN HOUSE
	4,000	SPECTRA GASES	SPECTRA GASES
GAS CABINET	1,830	SPECTRA GASES	--
VARIABLE LEAK VALVES	2,200	GRANDVILLE-PHILLIPS	--
OTHER GAS AND LN ₂ COMPONENTS	1,270	ASSORTED	--
NO & AIR CYLINDERS	420	MATHESON	--
OPTIONS			
MASS FLOW CONTROLLER*	6,825	SPECTRA GASES	--
EMERGENCY NO SHUT OFF SYSTEM	5,660	SPECTRA GASES	--
ROTAMETERS*	780	SPECTRA GASES	--

* -- replaces variable leak valves

REFERENCES

- [1] Aerodyne Products Corporation, "Design Evaluation Report for NO₂ Chemiluminescent Imaging System," Contract Number F33615-88-C-3013, June, 1990.
- [2] Aerodyne Products Corporation, "Supersonic Flowfield Measurements with NO₂ Chemiluminescent Radiation," Contract Number F33615-87-C-3012, August, 1987.
- [3] "Material Safety Data Sheet - Nitric Oxide," Spectra Gases Incorporated, Irvington, NJ: July, 1992.
- [4] Barron, F., Cryogenic Systems, 2nd ed., Oxford University Press, New York: 1985.
- [5] Anderson, J.D., Hypersonic and High Temperature Gas Dynamics, McGraw-Hill Book Company, New York: 1989.
- [6] Chapman, A.J., Heat Transfer, 4th ed., Macmillan, New York: c.1984.
- [7] Sonntag, R.E. and Van Wylen, G.J., Introduction to Thermodynamics Classical and Statistical, 2nd ed., John Wiley & Sons, New York: c.1982.
- [8] British Cryogenics Council, Safety Panel, Cryogenics Safety Manual: A Guide to Good Practice, 2nd ed., Mechanical Engineering Publications, London: 1982.

JACOBIAN UPDATE STRATEGIES FOR QUADRATIC AND
NEAR-QUADRATIC CONVERGENCE OF NEWTON
AND NEWTON-LIKE IMPLICIT SCHEMES

Daniel B. Kim

Aerospace Engineering and Engineering Mechanics
Graduate student of Dr. Paul D. Orkwis
University of Cincinnati
Apt # 408 Scioto Street
Cincinnati, Ohio 45219

Final Report for:
Summer Research Program
Wright Laboratory

Sponsored by:
Air Force Office of Scientific Research
Eglin Air Force Base, FL

September 1992

JACOBIAN UPDATE STRATEGIES FOR QUADRATIC AND
NEAR-QUADRATIC CONVERGENCE OF NEWTON
AND NEWTON-LIKE IMPLICIT SCHEMES

Daniel B. Kim
Graduate Student
Department of Aerospace Engineering and Engineering Mechanics
University of Cincinnati

Abstract

Evaluation of several Jacobian matrix simplification ideas for Newton and Newton-like implicit Navier-Stokes solvers were performed. It was found that simplifications to the Jacobian matrix can result in dramatic CPU time savings. The Jacobian simplifications were accomplished by approximating the entries of the Jacobian matrix. These approximations include updating only selected parts of the matrix with the most recently computed values of the conserved variables. The updating strategy of the matrix was based on a percentage of the maximum density residual or on a specific region of the flow-field. "Global freezing" of the Jacobian matrix for a specified number of sub-iterations was also tested. It has shown that both, partial updating of the Jacobian matrix and "global freezing" of the Jacobian matrix, can give quadratic or better convergence rates. The Jacobian simplification ideas were tested for a flat plate geometry.

JACOBIAN UPDATE STRATEGIES FOR QUADRATIC AND NEAR-QUADRATIC CONVERGENCE OF NEWTON AND NEWTON-LIKE IMPLICIT SCHEMES

Daniel B. Kim
Graduate Student
Department of Aerospace Engineering and Engineering Mechanics
University of Cincinnati

INTRODUCTION

The goal of this research was to improve the Newton's method solver for the Navier-Stokes equations originally developed by Orkwis and McRae [1,2]. The desirable part of this scheme is that it exhibits a quadratic or a near-quadratic convergence rate, which can produce iteration counts that are orders of magnitude less than those of the standard implicit relaxation solvers. The speed of the Jacobian matrix formation and the inversion of the resulting system directly influences the efficiency of the Newton's method. Therefore, procedures for improving the efficiency of these processes must be devised which will maintain the quadratic convergence property.

One way of improving the performance is to approximate the entries of the Jacobian matrix. Possible approximations include updating only selected parts of the matrix using the most recently computed values of the conserved variables, mixing higher and lower order matrix entries, and "global freezing" of the Jacobian matrix. Combination of the above can also be used, ie. partial updating the matrix and "freezing" selected matrix entries.

This leads to the question of exactly how "correct" the Jacobian matrix must be in order to obtain quadratic convergence. For example, the matrix entries from the freestream region above the bow shock wave formed by a supersonic flat plate can be altered without affecting quadratic convergence. Are there other regions in which the Jacobian matrix contribution can be slightly altered without convergence rate degradation? Also, must the order of accuracy of the contributions to the Jacobian matrix exactly match the order of the corresponding right hand

side terms or can they be reduced at selected locations? In summary, can a method be found that enjoys quadratic convergence rates and is computational efficient in the Jacobian matrix formation and inversion processes?

Matrix approximations can be used to significantly enhance the competitiveness of Newton's method if the number of matrix updates is small. If this number is small, the inverse of the new matrix can be computed quickly using the Sherman-Morrison formula [9]. When the Sherman-Morrison formula is applied $O(N^2)$ operations are required to compute the inverse for each row/column update of the original matrix as compared to $O(N^3)$ operations to recompute completely the inverse. It is important to note that the Sherman-Morrison formula was not used in this effort because the first goal of the research was to develop a strategy of updating the Jacobian matrix, and if the strategy was successful the Sherman-Morrison formula would then be applied. The ideal code would then require only one full inverse for the first iteration and minimal additional changes afterward.

The following sections describe this research in detail by first discussing the governing Navier-Stokes equations and the Orkwis and McRae Newton's method procedure. A discussion of the methodology for partially updating the Jacobian matrix and "freezing" the matrix follows. Finally, the results obtained with these methods are discussed and some concluding remarks are given.

GOVERNING EQUATIONS

The Newton's method solver has been developed for the solution of the steady, two-dimensional, laminar Navier-Stokes equations, shown below:

$$\frac{\partial \bar{F}}{\partial x} + \frac{\partial \bar{G}}{\partial y} = 0$$

where

$$\bar{F} = \begin{bmatrix} \rho u \\ \rho u^2 + p - \tau_{xx} \\ \rho uv - \tau_{xy} \\ (e + p)u - b_x \end{bmatrix} \quad \bar{G} = \begin{bmatrix} \rho v \\ \rho uv - \tau_{xy} \\ \rho v^2 + p - \tau_{yy} \\ (e + p)v - b_y \end{bmatrix}$$

$$\tau_{xx} = \frac{2}{3} \frac{\mu}{Re} \left(2 \frac{\partial u}{\partial x} - \frac{\partial v}{\partial y} \right)$$

$$\tau_{xy} = \frac{\mu}{Re} \left(\frac{\partial u}{\partial y} + \frac{\partial v}{\partial x} \right)$$

$$\tau_{yy} = \frac{2}{3} \frac{\mu}{Re} \left(2 \frac{\partial v}{\partial y} - \frac{\partial u}{\partial x} \right)$$

$$b_x = \frac{\gamma \mu}{Re Pr} \frac{\partial e_i}{\partial x} + u \tau_{xx} + v \tau_{xy}$$

$$b_y = \frac{\gamma \mu}{Re Pr} \frac{\partial e_i}{\partial y} + u \tau_{xy} + v \tau_{yy}$$

$$e_i = \frac{e}{\rho} - \frac{1}{2} (u^2 + v^2)$$

$$p = (\gamma - 1) \left(e - \frac{1}{2} \rho (u^2 + v^2) \right)$$

ρ is the density, u and v are the velocity components, p is the pressure and e is the total energy.

The viscosity, μ , is determined from Sutherland's law, Re is the Reynolds number, and Pr is the Prandtl number, a constant equal to 0.72. The governing equations are transformed into generalized coordinates and discretized using finite differences. The resulting set of nonlinear equations is then solved via Newton's method for systems of equations.

NUMERICAL METHOD

Orkwis and McRae's solution procedure is based on Newton's method for systems of nonlinear equations. These systems of equations have the form

$$F(\bar{U}) = 0$$

Newton's method is then

$$\left(\frac{\partial F}{\partial \bar{U}} \right) \Delta^n \bar{U} = -F^n(\bar{U})$$

A solution is obtained by forming the right hand side of the equation and the Jacobian matrix at the known n th iterate. The increment $\Delta^n \bar{U}$ is then found. The new value of $\bar{U}$ is

$$\bar{U}^{n+1} = \bar{U}^n + \Delta^n \bar{U}$$

The Newton's method requires a "close enough" initial guess. The scheme is modified to take this into consideration and be capable of using an arbitrary initial condition as follows

$$\left[\frac{I}{\Delta t} + \frac{\partial F}{\partial \bar{U}} \right]^n \Delta^n \bar{U} = -F^n(\bar{U})$$

The pseudo-time step in the above equation is based on the value of the residual using the equation

$$\Delta t^n = \Delta t^0 \frac{\|F(\bar{U})\|^0}{\|F(\bar{U})\|^n}$$

The above scheme has been shown to exhibit quadratic convergence once Δt becomes large.

MATRIX APPROXIMATION STRATEGY

As stated earlier, the majority of the CPU time required by the Newton's method is used in the Jacobian matrix formation and inversion processes. The first method of reducing the work in the inversion process used in this case was "global freezing" the inverted matrix for k iterations and then updating the inverse matrix [10]. The scheme becomes

$$\left(\frac{\partial F}{\partial \bar{U}}\right)^n (\bar{U}^{n,i} - \bar{U}^{n,i-1}) = -F^{n,i-1}(\bar{U})$$

where $i=1,\dots,k$

$$\bar{U}^{n+1} = \bar{U}^{n,k}$$

The iteration may be considered as the composition of one Newton step with $k-1$ simplified Newton steps. This procedure represents a simple way of generating greater than quadratic convergence rates. It is important to note that the full inverse matrix has to be solved at every k th iteration until the final solution is reached. It should be noted that the convergence rate is based on iterations with fully updated Jacobian matrix, since the work involved in the simplified Newton's step is minimal as compared to the exact Newton's step.

A second matrix approximation idea is "partial freezing" of the matrix. Partial freezing can be done based on the local residual or based on the specific region of the flow. In the case of "partial freezing" based on the residual, the update criteria is based on the changes in the solution, which are directly related to the local residual. The process involves setting a tolerance based on the maximum residual. This tolerance value decides which Jacobian matrix entries are updated and which are frozen. For example, whenever the local density residual is greater than the minimum residual plus a percentage of the difference between the maximum and the minimum density residual the point in the matrix is updated, those points not satisfying this criteria are frozen.

In the case of partial freezing of the Jacobian matrix based on the specific region of the flow, only the Jacobian entries from points that lie in the boundary layer are updated. The above methods are also applied as a composite method that combines the two approximations. It is important to note that none of the Jacobian approximation ideas has any effect whatsoever on the resultant accuracy of the final solution. Changes to the Jacobian concern only the convergence properties of the scheme not the accuracy of the final result, since all satisfy $F(\bar{U}) = 0$.

MATRIX INVERSE APPROACH

Efficiency can be enhanced if the number of Jacobian matrix updates is small since the resulting changes in the inverse of the matrix can be quickly computed using the Sherman-Morrison formula [9]. This formula states that if a matrix A is changed as per

$$A \rightarrow (A + \bar{U} \otimes \bar{V})$$

then the inverse changes as per

$$A^{-1} \rightarrow \left(A^{-1} - \frac{(A^{-1} \cdot \bar{U}) \otimes (\bar{V} \cdot A^{-1})}{1 + \bar{V} \cdot A^{-1} \cdot \bar{U}} \right)$$

To change an entire row, $\bar{V}$ would be chosen as a unit vector and $\bar{U}$ the delta changes to that row. Each row update then requires $O(N^2)$ floating point operations to compute the inverse as compared to $O(N^3)$ to recompute it completely.

RESULTS

This section describes the results obtained when the Jacobian matrix approximations were implemented in the Newton's method code developed by Orkwis and McRae[1,2]. In this work, the Jacobian matrix approximation methods were tested by calculating the supersonic flat plate test case used by Orkwis and McRae[1]. This work used the 40x40 stretched viscous grid shown in Figure 1. Slug flow initial conditions were used with $M=2$ and $Re=1.65 \times 10^6$ freestream conditions. The exact Jacobian matrix code was used as a benchmark. The cases tested were as follows

1. Fully update the matrix
2. Fully update/freeze the matrix
3. Partially update (residual) the matrix
4. Partially update (based on B.L.) the matrix
5. Partially update(B.L.)/freeze the matrix

Table 1 shows that the exact Jacobian matrix solution took 16 iterations and 390.2 CPU seconds to converge. Figure 2 shows the corresponding convergence rate graph. It is apparent that the convergence rate is quadratic.

For the second case, a combination of exact Jacobian and global freezing was tested. Figure 2 shows that for 10 iterations the convergence rate was linear and after 10 iterations it became quadratic. Based on the results shown in Figure 2, the Jacobian matrix for the first 10 iterations were fully updated, thereafter the Jacobian matrix was fully updated every 8 sub-iterations. Figure 3 shows the density residual after each full matrix update. The convergence rate was linear during the early iterations but afterward the convergence rate became cubic. Table 1 shows that this run took 11 iterations to converge as compared to 16 for the full update case. The CPU time was 284.0 seconds, which represents a 27% reduction.

Waiting to freeze the Jacobian matrix until the solution was near the quadratic convergence rate region was required. Freezing the matrix before that point gave an increase in the number of iterations and CPU time. When the matrix was globally frozen right after the first iteration the solution diverged. It is felt that in the linear convergence region the initial guess plays a dominant role, since the initial guess must be "close-enough" to the final solution for a Newton solver.

The third Jacobian approximation tested was partial updates of the Jacobian matrix. Here, the matrix was partially updated after the first iteration based on the percentage of the maximum density residual. In this case, at locations where the density residual was greater than the minimum residual plus 99% of the difference of the maximum and the minimum residual the matrix was updated. On average 2,432 entries were updated out of a total of 6,400 entries. Figure 4 shows that the convergence rate was not the smooth quadratic convergence rate of the fully updated matrix case. It took 20 iterations (4 more than the full updated case) and 485.14 CPU seconds (24.3% more than the full updated case).

Another Jacobian approximation tested was partial updating based on a specified region of the grid, in this case updating only the boundary layer. Figure 5 shows the convergence rate for this case. It is apparent that the convergence rate is much smoother than that obtained with partial updating of the matrix based on the density residual. The more dramatic result is that for this case the number of average entries being updated in the matrix is 1,440 out of 6,400. The number of iterations was 18 and 425.06 CPU seconds were required. Compared to the partial update Jacobian approximation, it is evident that what was updated is more important than the quantity of entries being updated. In the partial update case based on the density residual, 59.2% more entries were updated as compared to the above case, and more iterations and CPU time were required. The number of entries being updated in the matrix becomes very important especially once the Sherman-Morrison formula is implemented. If less entries in the matrix are updated the approach becomes more efficient.

The last case tested was a combination of the partial updates (in the B.L.) and Jacobian matrix freezing for a certain number of sub-iterations. In the linear region, the Jacobian matrix entries were updated only in the B.L. In the quadratic region the Jacobian matrix entries were frozen for eight sub-iterations. Figure 6 shows that the convergence rate was cubic. This case converged in 11 iterations and it took 275.394 CPU seconds. This case resulted in the greatest savings as compared to the fully updated case; ie 29.4%. If the Sherman-Morrison formula were implemented for this case the saving would be considerably greater. It is also important to note that only a flat plate geometry was tested. A more complicated wedge geometry is currently being tested.

CONCLUSION

The following conclusions have resulted from testing various Jacobian matrix updating strategies on the Newton's method solver developed by Orkwis and McRae[1,2]:

- * Full updating of the Jacobian matrix resulted in a quadratic convergence rate,
- * Full updating of the Jacobian matrix with global freezing resulted in a cubic convergence rate,
- * When partially updating the matrix it is more important that the "correct" entries be updated than a large number of entries,
- * The combination of partial updating of the matrix and global freezing gave the best CPU saving time as well as a cubic convergence rate,
- * More work is needed.

JAC. UPDATE STRATEGIES	ITERATION	CPU (second)	%SAVING
Full update	16	390.21 s	-
Full update/ Freeze	11	284.00 s	27%
Partial update (based on residual)	20	485.14 s	-24.9%
Partial update (based on B.L.)	18	425.06 s	-8.9%
Partial update (B.L.)/ Freeze	11	275.39 s	29.4%

Table 1. Comparison of the Jacobian matrix update strategies for flat plate

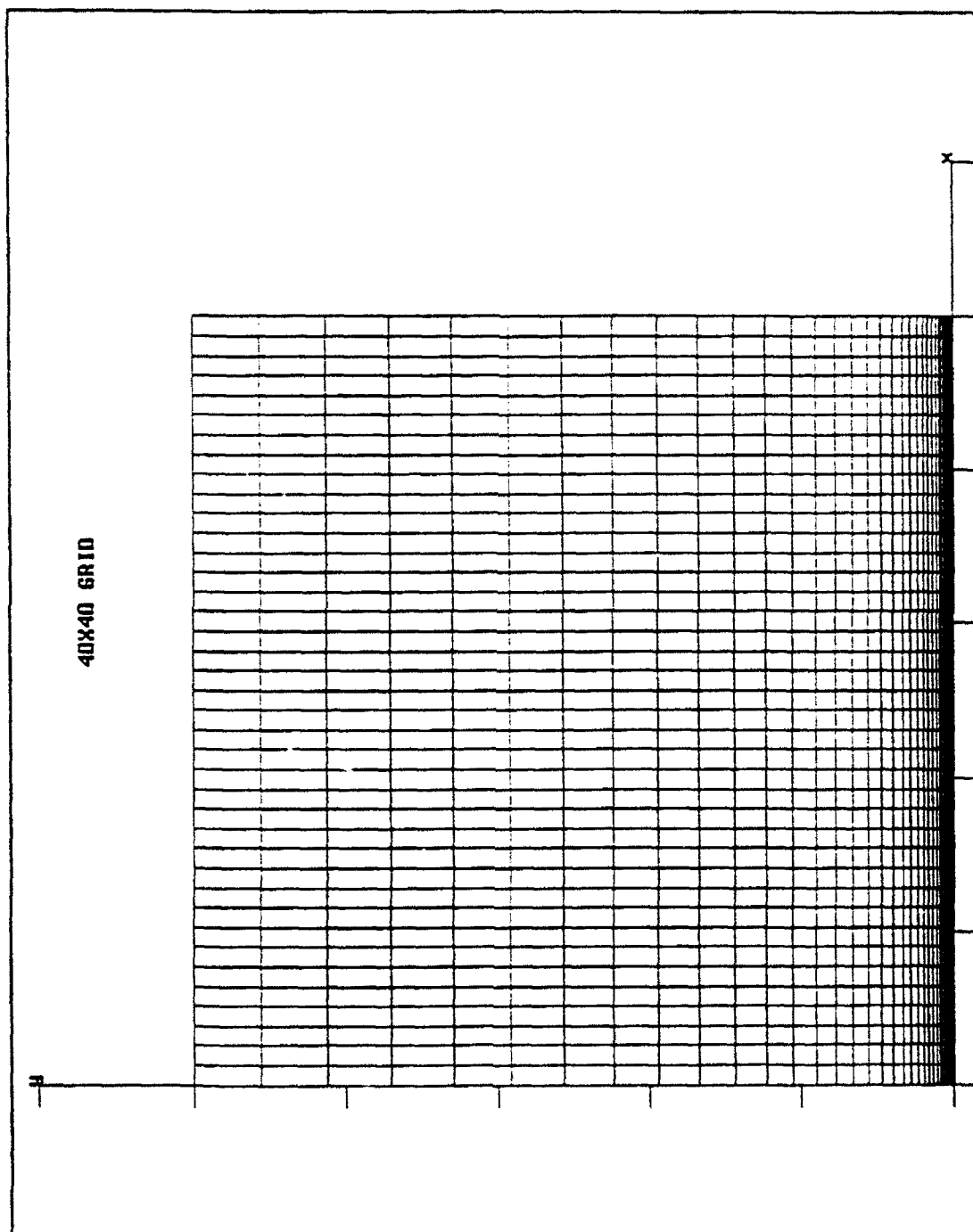


FIGURE 1

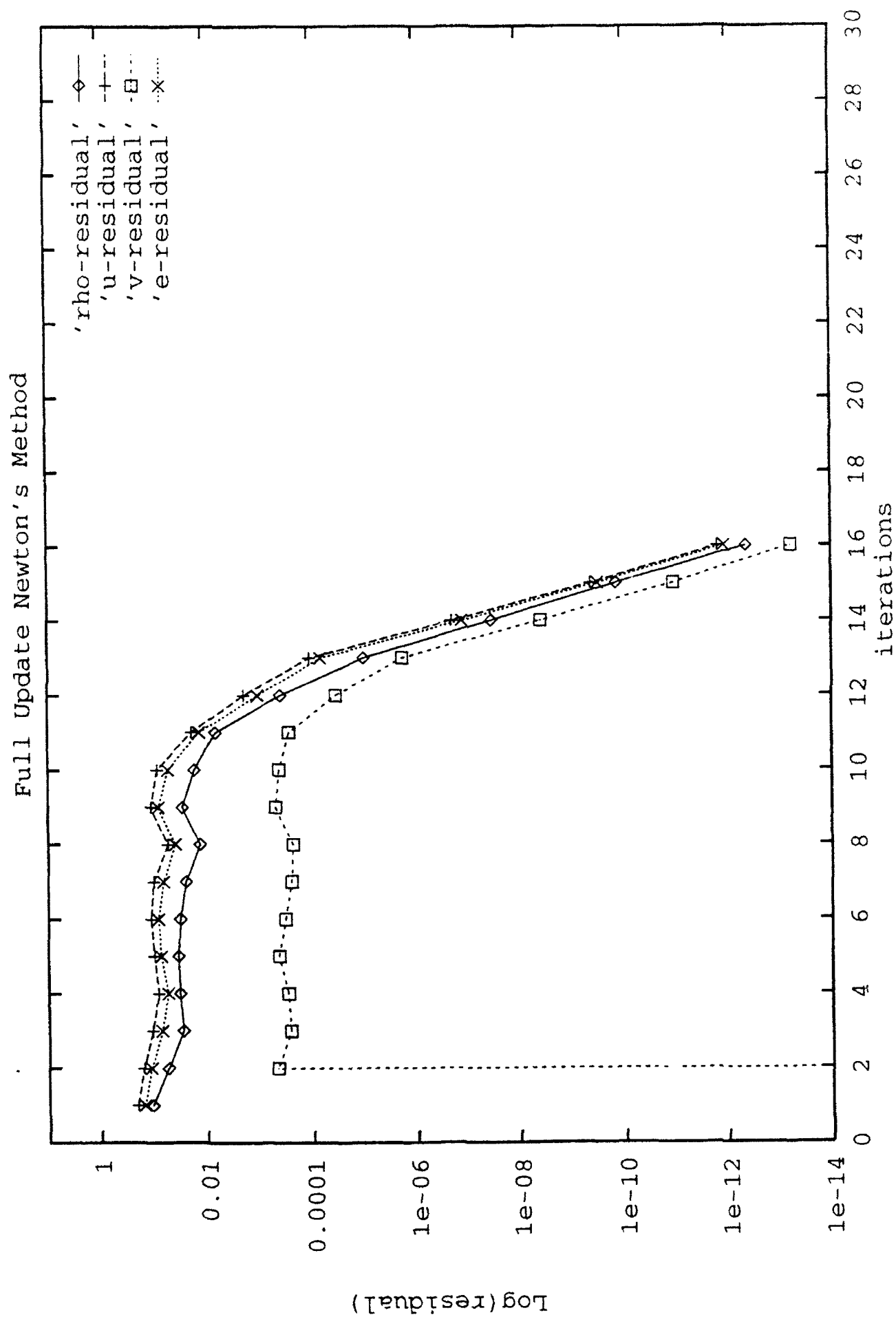


FIGURE 2:

Newton's Method: Freezing Jac.

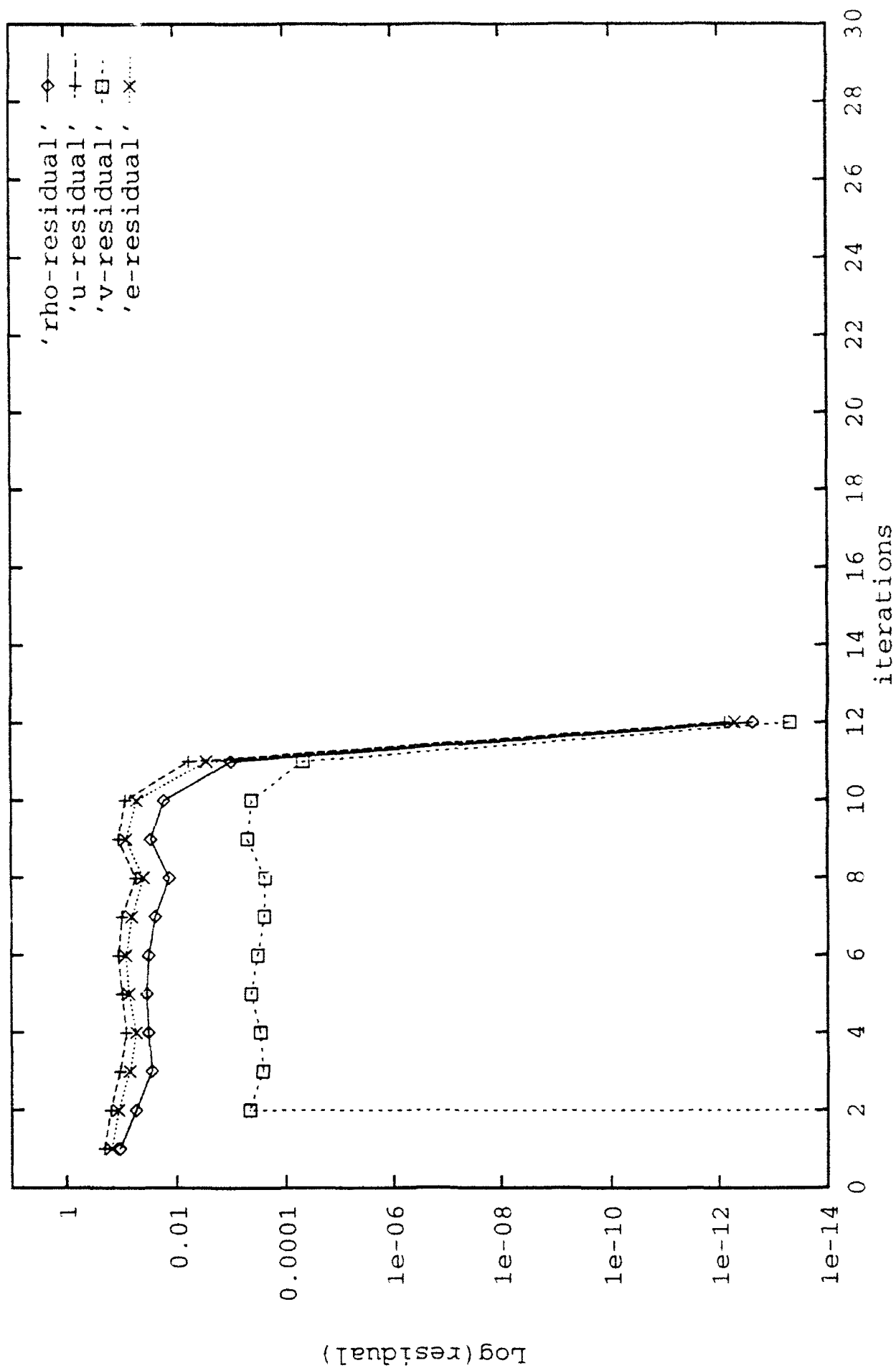


FIGURE 3:

Newton's Method: Partial Update Based on Residual

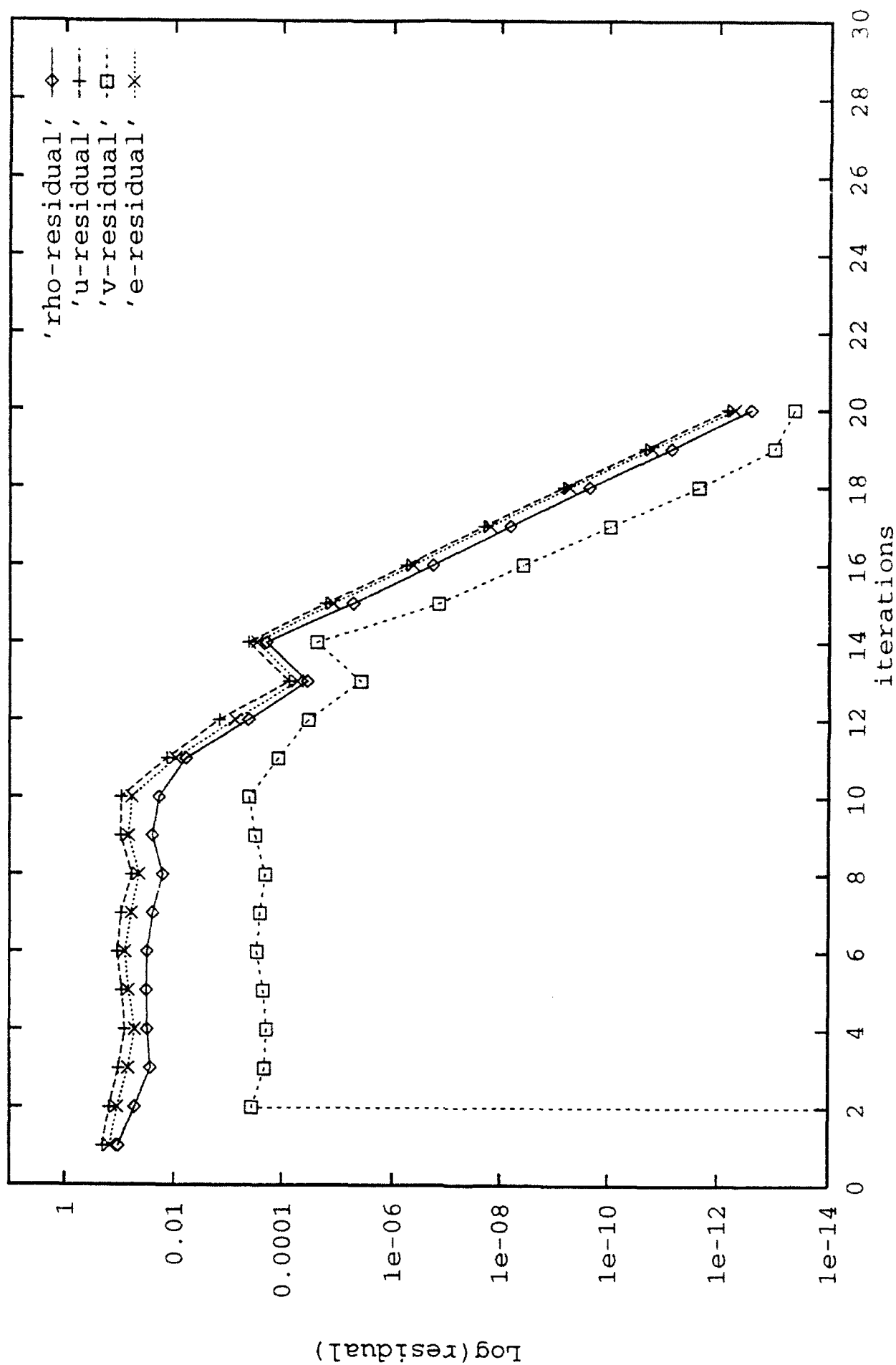


FIGURE 4:

Newton's Method: Partial Update

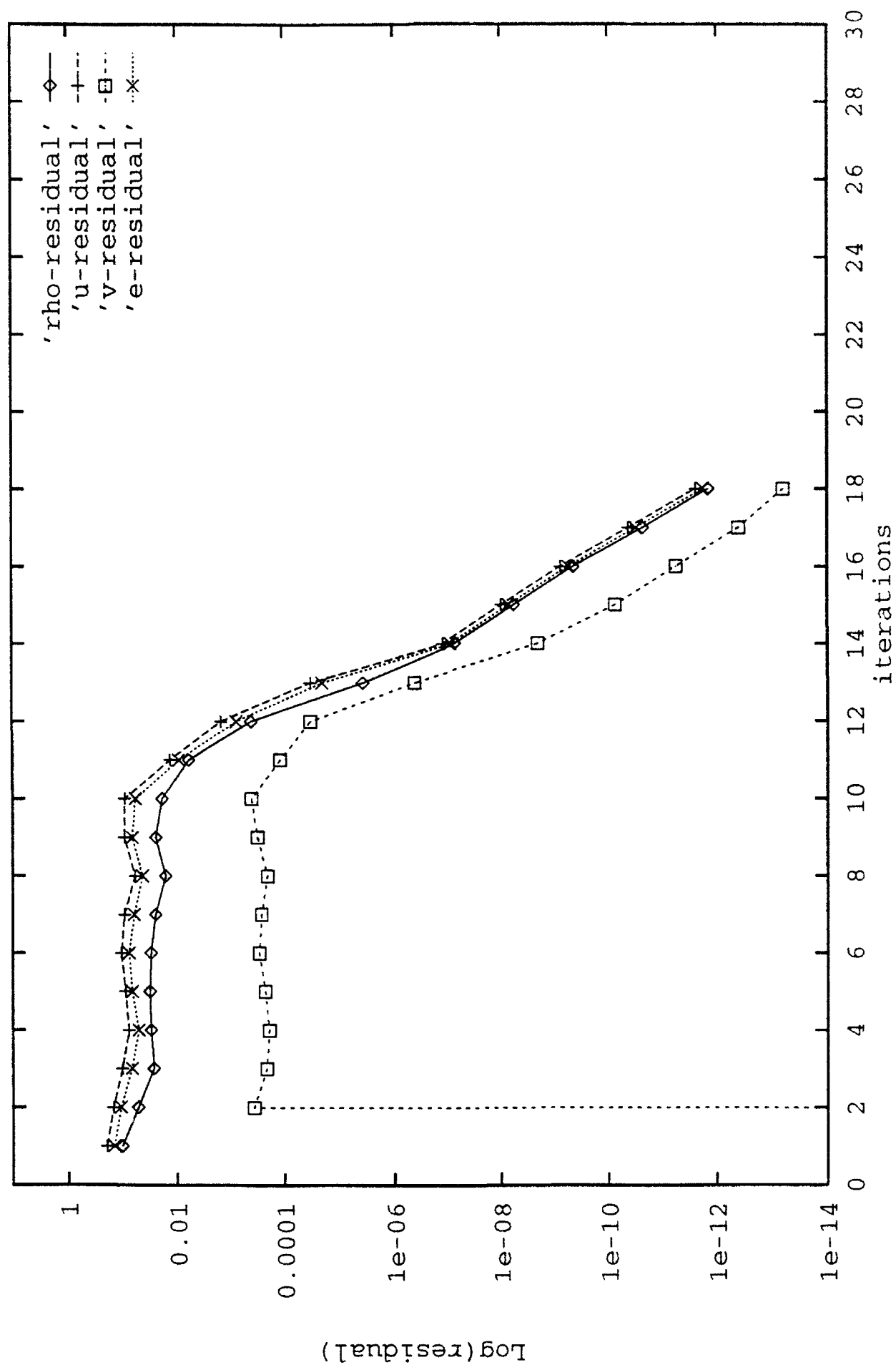


FIGURE 5 :

Newton's Method: Partial/Freeze

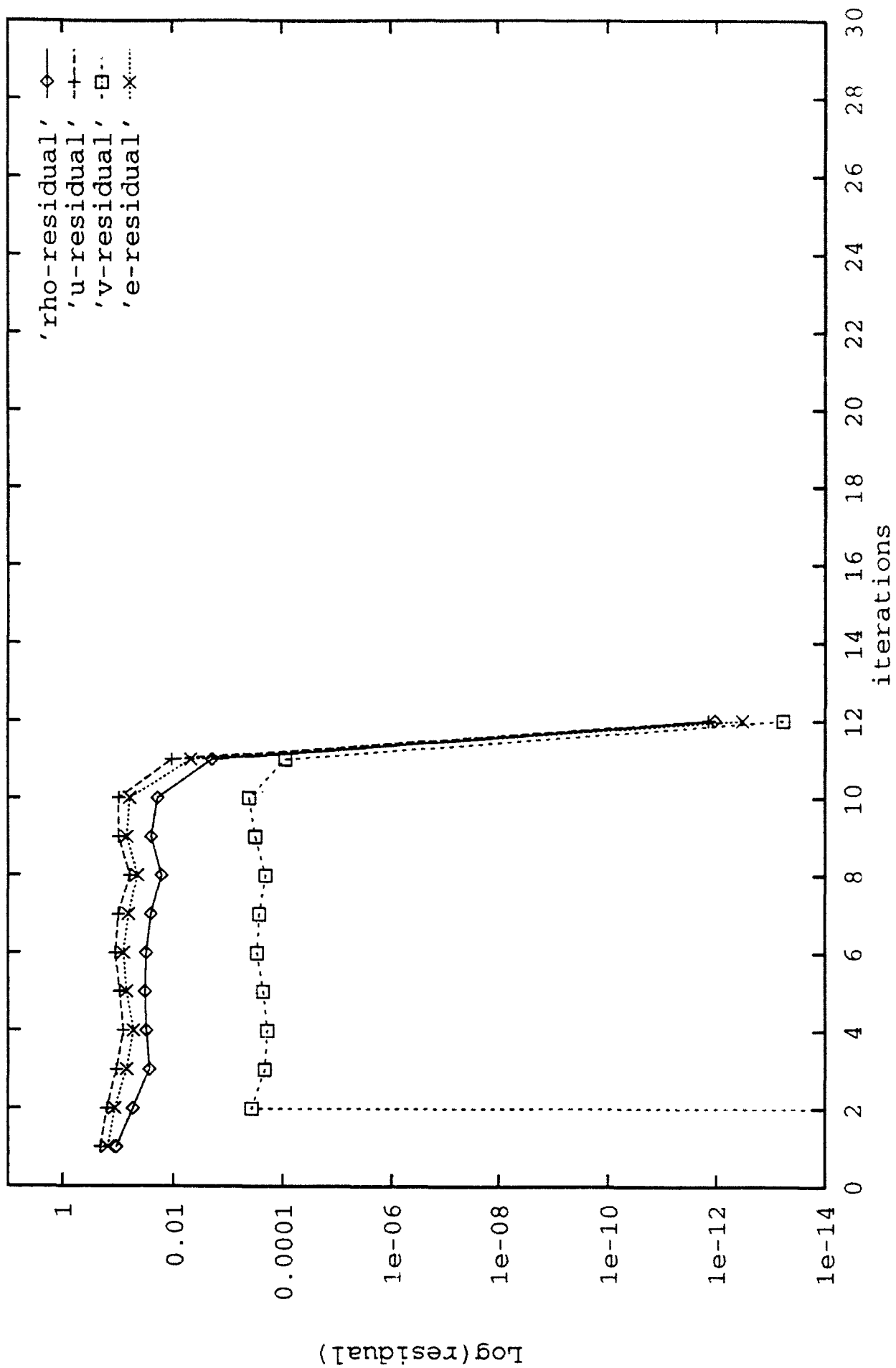


FIGURE 6 :

References

- [1] Orkwis, P.D. and McRae, D.S., "A Newton's Method Solver for the Navier-Stokes Equations," AIAA 90-1524, Seattle, Washington, June 1990, see also, "A Newton's Method Solver for High-Speed Viscous Separated Flow Fields," to appear in *AIAA Journal*, November 1991 .
- [2] Orkwis, P.D. and McRae, D.S., "A Newton's Method Solver for the Axisymmetric Navier-Stokes Equations," AIAA Paper 91-1554, Honolulu, Hawaii, June 1991, to appear in *AIAA Journal*.
- [3] Orkwis, P.D., "Newton Solver's for the Navier-Stokes Equations. " Final Report for U.S. Air Force Summer Faculty Research Program, September 1991.
- [4] Wigton, Laurence B., "Application of MACSYMA and Sparse Matrix Technology to Multielement Airfoil Calculations", AIAA Paper 87-1142, Honolulu, Hawaii, June 1987.
- [5] Venkatakrisnan, V., "Newton Solution of Inviscid and Viscous Problems", AIAA Paper 88-0413, Reno, Nevada, January 1988, see also. "Newton Solution of Inviscid and Viscous Problems", *AIAA Journal*, Vol. 27, No. 7, July 1989, pp 885-891.
- [6] Venkatakrisnan, V., "Preconditioned Conjugate Gradient Methods for the Compressible Navier-Stokes Equations", AIAA Paper 90-0586. Reno, Nevada, January 1990. see also. "Preconditioned Conjugate Gradient Methods for the Compressible Navier-Stokes Equations". *AIAA Journal*, Volume 29, Number 7, July 1991.
- [7] Venkatakrisnan, V., "Implicit Solvers for Unstructured Meshes." AIAA Paper 91-1537, Honolulu, Hawaii, June 1991.
- [8] Liou, M-S. and Van Leer, B., "Choice of Implicit and Explicit Operators for the Upwind Differencing Method," AIAA Paper 88-0624. Reno, Nevada, January 1988.
- [9] Press, W.H., Flannery, B.P., Teukolsky, S.A., and Vetterling, W.T., Numerical Recipes. The Art of Scientific Computing [FORTRAN Version], Cambridge University Press. 1989. pp 66-70.

Impact of Quasi-isotropic Composite
Plates by 1/2" Steel Spheres

John T. Lair
Graduate Student
Department of Civil Engineering

University of New Orleans
Lake Front Campus
New Orleans, LA 70148

Final Report for:
Summer Research Program
Wright Laboratory

Sponsored by:
Air Force Office of Scientific Research
Wright Patterson Air Force Base, Ohio

Impact of Quasi-isotropic Composite
Plates by 1/2" Steel Spheres

John T. Lair
Graduate Student
Department of Civil Engineering
University of New Orleans

ABSTRACT

Plates of quasi-isotropic composite materials are impacted with 1/2" steel spheres. The initial velocity and residual velocity of the projectile are recorded. Panel and plate initial and final weights are recorded.

Impact of Quasi-isotropic Composite
Plates by 1/2" Steel Spheres

John T. Lair

Acknowledgements

I wish to thank the Air Force Systems Command and the Air Force Office of Scientific Research for sponsorship of this research. Research and Development Laboratories has been most helpful in administrative matters and all other areas in which they could be of assistance.

I extend special thanks to Dr. Arnold Mayer and Dr. David Hui for their direction and support of this project. I am grateful to Greg Czarnecki, and Pat Pettit for their excellent technical advice and support. Ron Studebaker and Mark Morgan provided outstanding support in their operation of the test facilities.

Introduction

Graphite-Epoxy composite materials are composed of unidirectional layers of carbon fibers bonded together by epoxy. The anisotropic nature of composites in this case results in a material which although very strong in the inplane directions is relatively weak in the out of plane direction. The behavior of composite materials in high speed impact by projectiles differs from that of metals due to the

anisotropic nature of composites and differences in failure modes of the material. As Graphite-Epoxy composites are used in aircraft which can reasonably expect to encounter projectiles it is important to understand what damage modes exist and what the velocity thresholds for them are.

Methodology

In previous work it had been noted that there appeared to be several ways in which the composite material failed when impacted by a projectile. It was proposed that there exists a change of damage mode in composite materials dependent upon the velocity of the projectile upon impact. There was a certain amount of debate however as to why and at what velocities these changes in mode would occur and if there was two modes or three modes. The work of previous researchers had indicated the possible existence of more than one failure mode but due either to the variation in the experiment caused by the projectiles used [Pettit 1989] or variation in the material used [Altamirano 1991] it was difficult to determine at what velocities mode changes occurred. It was determined to use a projectile that presented the same surface regardless of orientation and plates of carefully controlled manufacture. If such changes in damage mode were to exist it was postulated that a change in weight of the plate would evidence such a mode change.

Test Setup

Graphite Epoxy composite plates of a [(0/90/+45/-45)₄]. layup were impacted by 1/2" steel spheres using a 1/2" smooth bore light gas gun and a 20 mm cannon with sabots to carry the projectile. The initial and residual velocities were measured by a variety of means. Break papers were used in all cases to measure velocity but these were supplemented by light screens at low velocities and a induction coil device at higher velocities. The break papers were thin papers with a continuous s-shaped line of conducting ink on them when the projectile broke the paper the circuit was broken and the time recorded. The light screens consisted of a line of light detectors opposite a light source so that when the projectile passed between them the circuit went low and was recorded. The induction coil devices were coils of transformer wire wrapped around a tube which caused a circuit to have a pulse when the steel projectile passed through. The velocities were determined by recording the relative times these devices triggered and comparing with the devices known relative position down range of the gun. Timing of the events were accomplished by use of a computerized transient analyzer. The three methods of detection were used as back checks against each other to guard against error due to equipment malfunction.

Analysis of Data

Initially a survey was made of plates which had been

impacted by a previous researcher [Altamirano 1991] and delaminated by Mr. Lair during a 1991 summer research program. This resulted in Figure 1. Area Affected by Delamination Versus Velocity by Deply Method, showing the change in delamination area versus velocity. Noting that the delamination is at a maximum at the lower velocity, it then decreases to a minimum around 2,000 ft/sec and finally rises after that at a shallower slope. Also note Figure 2, Observed Orientation of Delaminated Area Major Axis Versus Velocity, where the triangles are pointed in the direction of the ply the major axis of the delamination most closely matches, the hollow diamonds represent interfaces between plies oriented in the same direction and the squares represent delaminations that could not be associated with either the upper or lower ply of the interface. It is noted that interface 1 is closest to the impacted face of the plate. The meeting point of the projectile and the shock wave reflected off the rear face of the plate is also plotted here after previous research [Czarnecki 1991]. This chart supports Mr. Czarnecki's finding that the tensile shock wave reflected off the rear face of the plate does greatly affect the delamination between plies especially at lower velocities and that at the higher velocities the passage of the projectile has more effect. It was thought at the outset of the research that only these two failure modes existed. After completing 23 test shots at various velocities from 400 ft/sec to 6,000 ft/sec impact velocity

Figure 3, Normalized Change in Weight of Impacted Panel Versus Velocity, was created. This chart show that instead of two failure modes there must be three. Based on observation of the damaged panels the damage mode up to about 1,000 ft/sec impact velocity is primarily delamination of the plate followed by tensile failure of the delaminated plies. From about 1,000 ft/sec to about 3,500 ft/sec impact velocity the mode appears to be a punch though with delamination toward the rear face of the plate. At about 4,000 ft/sec impact velocity and above the damage mode appears as a fairly clean hole larger in diameter than the projectile with what appears to be uniform delamination through the thickness of the plate.

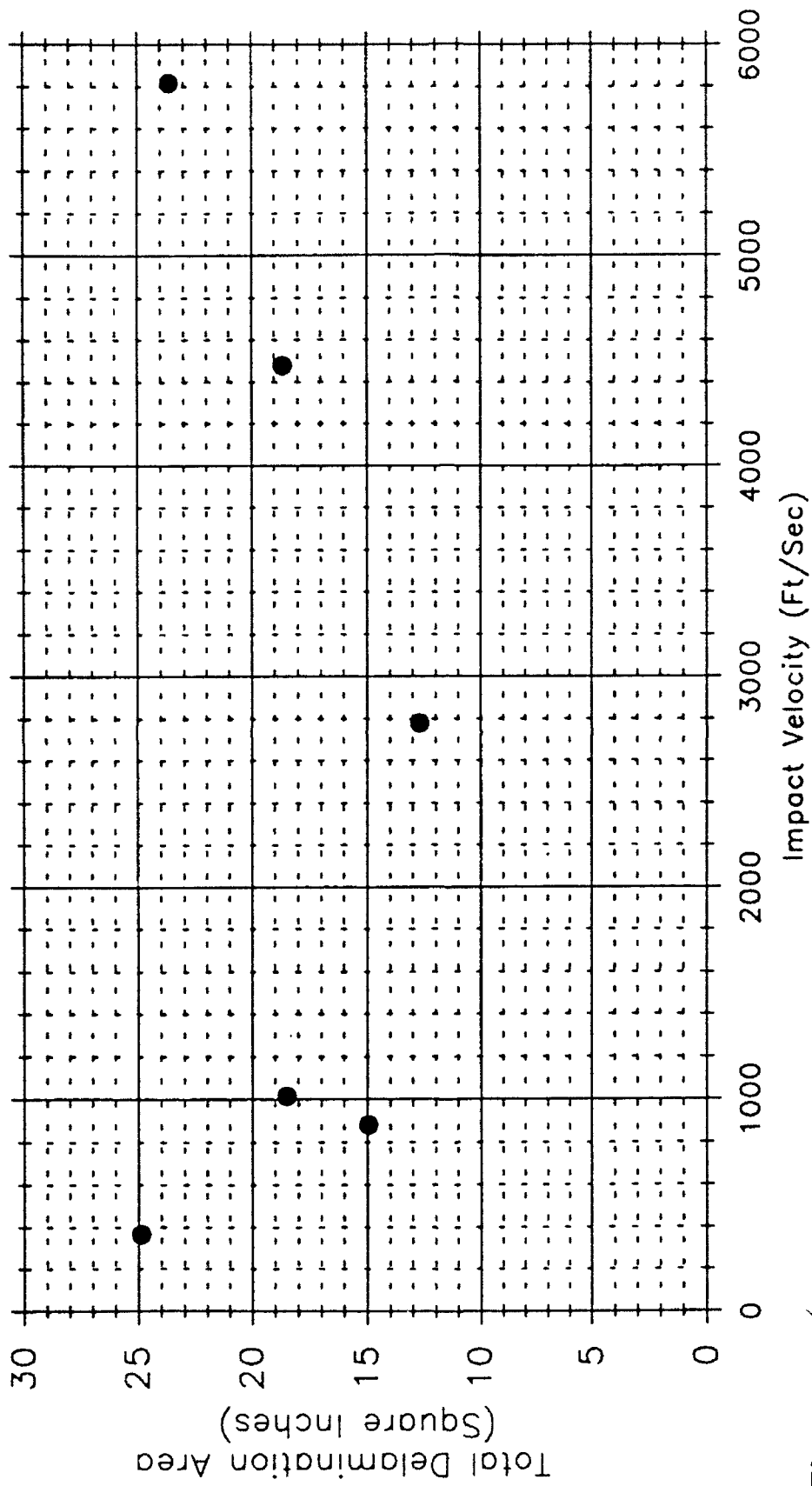


Figure 1
Area Affected by Delamination Versus Velocity by Deply Method

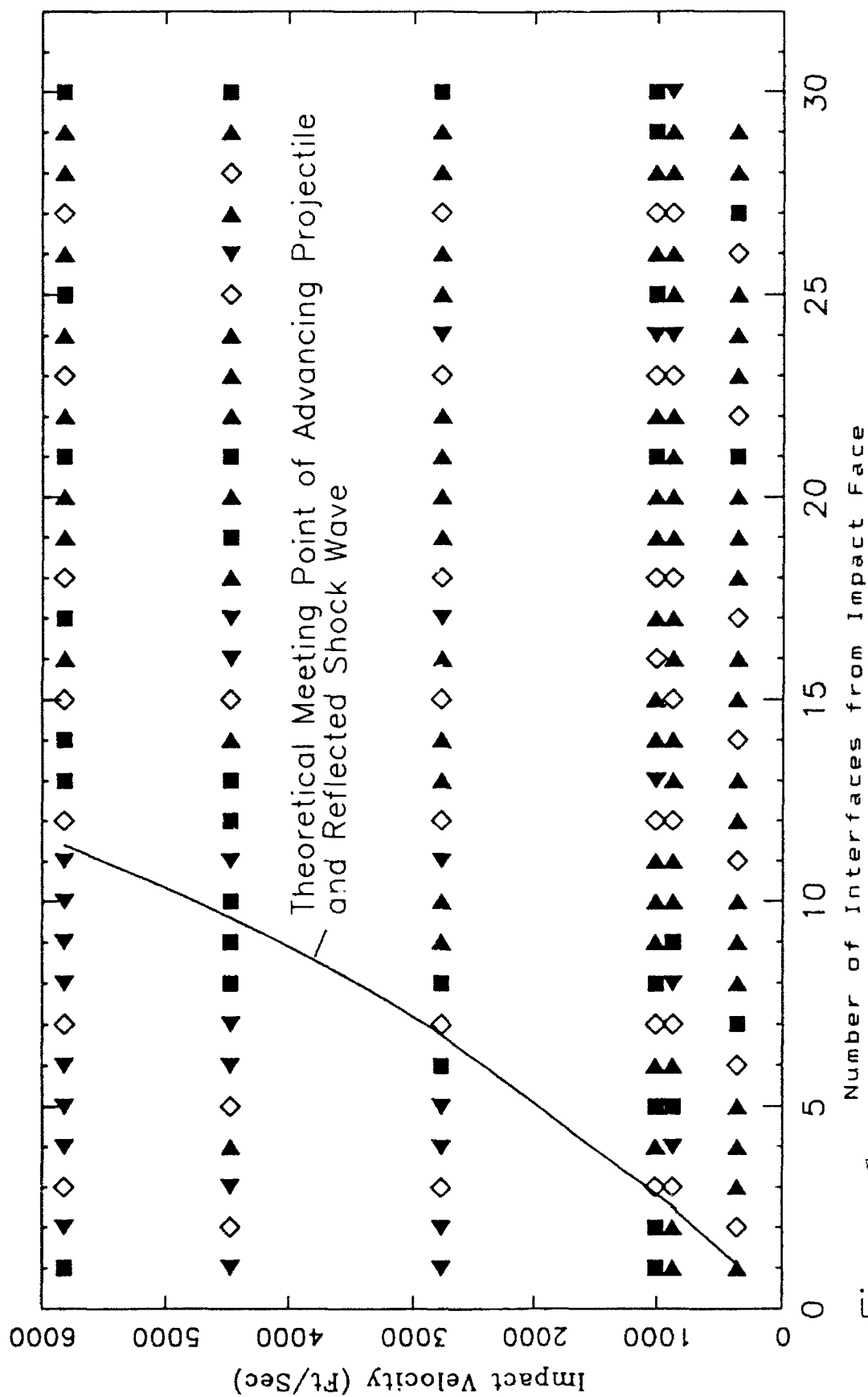


Figure 2

Observation of Orientation of Delaminated Area Major Axis Versus Velocity

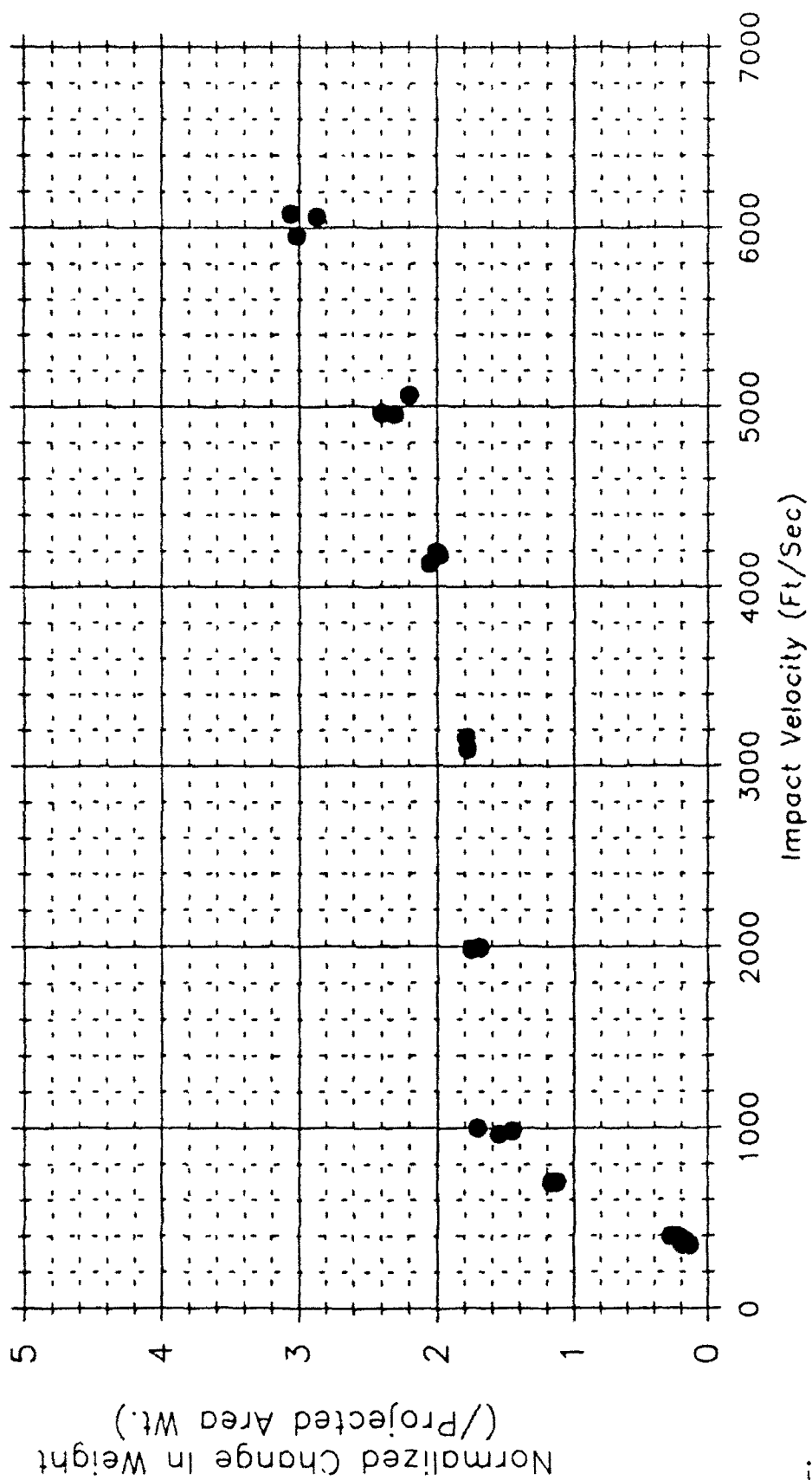


Figure
Normalized Change In Weight of Impacted Panel Versus Velocity

REFERENCES

Altamirano, Magna R., Experimental Investigation of High and Low Impact Energy Absorption of AS4/3502 Graphite/Epoxy Panels, University of New Orleans Master of Science Thesis, 1991

Czarnecki, Gregory J., A Preliminary Investigation of Dual Mode Fracture Sustained by Graphite/Epoxy Laminates Impacted by High-Velocity Spherical Metallic Projectiles, University of Dayton Master of Science Thesis, 1991

Pettit, P., Incendiary Functioning Characteristics of Soviet API Projectiles Impacting Graphite/Epoxy Composite Panels, WRDC/FIVST report, 6, September, 1989

UNIVERSAL CONTROLLER ANALYSIS AND IMPLEMENTATION

**Shawn H. Mahloch
Graduate Research Assistant
Department of Electrical Engineering and Computer Science
Arizona State University
Tempe, AZ 85287**

**Final Report for:
AFOSR Summer Research Program
Wright Laboratory**

**Sponsored by:
Air Force Office of Scientific Research
Bolling Air Force Base, Washington, D.C.**

August 1992

UNIVERSAL CONTROLLER ANALYSIS AND IMPLEMENTATION

Shawn H. Mañioch
Graduate Research Assistant
Department of Electrical Engineering and Computer Science
Arizona State University

Abstract

The universal controller, a method of controlling an "unknown system" via on-line deconvolution, was evaluated and implemented. This method of control was developed in [YBL] in contrast to the many other methods where controllers are designed based on the off-line approximate of the plant to be controlled. Preliminary simulations of the controller showed that on-line deconvolution was feasible and resulted in tracking of a reference command as close as the numerical accuracy to the instrumentation would allow. However, more in-depth evaluation using computer simulations showed that the controller consistently diverged. Numerical precision has been isolated as one of the causes of the divergence but it is believed that there exists another intrinsic cause of the divergence. A method has been developed to circumvent the controller divergence. There exists several controller variables that affect the controller performance, previously no methods to select these variables have been given. Thus based on computer simulation of many "unknown systems," methods to select these variables have been developed. Lastly, the controller was implemented to control a small motor. The controller diverged due to limited accuracy of both the motor input voltage and the measurement of the actual motor velocity.

UNIVERSAL CONTROLLER ANALYSIS AND IMPLEMENTATION

Shawn H. Mahloch

Introduction:

The universal controller, designed and referred to in [YBL] as a model reference control system, is a method of controlling an "unknown system" via on-line deconvolution. On-line deconvolution is used to generate control signals that will force the output of the plant to follow a reference trajectory. Specifically, the unit sample response of the "unknown system" is identified term-by-term at each sampling instant and uses the partial identification to calculate the next optimal control signal. The plant being controlled is said to be "unknown," but it must satisfy two constraints: (1) the plant must be minimum phase and (2) have a convergent unit sample response sequence (e.g. input-output stable). Since the controller transfer function is the inverse of the approximated plant transfer function, a minimum phase plant is required. The second constraint is necessary because the deconvolution generator, a finite memory application, can be implemented only if the plant has a convergent unit sample response. The approximated plant unit sample response can be measured as the output of the model error (em). See figures 1 and 2 for the block and schematic diagrams of the universal controller.

The controller must be started by applying a signal $A(z)$ referred to as a "starter signal." The choice of $A(z)$ can be made based on general properties of the "unknown system." Generally, the gain of $A(z)$, denoted as a_0 , is to be chosen small such that a small error sequence results. The model system is used to transition the system response since it is not desirable to have an instantaneous response to a reference command. For example, when applying the brakes on your car you do not wish to stop instantaneously, but rather at a transitioned pace. The rate of transition of the model system is determined by the parameter α and selected such that the controlled system will exhibit a desired characteristic. The forgetting factor discussed in [YBL] and [H] is not included in this application as its purpose was to limit the size of memory used in the deconvolution generator (size of memory was limited via software). It is assumed that the reader is familiar with the work done in [YBL], [H], and [DP] as all the details of the controller and previous work are not given here. The "unknown system" will hereafter be referred to as the plant for the sake of brevity.

Discussion of Problem:

When simulating the plant in [YBL] the process converged when the size of memory for the deconvolution generator was limited to 50 locations. As the size of memory was allowed to increase, the controller consistently diverged, see [H]. Further investigation by [DP] confirmed the problem as well as uncovered others. When a_0 , α , and the size of memory were varied both individually and in groups the process similarly diverged, sometimes faster than others. Initially it was thought that mathematical operations involving very large and very small numbers was

the cause of the divergence problem. However, upon careful examination of the numbers being mathematically operated on, it was determined that this was not the problem as all the numbers were within three to seven orders of magnitude of each other.

Numerical precision, how accurately a computer represents a floating point number, was also another possible cause of the divergence problem as the memory values stored in the deconvolution generator become consistently smaller in magnitude. To test this theory the software would be tested in different computing environments. Initially, the software was written in "C" and hosted on an Intel 80960KB processor. The software was then rehosted on a PC and on a VAX. Since "C" on the VAX and the Intel 80960KB have approximately the same numerical precision, no significant differences were observed. However, PC's have reduced precision of floating point numbers and it was suspected that the universal controller would diverge more rapidly. Simulation of the test case on the PC confirmed this theory. The software was then rewritten in Fortran and hosted on the VAX where quad precision floating point representation can be used. This implementation of the software improved the performance of the universal controller but did not prevent the divergence problem when memory was increased. Given the above results, numerical precision was deemed to be a certain constraint to the universal controller application.

To further verify numerical precision as a constraint, a numerical precision reduction algorithm was developed to mimic the effect of reduced precision. This process was tested in "C" on the VAX. When varying different degrees of precision, the universal controller diverged more rapidly as the precision was reduced. See [DP] for greater details of the numerical precision reduction algorithm. The above results further confirmed numerical precision as a constraint on the controller with no obvious solution apparent.

Methodology:

Assuming that the theory behind the universal controller was accurate, as it has been reviewed thoroughly several times, the investigation of the divergence problem began with a review of the previous work described above. After review of the previous results, several questions remained. (1) Since floating point numbers can be represented accurately to at least 15 decimal places (using "C" on the VAX and on the Intel 80960KB processor), is the error introduced due to these representations a cause of the divergence? If these types of errors are not the cause of the divergence, what is the cause? (2) The controller diverges primarily after the plant has been "significantly" identified. That is, typically the controller has forced the plant to follow the reference trajectory with a small model output error, but the divergence occurs when attempting to further reduce the error. Therefore, is it reasonable to stop identifying the plant (which corresponds to stop storing memory values in the deconvolution generator and thus truncating the unit sample response sequence of the plant) when the controller begins to diverge and still have an acceptable model output error? (3) What affect does the reference trajectory, sample time, and the parameters a_0 and α have on the divergence of the controller? The answers to these questions are what provided the motivation and direction of the project.

Analysis of Controller Divergence:

In an attempt to resolve the question regarding errors introduced due to the representation of floating point numbers, one needs to consider the accuracy of the numbers and the type of mathematical operations being performed. The output of the controller is the result of the deconvolution which involves many multiplication and addition operations (see figure 2). It is known that floating point numbers can be represented accurately to at least 15 decimal places on both the VAX and the Intel 80960KB processor board. Given that all the numbers involved in the deconvolution have 15 decimal place accuracy and are within three to seven orders of magnitude of each other, then the result of these mathematical operations will have 15 decimal place accuracy. Recall that in the deconvolution generator that the stored memory values are multiplied by past controller outputs. The memory values should be continually approaching zero whereas the controller outputs should be leveling out to a constant value. Assume that the past controller outputs, u_n , are such that $0.01 < |u_n| < 10$. Then with this assumption, errors due to numerical precision would not become a factor until memory values of magnitude $10e-13$ were reached. Representational errors would not become a factor until this small of memory value is reached because multiplying this value with the past control signal could possibly result in a number of magnitude $10e-15$ where this type of error may be significant. Thus in [DP] when reduced precision was mimicked, the controller failed much sooner because representational errors became a factor sooner (e.g. numbers of magnitude $10e-15$ could no longer be represented accurately). Similarly when quad precision was used, the divergence was postponed because numbers smaller in magnitude than $10e-15$ could be represented accurately. Given the above considerations, errors due to floating point representation are not the cause of the divergence until memory values of such small magnitude ($10e-13$, depending on computing environment) are reached.

When simulating the plant in [YBL] on the VAX (at least 15 decimal place accuracy), it is observed that the controller diverges when memory values of magnitude $10e-5$ are reached. Upon observing the numbers that are being mathematically operated on, these memory values could be added or multiplied by a number no smaller in magnitude than 0.01. Clearly, errors due to numerical representation are not a cause of the divergence of the controller in this case. The cause of the divergence was not able to be resolved, thus efforts were directed at modifications of the controller. The process of developing these modifications was to test many plants varying parameters individually and in groups, observe reoccurring patterns, and develop changes for the controller and methods for selecting the parameters.

A modification desired was one that would circumvent the divergence of the controller. Recall that initially the controller identifies the plant, thus during each iteration the plant is further identified and stored in memory. As the controller continues to identify, the memory values should become smaller in absolute value continually approaching zero, but eventually during the simulations the memory values began to increase in absolute value and thus the controller began to diverge. A question asked then was, can identification of the plant be stopped before the controller diverges and an acceptable model output error still be attained? One way to implement this would be to fix the size of memory used, but this would not be a general solution for any plant as the point of divergence would

vary. A solution that works for any plant is to monitor the convergence of the memory values and when divergence is detected, memory values would no longer be stored. The method of detecting when the controller begins to diverge requires some explanation. The memory values being stored are approaching zero, but the memory values typically oscillate between positive and negative values with decreasing amplitude of oscillation. Therefore one can not simply monitor for the smallest absolute memory value and stop storing values as soon as the next largest absolute memory value is detected. Typically, one must monitor for the occurrence of an increasing amplitude of oscillation and then save only the memory values up to the smallest oscillation.

This slight modification is advantageous in several ways. The number of memory locations used is now determined by the controller to yield the best result (an upper limit on the size of memory is still an option implemented in the application). The forgetting factor discussed in [YBL] and [H] is now eliminated. Lastly, the divergence problem has been eliminated, but now how accurately the controller forces the plant to follow a reference trajectory is an important concern. It is with this concern that methods for selecting the parameters have been developed.

Analysis of Controller variables:

Recall that the starter signal $A(z)$ given in [YBL] has a gain coefficient a_0 to be chosen. It should also be noted that a_0 is a gain in the feed forward path (see figure 2). When varying a_0 for the same plant, the point of divergence for the controller varied considerably. As the choice of a_0 affects the divergence so drastically, it tends to lead to a question addressed when approximating nonlinear systems, which deals with local and global convergence of the initial conditions (e.g. Can any a_0 be picked and still result in a converging identification of the plant, or must a_0 belong to a given range of numbers?). But for linear stable systems, local convergence implies global convergence. Therefore, the choice of a_0 should not affect the convergence of the universal controller, although it would affect the rate of convergence. While the explanation of a_0 and its affect of the divergence of the controller where not resolved, insight was gained from varying the choice of a_0 . Recall from [YBL] that $|a_0|$ should be chosen small such that a small error sequence results, but that a lower limit on $|a_0|$ exists due to controller implementation. For a small error sequence to result, a_0 should be selected such that $|a_0| < 1$, $a_0 \neq 0$. From many computer simulations, the choice of a_0 should satisfy two additional conditions. The parameter a_0 is the cause of the first plant output and thus directly effects m_1 , the first memory value (note that $m_1 = 1/e_m(1)$ and that the magnitude of a_0 inversely affects the magnitude of m_1). It has been consistently observed that if m_1 is approximately in the range $0.5 < |m_1| < 10$, then as a rule-of-thumb the best results (smallest model output errors) are yielded. Generally, if $|m_1|$ is large it is expected for the model output error to oscillate more before settling out towards zero; and if $|m_1|$ is small then it is expected for the model output error to gradually (without oscillations) approach zero. Both cases of a_0 being too small or large would require more memory locations to identify the plant. Typically, it has been observed that the universal controller diverges quickly if $|m_1|$ is large whereas if $|m_1|$ is small it still identifies but takes much longer to do so.

The magnitude of a_0 is not the only factor that affects the performance of the controller. The sign of a_0 apparently has a significant affect on how well the universal controller performs. For example, on the test case evaluated in [H] with $a_0 = 0.1$, the controller begins to diverge after the 36th memory value. If plant identification is stopped at this point, the model output error settles out at 5.0%. For the same case, if $a_0 = -0.1$, then the controller does not begin to diverge until after the 96th memory value. Similarly, if plant identification is stopped at this point, the model output error settles out at 0.004%. This phenomenon has been noted with all of the plants that were examined. An adhoc method of choosing the sign of a_0 , other than running several trials of the controller, has been identified. If one has information regarding the open loop step response of the plant, the sign of a_0 can be chosen. If from the open loop step response the plant output settles out at a positive value then a_0 should be chosen positive, otherwise a_0 should be chosen negative (this is with the assumption that the reference trajectory is chosen positive during plant identification, otherwise the opposite is true). In summary, a_0 should be chosen such that (1) $|a_0| < 1$, $a_0 \neq 0$; (2) $0.5 < |m_1| < 10$; and (3) with the proper sign as determined by the adhoc procedure described above.

The sampling rate also plays an important role in the performance of the controller, therefore the length of the transients of the plant to a command must be known. From many computer simulations of different plants it was determined that the best performances were typical of a sampling rate where the first 50-100 samples capture most of the "significant" transients. It was observed that if the system sampling rate was much slower or faster than this then the controller would quickly diverge.

The reference trajectory during plant identification and the parameter α also have slight affects on the universal controller. The reference trajectory should be chosen small when identifying the plant as this also has been observed to yield better results. After the plant has been identified then the reference trajectory can be chosen as desired. The parameter α is chosen to indicate how fast one desires the system to respond to a command. If the other parameters have been selected as set forth in this report, α can be selected as desired with only minimal affect on the performance of the controller.

Implementation:

The goal of this project was to implement the universal controller using a motor as the plant. The software for the controller was written in "C" and hosted on an Intel 80960KB processor. The speed of the motor was measured using an encoder and a HCTL-1000 general purpose motion control IC. The 80960KB processor board contained a prototype area where the HCTL-1000, decoding logic, timing logic, and a 12-bit DAC were mounted. Since the motor input voltage range was $\pm 10V$, the DAC was configured to the same voltage range which would indicate an output voltage resolution of 4.88mV. The minimum voltage to change the speed of the motor was observed to be about 2-3mV. The encoder on the motor outputs 500 square wave pulses per rotation and the HCTL-1000 counts the pulses that occur during an interval of 2048 μ S. The accuracy of the encoder was observed to be ± 1 count per interval. From the open loop step response of the motor it was determined that a_0 should be positive and the length

of the transients were about 0.8 seconds. A sampling period of $T = 0.01$ seconds was selected.

Upon testing, the controller failed to identify or control the motor (e.g. the controller diverge almost immediately). Before testing the universal controller on the motor there was concern regarding the errors being introduced from the measurement of the motor speed as well as the resolution of the DAC was not sufficient. It was thought that the result of these errors would result in a larger tracking error but that the controller would still work. One might ask why the voltage of the DAC was not limited to $\pm 5V$, thus giving a voltage resolution of 2.44mV, and to just operate within this range. The answer to this is that the motor has a $\pm 1V$ deadzone. Therefore the $\pm 5V$ DAC output range was not acceptable as it was desirable to keep the control signal out of the deadzone. Initially it was thought that the deadzone phenomenon was the cause of the failure. The reason for this is that when a few of the control signals fell within the deadzone then the plant output goes to zero and the plant is then not operating linearly as expected, thus introducing errors into the identification of the plant. After further evaluation and choosing the starter signal sufficiently large such that no control signals fell within the deadzone, the controller still diverged, thus indicating that the deadzone may not be the cause of the failure.

The investigation quickly turned to the accuracy of which the plant output can be measured and the accuracy of the motor input voltage versus the computed value. The input voltage to the motor and the measured motor velocity have at most three significant digits of accuracy. Since it was difficult to ascertain any meaningful conclusions regarding the accuracy when using the motor as the plant, answers were sought using computer simulation. When simulating an "unknown system" and truncating all numbers to three significant digits of accuracy, the controller diverged almost immediately. When testing the same "unknown system" with four and five significant digits of accuracy, the model output error was approximately 50% and 10% respectively. Thus, to get the controller to identify and force the motor to follow a reference trajectory, the accuracy of both the motor input voltage and measurement of the actual motor velocity will have to be at least 5 significant digits. It was observed that the controller variables (especially a_0) have a much greater affect on the controller performance with reduced accuracy. To achieve a model output error of 1% or less, it is estimated that 8 significant digits of accuracy are required.

Conclusion:

In summarizing the results gathered from this project it seems that the universal controller would have limited application due to the accuracy requirements necessary of the output and input to the plant. While further evaluation is needed to completely confirm this theory, it is believed that without sufficient accuracy that too much error is introduced into the on-line identification of the plant.

Before this investigation began it was believed that the error introduced due to how a computer represents floating point numbers was the cause of the divergence problem. During the simulations when extremely small memory values on the order of $10e-13$ (depending on computing environment) are achieved, then it is agreed that the error introduced due to the effect of how the number is represented is an issue. However, during many simulations the controller diverged well before memory values of such small magnitude were reached. Thus, it is believed that numerical precision is not the cause of the divergence problem, but rather that there is some other intrinsic cause of the problem. A method was developed to circumvent the divergence problem by detecting when the controller begins to diverge and saving only the converging memory locations. Methods have also been developed to select the controller variables, as the selection of these variables significantly affects the controller performance. Computer simulations, using the above changes, have shown that an "unknown system" can still be forced to follow a reference trajectory, achieving extremely small model output errors.

References:

[YBL] Hsi-Han Yeh, Siva S. Banda and P.J. Lynch, "Control of Unknown Systems Via Deconvolution," Journal of Guidance, Control, and Dynamics, Publication of American Institute of Aeronautics and Astronautics, May-June 1990, Vol. 13, No. 3.

[H] Thomas Hummel, "Universal Controller Computer Simulations," Technical Report, Flight Dynamics Laboratory, Wright-Patterson AFB, Ohio, January 1992.

[DP] Thomas F. Dermis and Capt. Donald L. Pagoda, "Universal Controller Implementation," Technical Report(Draft), Flight Dynamics Laboratory, Wright-Patterson AFB, Ohio, May 1992.

BLOCK DIAGRAM OF UNIVERSAL CONTROLLER

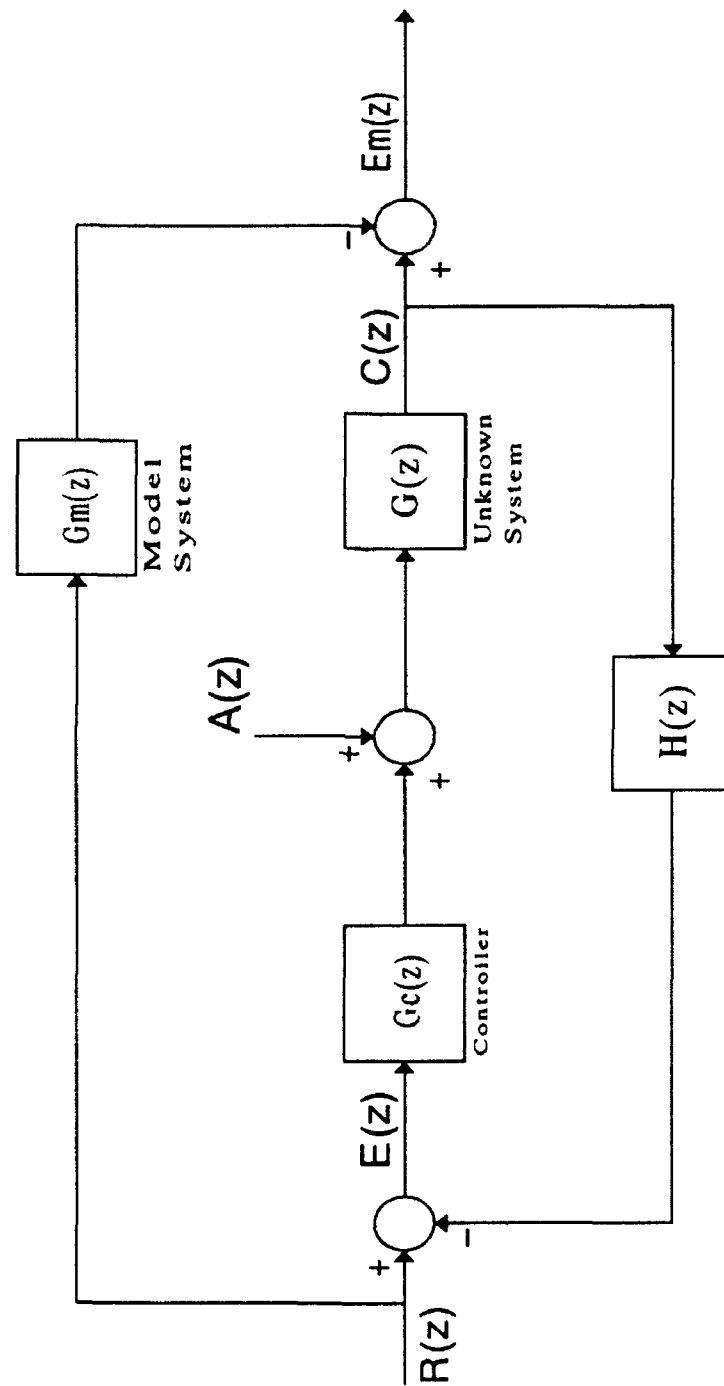


Figure 1

Schematic Diagram of Universal Controller

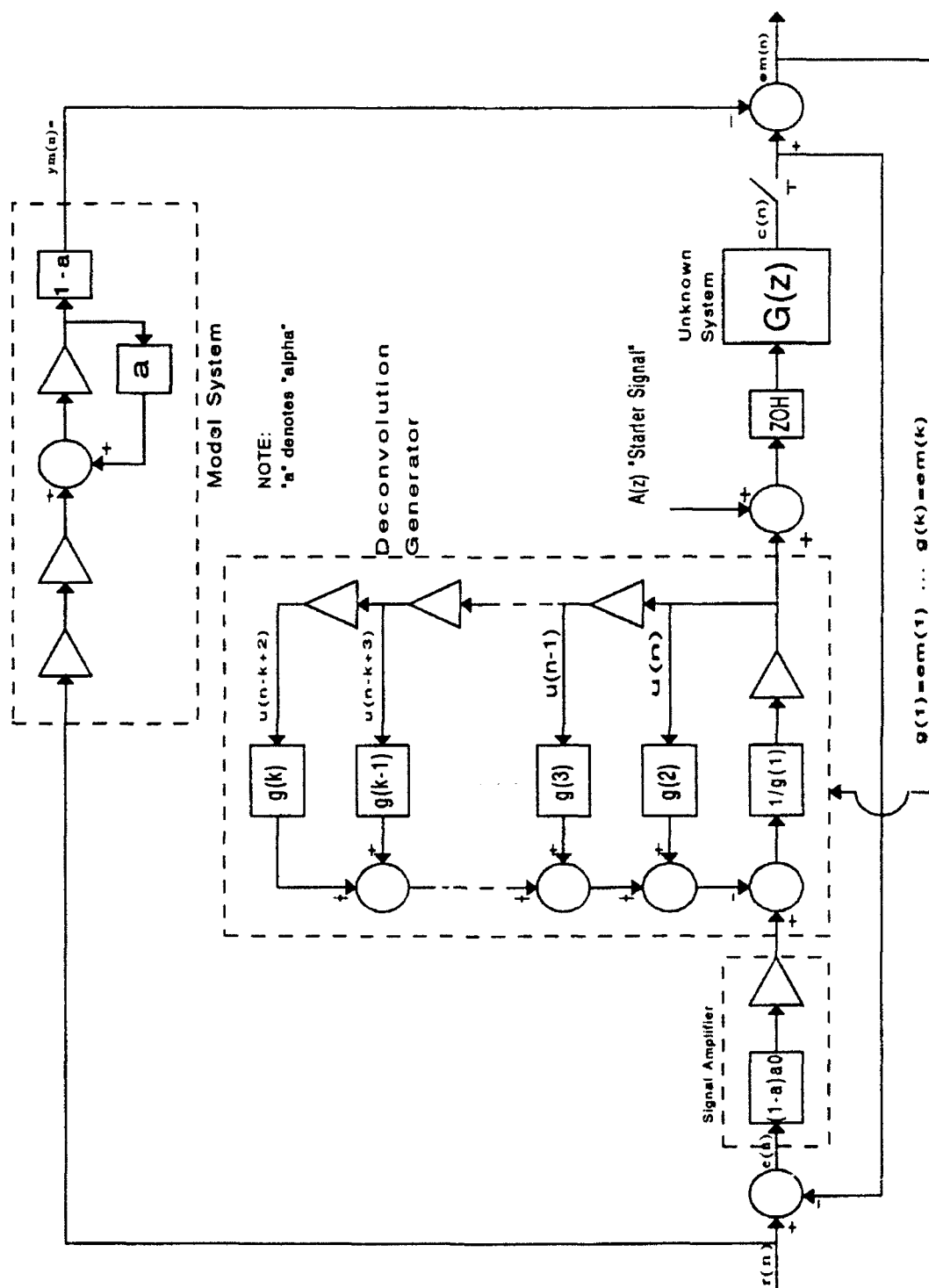
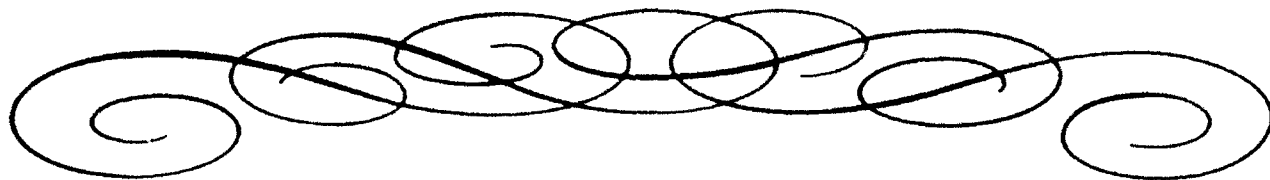


Figure 2



Process Migration Facility for the QUEST Distributed VHDL Simulator

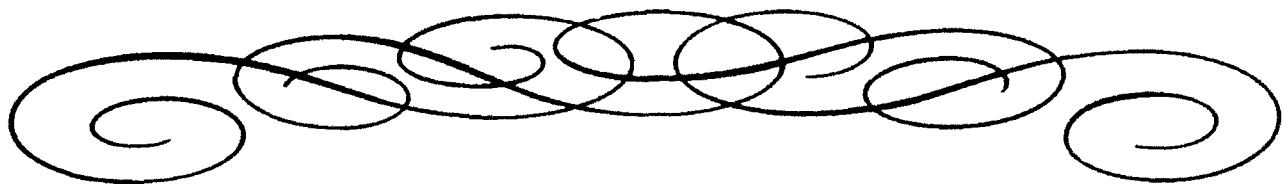
Dallas J. Marks
Graduate Student
Department of Electrical and Computer Engineering

University of Cincinnati
Cincinnati, OH 45221

Final Report for:
Summer Research Program
AAA/2

Sponsored by:
Air Force Office of Scientific Research
Wright-Patterson AFB, Dayton, OH

September 1992



Process Migration Facility for the QUEST Distributed VHDL Simulator

Dallas J. Marks

Graduate Student

Department of Electrical and Computer Engineering

University of Cincinnati

Cincinnati, OH 45221

Abstract

The QUEST VHDL Simulator is a DARPA-supported distributed simulator written for execution on an ES-Kit multiprocessor by researchers at the University of Cincinnati [5]. A recent Air Force funded project ported the simulator and its compiler to a network of Silicon Graphics 4D workstations shared among several research groups at Wright Patterson Air Force Base. High performance simulation is achieved by the distribution of VHDL simulation objects [10] and the use of local shared memory, ethernet, and Scramnet, a high-speed shared memory networking system. [4] The QUEST VHDL simulator was needed due to the sheer size and complexity of simulations soon to be required by the Air Force. The Cockpit Avionics office will use the QUEST VHDL simulator during the design of specialized hardware for cockpit display generators to be used in next-generation transports and fighters and in a retrofit to the Advanced Tactical Fighter [11]. A full simulation will often take several hours or even days to complete. Because the machines are shared, a long-running simulation creates unacceptable machine loading for other groups in the lab that must perform real-time flight simulations.

The thrust of summer research has been focused toward the design and development of a graphical network management tool to allow dynamic reconfiguration of the network. When complete, this work will provide the lab with the ability to control and restrict the machines available to the VHDL simulation. Specifically, this utility has the ability to checkpoint a running process on one machine, migrate it to another machine already in use, and restart it without any apparent interruption to the user or other simulation objects.

This research draws from similar work at the University of Wisconsin with the Condor Distributed Batch System [6]. However, the Condor system can only migrate single-process, non-communicating programs. In addition, Condor provides no support for shared libraries, which are quite useful in reducing the size of simulation objects.

This report highlights background work and summarizes the research completed during the summer research period.

Process Migration Facility for the QUEST Distributed VHDL Simulator

Dallas J. Marks

1.0 Introduction

The QUEST VHDL Simulator is a DARPA-supported distributed simulator originally written for execution on an ES-Kit multiprocessor by researchers at the University of Cincinnati [5]. Distribution and the resulting parallel execution of multiple objects provides greater performance than can be achieved by running simulations on a single processor. Distribution is based on the granularity of VHDL objects. A distributed VHDL simulator is needed also by the Air Force due to the sheer size and complexity of planned designs. The Cockpit Avionics office at Wright-Patterson AFB will use a VHDL simulator during the design of specialized hardware for cockpit display generators to be used in next-generation transports and fighters and in the retrofit to the Advanced Tactical Fighter [11]. The QUEST simulator has been recently ported to a network of Silicon Graphics 4D workstations. The design and implementation of this simulator has been outlined in [3, 10]. Its main features are the use of shared libraries for reduction in code size, and the use of a unique message-passing interface that transparently exchanges messages between objects using local shared memory; ethernet; or the high-performance Scramnet, a high-speed shared memory networking system [4].

The current simulation environment is a network of Silicon Graphics 4D workstations, shared by various groups within the Cockpit Avionics Office AAA-2 and XPK Joint Cockpit Office. Other work in the office involves the real-time simulation of cockpit environments. These real-time simulations require the full support of their host processors, meaning they cannot co-exist with the VHDL simulation objects. It is therefore necessary to be able to restrict the VHDL simulation's access to certain machines on the network during real-time simulation.

The thrust of the summer research has been the design of a network configuration tool to allow users of the network to control which machines are available to run VHDL simulations. Currently, the network configuration tool allows configuration of the network while the simulator is dormant; this work has been completed. Work is ongoing to support dynamic reconfiguration of the network during a VHDL simulation. Dynamic reconfiguration requires the capability to transparently move running process from a restricted machine to an available machine without side-effects or dependency on the user of the simulator. This capability is known as *process migration*.

2.0 Background and Related Work

The basis for this work is found in recent work at the University of Wisconsin in the Condor Distributed Batch System, [2, 6, 7, 8], designed by Litzkow. Condor was developed for the explicit purpose of executing single-process, long-running, computationally intensive jobs. The Condor system resides outside the Unix kernel and allows the *checkpointing* and *migration* of Unix processes.

Checkpointing refers to the halting of a process and storing its current state. This is achieved under Unix by signaling the running process which then performs a core dump and halts. The resulting core file (also known as a *core image*) contains the process state at the time of termination. The process state consists of the stack data, and other information. A new executable file, called the *checkpoint file*, is created. It is composed of code from the original executable and data retrieved from the core image. This process is illustrated in Figure One.

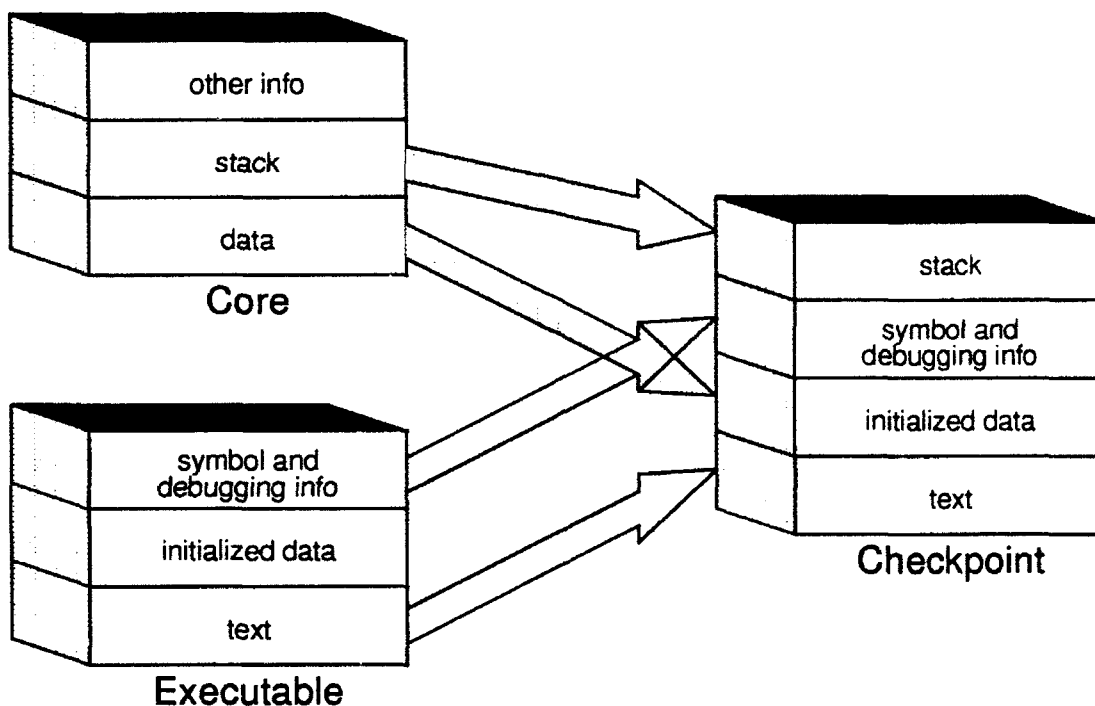


Figure 1: Creating a Checkpoint File

Migration refers to the moving of this modified executable file to another workstation and restarting its execution. Because of the modifications made during checkpointing, the process will continue executing where it was previously halted; the user perceives no break in processing.

Some modifications are required to allow an executable program file to be migrated. Fortunately, these modifications are minor. The main additions are two functions. The first is executed at startup to redefine a Unix signal. This signal will be used to inform the program to terminate and dump its core. The second function initializes the executable file upon restart after migration to start execution where the program was interrupted. These functions will have little effect on code size of simulation objects and their use involves only relinking of simulation code, not recompilation.

Condor performs its work outside the Unix kernel to enhance portability. Although the formats for executable and core files vary among different architectures and systems, Condor has been ported to ten platforms, including the Silicon Graphics 4D. However, the Condor system has its limitations. Some Unix features such as multi-process applications, signals, interprocess communication, and shared libraries are currently unsupported [9]. However, it is necessary that these features be added to support the QUEST VHDL simulator.

3.0 Summary of Recent Work

It is the primary goal of this research to implement a process migration facility that allows both the user of the VHDL simulator and other users of the network to have full and unrestricted access to run their respective applications. Control of the migration will be achieved through the use of a graphical user interface on the simulation host machine and by remote shell execution for other users. Once changes in the network configuration are made, the migration utility will notify the simulation (through its main, or *creator* object) using the simulator's shared memory interface. The creator will inform all objects to discontinue communication to and from objects that exist on the restricted machines. When all necessary objects are migrated from restricted machines to available machines, communication to these objects will resume normally.

An important criteria of the process migration scheme will be the speed of migration. The temporary inability to communicate with migrating objects will result in time rollbacks and recomputation of simulation events. However, unlike the Condor project, periodic checkpointing is not necessary; a simulation object will only be checkpointed prior to migration.

A special concern for this project will be the existence of *shared libraries* in the simulation objects. Shared libraries can be used to reduce the code size of simulation objects, improving performance. This will also aid migration performance, as the size of executable files will be reduced, requiring less transmission time over ethernet. However, shared libraries are currently unsupported in Condor. Implementation of the migration utility will involve thorough investigation into how shared libraries are supported on the Silicon Graphics machines. Each object will have its own private data for the shared libraries that is stored in the core image and must be properly retrieved.

An additional concern is the *interprocess communication* that occurs between simulation objects during execution. The Condor project does not support communicating processes. The Condor package has been enhanced at the University of Cincinnati in a specialized version of the apE graphics environment [1], allowing stages of the apE pipeline (each implemented as a Unix process) to not only reside on multiple machines but migrate freely among them under the user's control. Support for communication, namely Unix sockets, has been added. This work will be incorporated into the migration utility.

Secondary goals have included the use of C++ to design reusable and extensible classes that can perform checkpointing and migration. An object-oriented design will allow this work to be extended within the VHDL simulator or other projects. Special care has been taken to follow POSIX standards to enhance portability.

During the summer research period, the graphical user interface has been completed and is currently used to provide startup configuration data to the simulator. The interface identifies each machine on the network by an icon. The user may choose which machines that the simulator may use by selecting the proper icons. Ultimately, with the addition of process migration, the user will be able to modify this configuration dynamically during execution of the simulator. If a machine must be made available to others in the lab, the user will be able to click on the appropriate icon, causing all simulation objects to be removed from that machine and placed on other machines that are available.

The basic techniques and pitfalls of process migration have been investigated during the summer research period and a plan has been formed to complete the migration utility during the upcoming year. After a rudimentary migration facility is complete, attention will focus on the optimization of object relocation. Criteria such as current machine load, number of CPUs, CPU speed, and availability of other performance enhancements such as Scramnet are all valid parameters to be utilized to further optimize the simulator performance.

4.0 Conclusions

The addition of a process migration facility will greatly enhance the functionality of the QUEST VHDL simulator and allow full and unrestricted use of laboratory facilities to all its members. Checkpointing and migration capabilities can provide the basis for future development of the QUEST simulator. Periodic checkpointing can be used to provide *fault tolerance* capabilities. The ability to recover from faults such as power outages or hardware failure could save hours of computation. Migration provides the groundwork for future work in *dynamic load balancing*; migrating simulation objects from machines with heavy computation loads to machines with light or no computation loads during simulation execution. Because these underlying mechanisms rely on the Unix kernel, they are highly portable. The migration tool could be easily ported to support future versions of the QUEST VHDL simulator.

5.0 Bibliography

- [1] M. Ankola, "Implementation of Process Migration in apE," Master's Thesis, University of Cincinnati, 1992.
- [2] A. Brickner, M. Litzkow, M. Livney, "Condor Technical Summary," University of Wisconsin, October 1991.
- [3] D. Charley, T. McBrayer, D. Hensgen, P. A. Wilsey, M. Ankola, "Distributed Simulation on a Reconfigurable Network Using Non-Uniform Message Passing," *Proceedings of the Fifth ISMM Parallel and Distributed Computing Systems Conference*, 1992.
- [4] D. Charley, T. McBrayer, D. A. Hensgen, P. A. Wilsey, "High Speed Communication for Simulation of Large VHDL Models," *VHDL International Users Forum Fall 1992 Conference and Exhibition*, October 1992.
- [5] P. Chawla, H. W. Carter, P. A. Wilsey, "An Investigation of the Performance of a Distributed Functional Digital Simulator," *32nd Midwest Symposium on Circuits and Systems*, pp. 470-473, 1989.
- [6] M. Litzkow, M. Livny and M. W. Mutka, "Condor- Hunter of Idle Workstations," *Proceedings of the Eighth International Conference on Distributed Computing Systems*, 1988.
- [7] M. Litzkow and M. Livny, "Experience with the Condor Distributed Batch System," *Proceedings of the IEEE Workshop on Experimental Distributed Systems*, 1988.
- [8] M. Litzkow, "Remote Unix Turning Idle Workstations Into Cycle Servers," *Proceedings of the Summer 1987 USENIX Conference*, 1987.
- [9] M. Litzkow and M. Solomon, "Supporting Checkpointing and Process Migration Outside the Unix Kernel," *Proceedings of the 1992 Winter USENIX Conference*, 1992.
- [10] T. McBrayer, D. Charley, P. A. Wilsey, D. A. Hensgen, "A Parallel Optimistically Synchronized VHDL Simulator Executing on a Network of Workstations," *VHDL International Users Forum Fall 1992 Conference and Exhibition*, October 1992.
- [11] J. Myers, "Project Update: Design of an Airborne Graphics Generator," *VHDL International Users Forum Fall 1992 Conference and Exhibition*, October 1992.

THIS PAGE INTENTIONALLY LEFT BLANK

Preliminary Work on the Design of an Image Algebra Coprocessor

Trevor E. Meyer
Department of Electrical Engineering and Computer Engineering

Iowa State University
Ames, IA 50010

Final Report for:
Summer Research Program
Wright Laboratory

Sponsored by:
Air Force Office of Scientific Research
Bolling Air Force Base, Washington, D.C.

September 1992

PRELIMINARY WORK ON THE DESIGN OF AN IMAGE ALGEBRA COPROCESSOR

Trevor E. Meyer
Department of Electrical Engineering and Computer Engineering
Iowa State University

Abstract

This paper discusses preliminary work in the design of a coprocessor for performing image algebra generalized convolution operations. Based on a decomposition of the generalized convolution which simplifies the sliding window procedure into shifts of the source image, two processor designs have been proposed. One design is a systolic array which processes regions or even complete images in parallel. The other design is a serial pipeline which processes an image a pixel at a time. At least one of the two designs is in the process of being simulated.

This paper presents the current progress in the design of image algebra C++, the target language for this coprocessor. We then look at the two designs, comparing and contrasting, and discuss a few important details that have yet to be resolved. Finally, we examine some existing architectures and discuss features that might be useful in the design of this coprocessor.

PRELIMINARY WORK ON THE DESIGN OF AN IMAGE ALGEBRA COPROCESSOR

Trevor E. Meyer

1. Introduction

Image algebra is an intuitive language for writing image processing algorithms. Unfortunately, due to the large numbers of computations required by most of these algorithms, implementations of image algebra in software tend to be slow. In addition, unless the language in which the algebra is imbedded is a parallel language on a parallel processor, none of the natural parallelism present in image algebra can be exploited.

One solution to both these problems is to design a parallel coprocessor which performs image algebra operations. This paper presents some of the preliminary work towards this goal. We first discuss the C++ language and a new implementation of image algebra in that language. Then we examine several processor designs based on a decomposition of the generalized convolution, a key operation in many algorithms. Finally, we look at some similar systems that have been implemented and examine features of these systems that are relevant to this design.

2. The C++ Language

The C++ language is a popular example of what is known as an "object-oriented" language. Object-oriented refers to the ability to define and manipulate objects as if they were built-in data types. This greatly increases the utility of the language while simultaneously increasing code readability. Large programming projects involving multiple programmers, in particular, are well suited to the object-oriented paradigm. An object-oriented program tends to be easier to maintain and modify, even by future programmers who did not take part in the writing of the original code.

C++ is an object-oriented version of C, which means it is C with extensions that support objects. In fact, early versions of C++ were compiled by converting the object-oriented code into standard C using a preprocessor, then compiling the resulting C code. Thus, with a few minor exceptions, C is a subset of C++, and most C programs will compile in a C++ environment. This similarity has definite advantages, as it allows easy migration for C programmers to an object-oriented language.

The primary difference between C and C++ is the ability to define objects in C++ [9]. An object is similar to a C structure type. Like a structure, an object contains a collection of data fields, which may be any C type, or even other structures or objects. The important difference between structures and objects is that objects may also contain functions. These functions are referred to as "member functions" of the object, and are usually used to manipulate

the data fields within the structure. In most cases, one or more of the data fields are protected, and thus cannot be directly accessed by the user. The member functions provide safe ways to set and retrieve these protected data items. The advantages are clear. Objects can be written and collected into libraries, which are then accessed by different parts of a program. However, there is no need to worry about one misbehaved routine modifying the data in an illegal way, because the data is only accessible through the member functions. On large projects with multiple programmers, such data management is a necessity. By using well written and debugged objects, programmers do not have to be concerned that someone else's code is incorrectly using or modifying shared program resources. Policing of data is entirely taken care of by the objects themselves.

The other big difference between C and C++ is the ability of C++ to handle operator overloading. Operator overloading refers to the process of defining multiple meanings to a single operator, based on the types of its operands. Operator overloading is already built into the C language, in that, for example, the "+" operator can add both integers and floats. Unfortunately, C does not allow users to overload operators; the only overloading in C is that which is already built in. C++, on the other hand, supports overloading for user defined objects. Thus a user may define an object, then overload the "+" operator so that these new objects can be added with code of the form "object + object." In C, on the other hand, this addition of new data types must be performed by a function call. Although the same result is obtained in either case, the overloaded example is easier to understand while reading the code.

3. Image Algebra C++

A team at the University of Florida is working on an implementation of image algebra in C++. Image algebra is a machine-independent language designed as a tool for describing image processing algorithms. It has previously been imbedded in FORTRAN, Ada and C. These implementations, however, differ markedly from the C++ version. Image algebra FORTRAN, C and Ada were created by adding extensions to the language and then creating a preprocessor or compiler program that would recognize these extensions and could translate them into ordinary code, which is in turn compiled to machine code [5][10]. The problem with this approach is that it is not easily extended by users and it is not portable between systems. By the first point, we mean that users cannot modify or add to the language without rewriting the compiler or preprocessor program. This is a definite disadvantage for a new, dynamic language such as image algebra, which is continually undergoing changes and extensions. By the second point, we mean that an image algebra FORTRAN program, for example, can't be compiled on a machine that does not have the image algebra FORTRAN preprocessor installed. These languages are limited to certain types of machines to which the compilers or preprocessors have been ported, and that number is still fairly limited. Even the comparatively small image algebra FORTRAN preprocessor has proven difficult to port to systems other than the Sun for which it was written. Even changes in the operating system on the Sun itself have caused problems.

as we discovered when installing it at Eglin. Several sections of the code had to be modified before the program functioned properly.

The image algebra C++ will not suffer these deficiencies because it will not be implemented using a preprocessor or compiler program. Rather, the image algebra data types and operators will be defined using an object library. Each image algebra operand and its associated operators will be described by an object in the library. Thus, an image algebra C++ program is simply a standard C++ program which references the image algebra object library, and so may be compiled using any C++ compiler. For this reason, image algebra C++ will be portable to any system which supports C++, with little or no effort required beyond copying the object library onto the filesystem. Additionally, because the image algebra objects will be accessible in the library, they may be modified and extended as new developments occur in image algebra research. If a new type of operation is proposed, for example, it can be retrofitted to the objects in the library, perhaps not a trivial task but far simpler than rewriting a compiler.

Currently, image algebra C++ is in the very earliest stages of development. Work is being done to define the lowest level operands and the operations between them. Coordinate types, which may be eliminated in the future, are defined as numbers, and point types are defined as an ordered set of coordinates. Various operations between points, between coordinates, and between points and coordinates have been coded. More importantly for this project, an architecture for pointsets, and for images based on those pointsets, has been suggested. This architecture is similar to that developed for the image algebra Ada project, as described in [11]. Pointsets and images will be discussed further in a later section of this paper.

4. Image Representation in Image Algebra C++

One of the most important decisions to be made when embedding image algebra in a computer language is the way in which images will be represented. This choice is crucial to the performance of the language because images are the basic data types upon which the language will operate. One would like an image representation that is simple for the user of the language to manipulate, and that is simple for the computer to manipulate internally. On the other hand, one wants the image representation to be as flexible as possible, so that many different types of images can be represented. For example, an image may be defined over a domain with dimension greater than two, and the domain might not be rectangular. Similarly, image values could be vectors instead of simple numbers. Obviously, there will be trade-offs. The more general the image representation, the more complex the internal representation becomes.

As a compromise, an image representation can be formulated that varies in complexity. It will be optimized for simple, rectangular images, but at the same time be flexible enough to represent, for example, a vector-valued image defined over the volume of a sphere. Such a representation was developed for the image algebra Ada and is discussed in [11]. Although the image algebra C++ image representation scheme had not yet been formalized when this paper was written, personal communication with the members of the project indicates that it will be similar to that described in [11].

In image algebra, images are defined relative to a pointset and a value set. The pointset is a subset of euclidean space over which the image is defined. The value set is the set of all values in the image. The image itself, then, is a function that maps a point in the pointset, known as the "pixel location," to a value in the value set, known as the "pixel value." Thus to understand the proposed representation of images in image algebra C++, we must understand the representation of points, values and pointsets.

A point is defined in image algebra C++ to consist of an ordered set of coordinates. These coordinates are each of type coordinate, which in turn is simply an integer. Coordinate types may or may not be used in the future. If they are eliminated, they will be replaced by simple integer types. There are also plans to allow real-valued coordinates, though this work was not formalized before this paper was written.

A pointset will be defined using a special data structure described below. The goal of the pointset representation is to define the domain of an image in scanline order, and in a way that allows the most compact representation for rectangular images with two-dimensional domains. The data structure is straightforward, and is diagrammed in Figures 1 through 4.

The pointset will consist of a set of intervals. Each interval represents a slice across the image. These slices will thus be of dimension $N-1$, where N is the dimension of the image. Each $N-1$ -dimensional slice is then itself divided into slices of dimension $N-2$, and so on. The process continues recursively until no more slices may be taken.

To illustrate the procedure, consider the two-dimensional domain shown in Figure 1. The domain consists of two regions, a 6-sided polygon with all angles equal to 90 degrees, and a small rectangle. Note that while these regions are disconnected, they represent a single pointset and thus could serve as the domain of a single image.

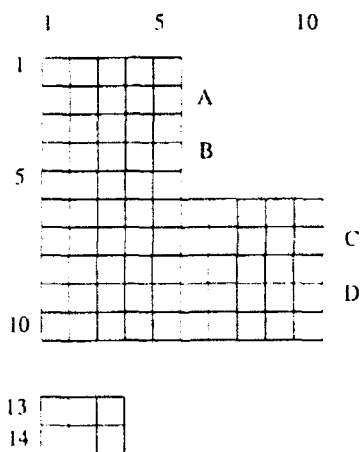


Figure 1. Example image domain.

In Figure 2, the above domain has been divided into one-dimensional slices at integer intervals. Integer intervals will suffice to fully describe the pointset as long as it is defined over integer coordinate values. If real-

valued coordinates are used, the process is significantly complicated. As can be seen in the figure, a total of 12 slices are generated, one slice at each integer coordinate in the vertical direction except 11 and 12, which are outside the domain. Note that no slices are needed at coordinates beyond the extent of the shaded regions.

Thus the region shown in Figure 1 is fully described by the 14 slices shown in Figure 2. Further slicing of these slices will generate zero-dimensional points, and is thus not necessary. If the original region had been a three-dimensional volume, each slice would give a two-dimensional plane. Each of these planes would then be sliced as in Figure 2, giving sets of one-dimensional lines. This collection of lines would fully describe the original volume.

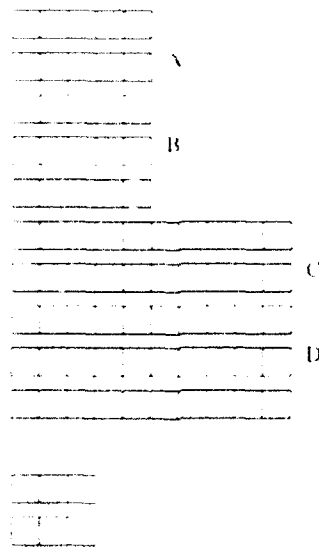


Figure 2. Image domain divided into slices.

Examining the slices shown in Figure 2, one sees that many of the slices cover the same range of horizontal coordinates. Slices A and B, for example, both traverse the figure from horizontal coordinate 1 to coordinate 5. The slices are, in fact, identical save for their vertical coordinate values, 2 for A and 4 for B. Slices C and D are related similarly. This repetition offers an opportunity to compress the representation of some image domains, particularly "boxy" domains such as that shown in Figure 1. Consider Figure 3. In this representation, the duplication is removed by storing each unique slice only once. Pointing to each slice is a data record which contains the range of vertical coordinates for that slice. Thus the slices in the upper portion of the large block are replaced by a single slice five units long accompanied by a data record indicating that slice occurs in the vertical range of 1 to 5. Similarly, the lower part of the large block is replaced by a slice ten units long and a data record with the vertical range of 5 to 10. Notice also that, while this representation can store any domain of arbitrary size, shape and dimension, it will be most compact for rectangular, two-dimensional domains, which are the most commonly used.

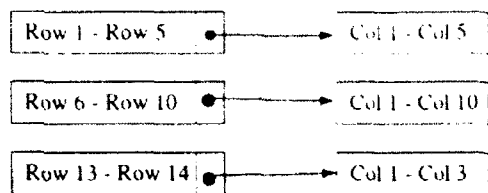


Figure 3. More compact representation of image domain in Figure 1.

This slice representation stores information about the image domain, but not about its range. In fact, the values of an image may be stored using a simple vector representation, regardless of the shape or dimension of the domain. Notice that the slices are stored in scanline order, taken from top to bottom and left to right. As we may assume this always to be the case, image values corresponding to each slice are stored in the same order, truncated together in a single long vector. In order to find a value associated with a given pixel location, one simply moves through the vector slice by slice until the value is located. One large advantage to this representation is that for flat rectangular images, the values are stored in the same way as two-dimensional data arrays, and thus may be manipulated as such. Figure 4 shows a simple image and its domain representation and value vector.

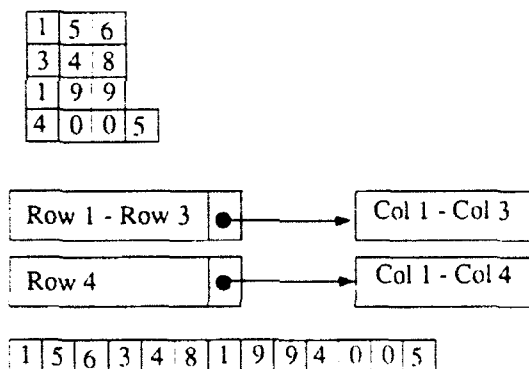


Figure 4. An image and its representation in iac++.

5. An Image Algebra Coprocessor

A failing of any software implementation of something as computationally expensive as many image processing algorithms is slow speed. The delays can become prohibitive for complex procedures on large image arrays. Unfortunately, many real world algorithms are very complex and many real world data sets contain huge images. What are some solutions to this problem? Assuming the code is already optimized as much as possible, the obvious answer is to get a faster computer. Improving the speed of a large, general purpose computer is often very expensive, and in many cases infeasible.

A more practical solution that is being investigated by a number of researchers is the development of special purpose image processing hardware. Usually this hardware takes the form of a coprocessor that runs under the control of a general purpose host computer. The host computer executes the body of the image processing program, but actual computation intensive steps in the algorithm are sent to the coprocessor, which is optimized to perform these operations. When the coprocessor has completed its task, it signals the main computer and sends the result over a shared bus or i/o channel. In this section we will look at two proposed architectures for such a coprocessor, optimized specifically for image algebra operations. These architectures are based on a decomposition of the generalized convolution studied by Dr. Patrick Coffield at Eglin AFB. In the next section, we'll look at some other architectures in use elsewhere which were examined and compared this summer to Dr. Coffield's architectures.

Dr. Coffield noticed that the generalized convolutions operation may be expressed as a series of shift-accumulate steps. For each element in the template, the image is shifted, combined with the corresponding template value on a pixelwise basis, and then incorporated into an accumulator. The term "incorporated" is used because the operation for combining images is not necessarily addition, but depends on the properties of the convolution operator. The derivation is explained in detail in [2].

The advantage of this decomposition, though it appears at first only to complicate matters, is that it reduces the moving mask convolution type operation, so common in image processing, to a few shifts and accumulates. In most cases the number of shifts will be small, as one shift is required for each location in the template, and most common templates are 3x3 or smaller. The accumulates, on the other hand, could be performed by a simple ALU, as can the combination of the shifted image with the template value. Thus the algorithm suggests a straightforward implementation, shown in Figure 5 and discussed in more detail in [2].

The processor in Figure 5 is a pixel-slice processor. That is, it is capable of processing a single pixel value in the image. A physical implementation would probably consist of a grid of these pixel processors etched on a single VLSI chip. Ideally, we would like the size of the grid to equal or exceed the largest image size we planned to process. In reality, the dimensions would probably be limited by current VLSI technology. A smaller processor would process a "window" of the image, with circuitry added to handle edge effects. Even with a relatively small window of, for example, 100x100, large processing speedups could be realized in some cases.

The main problem with this type of architecture is in fact a general problem that arises in any massively parallel system, and that is the i/o network. How do we design circuitry that allows us to not only feed a square region of an image into the inputs of this processor array, but also shift the image in four different directions? If there is a simple solution, it currently eludes us. The most promising route is that taken by the Cytocomputer to shift a three-pixel-wide strip of an image past a 3x3 mask [6]. The image values are stored in a long shift register, part of which "lies beneath" the mask. Using dead space in the shift register, a continuous linear shift of image values though the register maintains the correct 3x3 region of the image beneath the mask. Thus, in affect, the mask is slid across the image from left to right. When it reaches the right side, it moves down a row and starts back at the

left. Figure 8, in section 6.4, illustrates this process on the Cytocomputer. The Cytocomputer is discussed in detail in [6], and will be covered along with other architectures for image processing in the next section.

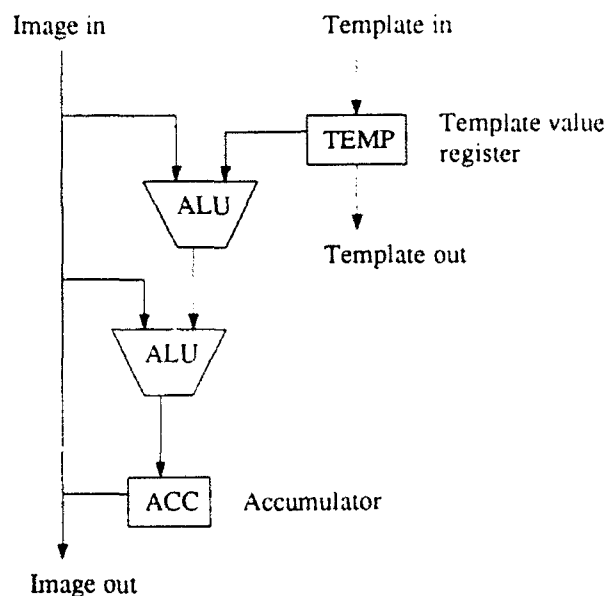


Figure 5. Image algebra pixel slice processor.

Scientists at Oak Ridge, working to implement Dr. Coffield's decomposition in hardware, have proposed a different approach. Their system, shown in block diagram in Figure 6, is based on a image pipeline structure. Rather than processing individual pixels in parallel, this architecture will read in two images in a serial pixel stream. In this way, the parallel i/o bottleneck effectively vanishes. The two image streams are directed into an ALU which performs pixelwise mathematical operations. The resultant stream then enters a FIFO buffer with built-in delay sections that function in a manner analogous to the shift register discussed above and in [6]. From the FIFO, data is fed into arithmetic units which perform the template operation and accumulate the result. In addition, this design provides optional hardware for handling border effects, for sorting local neighborhoods, and for histogram generation. One suggested physical realization of this circuit is through the use of standard digital signal processing chips. This approach offers the advantages of simple prototyping and a flexible final design against the shortcoming of slower overall image throughput.

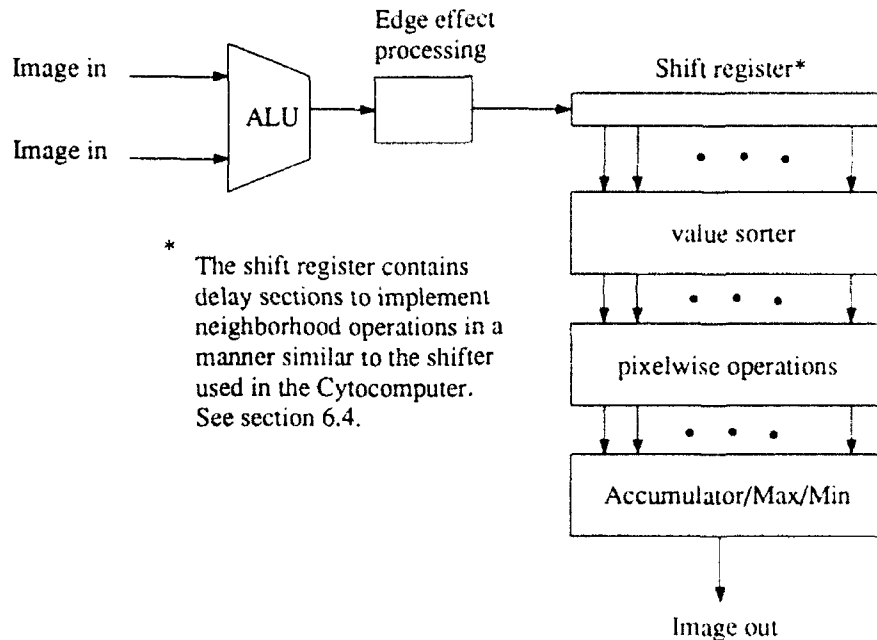


Figure 6. Oak Ridge implementation of image algebra processor.

The two approaches detailed above are at practically opposite ends of the spectrum, surprising since they describe the same equations. As of the end of the summer, which approach to take remained a subject of debate. Lacking experience in this area, we looked at several related architectures that are built and functioning. The more interesting of these designs will be detailed in the next section.

6. Related Architectures

During the summer, several existing image processing architectures were studied. These architectures either performed operations similar to those of the image algebra generalized convolution, or involved hardware similar to the hardware from Oak Ridge or the hardware detailed in Dr. Coffield's paper. Among those architectures worthy of note are the Connection Machine, MasPar MP1, the Cytocomputer and the Parallel Recirculation Pipeline (PREP). Other papers studied did not lead directly to insight on this problem and thus are not included here.

6.1 The Connection Machine

The best known of the four image processing architectures to be discussed is the Connection Machine, from Thinking Machines. Actually, the Connection Machine is not designed specifically for image processing, as are the last two machines discussed in this section. Rather, this computer is a general purpose Single-Instruction-Multiple-Data (SIMD) fine-grained parallel processor [1 pp. 335-343]. Single-Instruction-Multiple-Data means that all

processors in the machine execute identical instructions on possibly differing local copies of data. Fine-grained parallel means that the machine consists of a large number (65,536) of small, relatively weak processors. Thus, the power of the machine comes not from the power of the individual processors, but rather from the combined power of the whole. In a SIMD computer, the goal is to split a problem into many tiny, identical pieces that can be spread over all the processors. Since many image processing operations consist of individual pixel operations or small neighborhood operations replicated at every pixel in a large image, SIMD machines are ideally suited for image processing tasks.

The Connection Machine resembles the architecture in Dr. Coffield's paper in that both are fine-grained massively parallel SIMD computers. Because the Connection Machine is a general purpose computer, many features of that machine are not needed in an architecture optimized for generalized convolution. One feature notably lacking in Dr. Coffield's paper is an interconnection network between processors. This network, the design of which is so important in general purpose parallel processors [8 pp. 325-353], is not needed because of the way Dr. Coffield has decomposed the convolution operator. Rather than use a communication network between processors to communicate values needed in neighborhood operations, on which the convolution operation is based, Dr. Coffield's circuit will shift the image in a separate section of memory. Thus the communication network normally associated with the processor array exists instead implicitly in the memory architecture. Whether such a design, especially for large images, is feasible or not requires further study. The Oak Ridge team took a different approach, possibly to avoid this problem. Another modification in the design might lead to a SIMD architecture, like Dr. Coffield's suggestion, with the addition of an interprocessor communication network instead of a shiftable memory array. Unfortunately, the network topology used in the Connection Machine, a 12-dimensional hypercube, does not lend itself well to most image processing problems. The MasPar MP1, discussed below, uses a network better suited to this problem, and thus proves more worthwhile for further investigation.

6.2 The MasPar MP1

Another massively parallel computer, similar to the Connection Machine, is the MasPar MP1 from MasPar Computer. Like the Connection Machine, the MP1 consists of a large array of small processors, operating in a fine-grained, SIMD fashion [3]. The chief difference between the machines is in the interprocessor connection scheme.

As mentioned above, the processors in the Connection Machine are connected by a 12-dimensional hypercube. While this topology is, overall, one of the most flexible, it is not well suited for tasks laid out in two-dimensional rectangular domains, such as almost all image processing operations. The MP1 processors, on the other hand, are laid out in a two-dimensional array with 8 nearest neighbor connections. In addition, hardware is provided to simplify more complex routes through the network. This architecture, then, seems ideally suited to image processing.

In fact, an investigation into the implementation of image algebra on this computer has been made by the author [7]. The MP1 is in fact ideal for performing pixelwise operations between images, and for performing

convolution-type operations with 3x3 or smaller neighborhoods. Larger neighborhoods are easily implemented on the machine but may incur a performance penalty, as routing hardware is shared by multiple processors. This architecture, then, is worth further investigation if a redesign of Dr. Coffield's initial circuit is considered. A network similar to that in the MasPar could solve all the problems of shifting image data across processors. In fact, one could simply add a level to the design, a "dummy-processor" whose only purpose would be to store data and to send and receive messages on the network. This upper array of processors would be controlled by a single, master processor whose job is to load data from memory into the array. Of course, a design like this begins to leave the domain of small, one or two VLSI-chip implementations and enters the arena of desktop or larger hardware. Whether this is a worthwhile trade-off needs to be considered.

6.3 The Parallel Recirculation Pipeline (PREP)

The PREP architecture is similar to the architecture suggested in Dr. Coffield's paper. The PREP consists of a rectangular array of cells, called PREP cells, attached to a host processor [4]. The cells are connected in columns, and each column is independent of the others, though they may pass values. A PREP machine may consist of an arbitrary number of columns, the columns containing an arbitrary number of cells. To extend a given machine, one can either add cells to the column or add columns. The host processor then divides a given image operation between the available columns and the cells in those columns.

The similarity between the two architectures is evident when one examines the internal block diagram of a PREP cell, as shown in figure 7. The cell is essentially a general purpose arithmetic pipeline, the image pixels entering at the top and the resultant pixels exiting at the bottom. Like Dr. Coffield's architecture, various ALUs are available to perform operations on the pixel data as it moves through the cell. Unlike Dr. Coffield's architecture, which is optimized to perform a specific type of neighborhood operation, the PREP cells are general purpose, with several different data paths. Both architectures are extensible, however, with the processing time for an image depending mainly on the number of available processors. One item to note, however, is that the PREP is designed to handle only 2x2 neighborhood operations efficiently. To do larger neighborhoods, one must first decompose them into 2x2s. The image algebra coprocessor, on the other hand, is designed to perform equally well with any neighborhood size up to 7x7 or more.

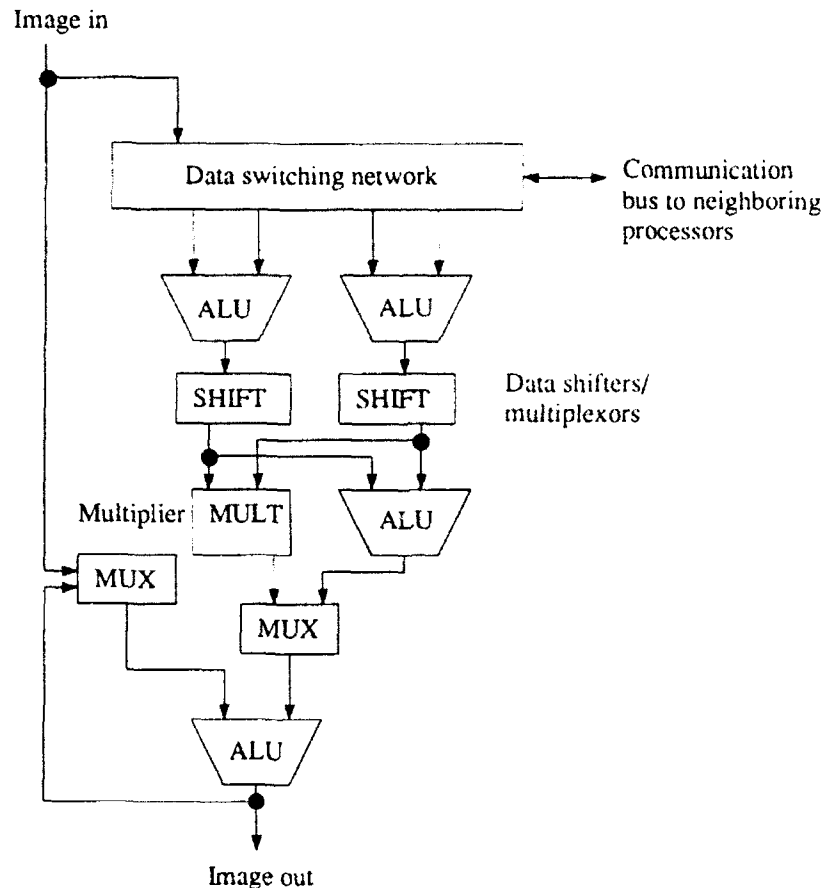


Figure 7. Simplified diagram of a PREP cell.

6.4 The Cytocomputer

The Cytocomputer is a single pipeline of cells bearing some resemblance to one column of the PREP [6]. Each cell is designed to perform a 3x3 neighborhood operation, and an arbitrary number of cells can be added to the pipeline. Additional cells speed up more complex operations by allowing further processing without having to reload the image. See Figure 8. In some ways the Cytocomputer resembles the processor proposed by the team at Oak Ridge, in that both process images as a serial streams of pixels, and both implement neighborhood operations by shifting image data.

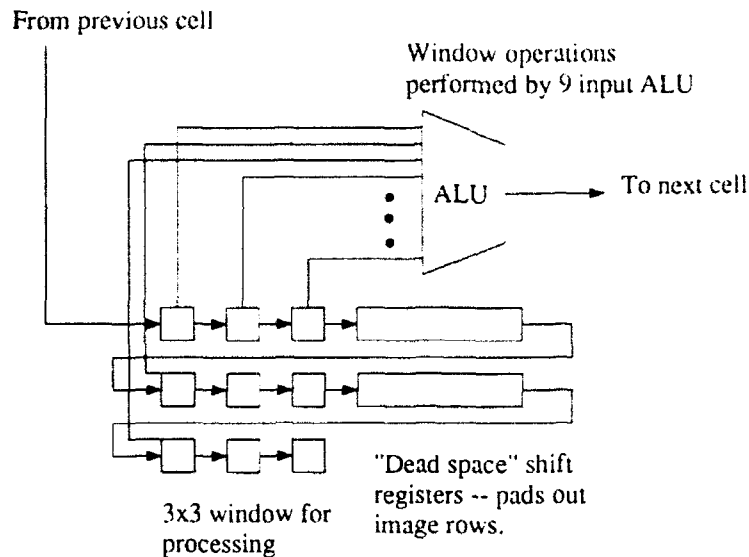


Figure 8. One cell in the Cytocomputer.

The shifter in the Cytocomputer is worth looking at in more detail. As can be seen in the figure, it is a serial shift register, and thus shifts the image in one dimension only. In order to process a two-dimensional neighborhood of an image, the register must be long enough so that several rows (three in this case) are in the register at one time. Then, by varying the length of the register so that it is just over two image rows long, a correct 3x3 neighborhood will be maintained in the active register cells as the image is shifted through. Active cells refers to the cells which are connected to the ALU. The other cells in the register can be thought of as dead space designed to pad out the image data in the register. This register design is similar to the memory design suggested for the image algebra processor developed by Dr. Coffield. In that case, however, the active cells would refer to those portions of the image currently feeding a processor, while the dead space would hold portions of the edges that extend past the active region of the processor array. Of course, a large image will require a very large register. If the idea is developed further, in depth investigation of the Cytocomputer shift register is a good place to start.

7. Interfacing Image Algebra C++ to the Coprocessor

In designing a coprocessor, an important consideration is the ease of integrating the coprocessor with the software that is destined to use it. In this case, we want to look at the modifications to image algebra C++ that are required in order for image algebra programs to access the coprocessor. Ideally, we would like the interface to be as unobtrusive as possible: in the best scenario, the user could compile one piece of code and then use it with or without the coprocessor, without modification.

The image algebra C++ differs significantly from previous image algebra implementations in that it is not implemented with a compiler or preprocessor. Rather, image algebra data types and operators are coded into

objects and bundled into an object library. This makes the task of interfacing much easier. Rather than code run-time tests that search for the coprocessor, and then include hooks in the code to send operations to the coprocessor rather than the main processor, we can simply create two object libraries. One library will include objects whose operators are coded to work with the coprocessor. The other library will contain code that runs on the main processor. The user then includes the appropriate library when compiling. The source code written by the user will not have to be changed in any way. The one disadvantage of this approach is that a program will need to be recompiled if the availability of the coprocessor changes. This situation can be avoided by coding objects that check for the coprocessor and then execute the corresponding code. The extra work involved, however, may well be more trouble than it is worth.

It was decided during meetings with Dr. Coffield that the first coprocessor design will only handle operations on two-dimensional rectangular images. This assumption simplifies the hardware interface and the data conversion between image algebra format and coprocessor format. If images are implemented in image algebra C++ as planned, the image values will be stored in memory in raster scanline order. The dimensional information will consist of two numbers, the number of rows and the number of columns. Note that this simplification of the pointset data structure down to two numbers is due to the fact that we are only processing rectangular images. In the case of a rectangular image, the first item in the pointset structure is simply the number of rows, since all rows are the same length. This entry then points to the column entry which contains the number of columns. See Figure 9.

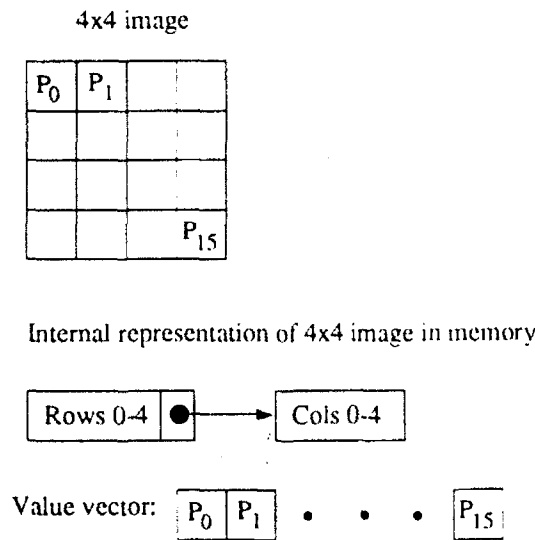


Figure 9. A rectangular image and its representation in image algebra C++.

Once the coprocessor is loaded with the number of rows and columns, it can read the image values from memory in scanline order, possibly using a DMA controller to speed memory access. No specifics can be decided

upon until the processor board is designed, however the interface need not be complex. Command codes will be needed to inform the coprocessor of the type of operation, when to read row and column numbers, and when to read the data. Similar codes will be returned by the coprocessor indicating the completion of the operation. Further details must wait until the hardware and software sides of the project are complete.

8. Conclusion

Though important preliminary work has been completed, much remains to be done. No concrete implementation of a coprocessor can be realized until the image algebra C++ is complete. In the meantime, further design of components, such as the memory and i/o system, and extensive simulation of the proposed designs should provide a better measure of their feasibility. Much can be learned from existing systems, such as those few described here, by looking at their features and shortcomings. An image algebra parallel coprocessor is a difficult but worthy goal; with one could at last realize both the elegance of the image algebra and the speed of parallel processing.

Bibliography

- [1] George S. Almasi and Allan Gottlieb. *Highly Parallel Computing*. Benjamin/Cummings, Redwood City, CA, 1989.
- [2] P. C. Coffield. Architectures for processing image algebra operations. In *SPIE proceedings, Image Algebra and Morphological Image Processing III*, volume 1769, pages 178-189, San Diego, CA, July 1992.
- [3] MasPar Computer Corporation. MasPar MP-1 standard programming manuals part number 9300-9001-00. Technical report, MasPar Computer Corporation, July 1990.
- [4] Honeywell. A programmer's guide for the parallel recirculating pipeline (PREP). Technical report, Honeywell.
- [5] IVS Inc. *Image Algebra FORTRAN Version 3.0 Language Description and Implementation notes*. IVS Inc., Gainesville, FL., 1989.
- [6] Robert M. Loughheed and David L. McCubbrey. The cytocomputer: A practical pipelined image processor. In *IEEE Proceedings of 7th Annual International Symposium on Computer Architecture*, 1980.
- [7] T. Meyer and J. L. Davidson. Image algebra preprocessor for the MasPar parallel computer. In *SPIE proceedings, Image Algebra and Morphological Image Processing II*, volume 1568, pages 125-136, San Diego, CA, July 1991.
- [8] Harold S. Stone. *High-Performance Computer Architecture*. Addison-Wesley, Reading, MA, 1990.
- [9] Bjarne Stroustrup. *The C++ Programming Language*. Addison-Wesley, Reading, MA, 1991.
- [10] J. N. Wilson. Introduction to image algebra Ada. In *SPIE proceedings, Image Algebra and Morphological Image Processing II*, volume 1568, pages 101-112, San Diego, CA, July 1991.

- [11] M. F. Yoder. Images and image domains in Image Algebra Ada. In *SPIE proceedings, Image Algebra and Morphological Image Processing*, volume 1350, pages 274–282, San Diego, CA, July 1990.

THIS PAGE INTENTIONALLY LEFT BLANK

**DETECTION AND ADAPTIVE FREQUENCY ESTIMATION
FOR DIGITAL MICROWAVE RECEIVERS**

Faculty Associate : Arnab K. Shaw
Assistant Professor
Electrical Engineering Department
Wright State University
Dayton, OHIO-45435

Graduate Associate : Steven Nunes
Electrical Engineering Department
California State University, Chico

Final Report for :
Summer Research Program
Wright Laboratory, Avionics Directorate

Sponsored By :
Air Force Office of Scientific Research
Bolling Air Force Base, Washington, D. C.

October 1992

DETECTION AND ADAPTIVE FREQUENCY ESTIMATION FOR DIGITAL MICROWAVE RECEIVERS

Arnab K. Shaw

Assistant Professor

Electrical Engineering Department

Wright State University

Steven Nunes

Graduate Student

Electrical Engineering Department

California State University, Chico

Abstract

Detection of presence of targets and estimation of Angles-of-Arrival (AOA) are two of the important tasks of a digital EW receiver. In this report, the time-domain detection problem and an adaptive frequency/AOA estimation scheme have been studied. For the detection problem, detection thresholds have been derived for square-law detectors in cases of single and multiple observation samples. The adaptive frequency/AOA estimation scheme has been studied to analyze its behavior at various noise and input conditions. Simulation results indicate that the use of higher estimation model order than the actual order improves the bias and variance of the estimates at the cost of longer convergence time.

**EFFICIENT ANALYSIS OF PASSIVE
MICROSTRIP ELEMENTS FOR MMIC'S**

Krishna Naishadham
Assistant Professor

and

Todd W. Nuteson
Graduate Associate

Department of Electrical Engineering
Wright State University
Dayton, OH 45435

Final Report for:
AFOSR Summer Research Program
Wright Laboratory, WPAFB

Sponsored by:
Air Force Office of Scientific Research
Bolling Air Force Base, Washington, D.C.

September 1992

EFFICIENT ANALYSIS OF PASSIVE MICROSTRIP ELEMENTS FOR MMIC'S

Krishna Naishadham, Assistant Professor
Todd W. Nuteson, Graduate Associate
Department of Electrical Engineering
Wright State University

ABSTRACT

Passive microstrip elements such as meander lines are analyzed by the full-wave, space domain moment method at microwave and millimeter wavelengths. Redundant calculations in the moment matrix are eliminated by utilizing various symmetries. At lower microwave frequencies quasi-static approximations of the Green's functions are invoked to simplify the analysis. The Green's functions are in general Sommerfeld-type integrals, which are computationally intensive. In this paper, closed-form analytical approximations of the integrals, recently developed by Chow et al., Aksun and Mittra, are utilized to increase the efficiency of the algorithm such that a circuit of moderate electrical size can be analyzed in reasonable time. Sample computed results are presented for the scattering (S) parameters of meander lines and multi-turn spiral inductors. Computed results compare reasonably well with measurements on a vector network analyzer.

BUILT-IN SELF-TEST DESIGN OF PIXEL CHIP

R. Frank O'Bleness
Graduate Student
Department of Electrical Engineering
Wright State University
Dayton, OH 45435

Final Report for:
Summer Research Program
Wright Laboratory

Sponsored by:
Air Force Office of Scientific Research
Bolling Air Force Base, Washington, D.C.

September 1992

BUILT-IN SELF-TEST DESIGN OF PIXEL CHIP

R. Frank O'Bleness

1. INTRODUCTION

Presented within is a progress report summarizing work completed (and subsequent conclusions drawn) over the twelve weeks of our research period at Wright Laboratories, Electronics Directorate, Wright Patterson Air Force Base, Dayton, Ohio.

A brief preliminary evaluation of the COMPASS design tools is presented in section 2. Section 3 is devoted to an evaluation of the ALU and Linear Tree elements of the PIXEL graphics chip and their suitability to Built-In Self-Test techniques. Section 4 discusses an initial plan to implement BIST to the ALU and Linear Tree elements. Section 5 discusses the updated BIST plan. Section 6 offers a brief conclusion and presents an overview of the Circular BIST study.

2. COMPASS DESIGN TOOLS

COMPASS offers a wide variety of powerful design tools for both Standard Cell and Gate Array design methodologies. As our focus was intentionally limited to Standard Cell design, all discussion of design tools is made within the context of the Standard Cell design methodology.

The Logic Assistant is the main tool for Standard Cell design and layout, but the Terminal Shell is equally important in regards to simulation and timing. Within the Logic Assistant are a number of powerful tools including the Design Assistant, the Timing Verifier, the HDL Assistant, and the Test Assistant. Exploration into COMPASS' design tools was made primarily in the Test Assistant, but enroute many other aspects of the Logic Assistant were utilized.

COMPASS is complete with an extensive set of Standard Cell libraries, including both technology specific and portable technology libraries. In the Logic Assistant, through the use of the Portable and Synthetic libraries, we completed a datapath design, and with the use of the State Machine library we completed a state machine design from a state diagram representation. Both design examples were constructed within the graphical design environment, although the latter could have been completed in the Shell.

Also in the Logic Assistant, we built an 8-bit Adder/Subtractor. Shown in figure 1, this is a 93 gate combinational circuit that is used as a test bench for combinational fault simulation.

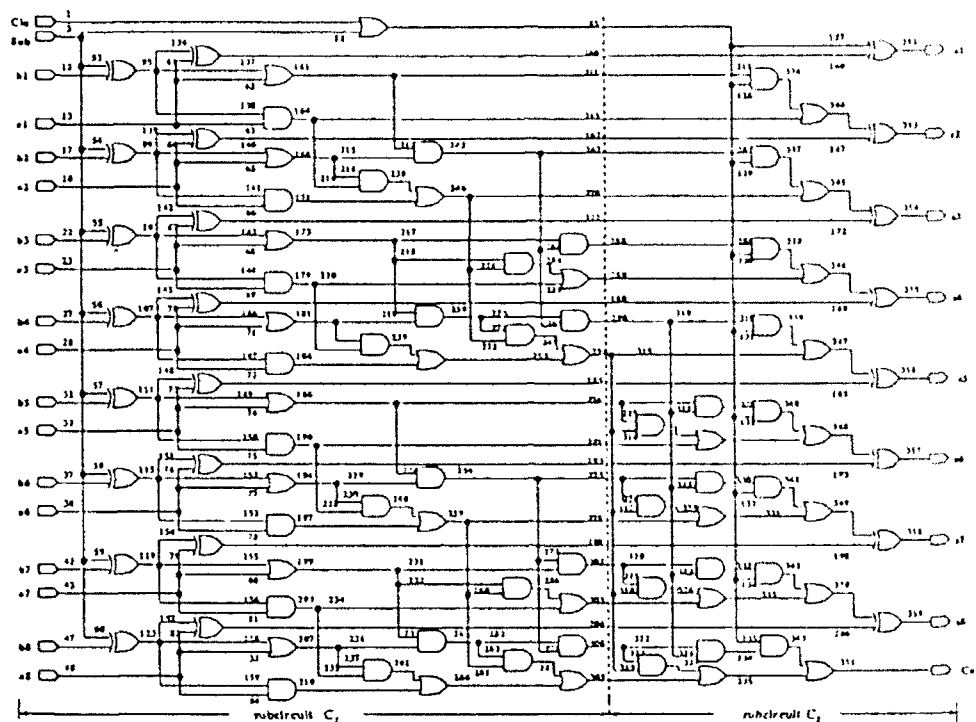


Figure 1. 8-bit Adder/Subtractor

In the Shell, using qsim, we attempted to run COMPASS' PROFAULT fault simulator. This was aborted as PROFAULT was not under a licensing agreement. Upon advice from the COMPASS Hotline, we attempted to run another fault simulator, MACH1000, from the Nellist Utilities in the Shell and met with a similar response.

COMPASS' Test Assistant offers four main test methodologies which may be used independently or in conjunction with one another. These Methodologies are: Boundary Scan Test (BST), MUX Isolation, Threshold Detection, and Full Scan Design. Within COMPASS, BST and Threshold Detection are mutually exclusive.

COMPASS Boundary Scan Test is compatible with IEEE Std. 1149.1. The BST13E000d library contains a large number of BST register cells and controllers. The controllers are dependent on which BST instructions are supported. Mandatory functions are EXTEST, SAMPLE/PRELOAD, and BYPASS. Optional instructions include INTEST, IDCODE, RUNBIST, and SCANM.

Boundary Scan Test was incorporated into the 8-bit adder/subtractor above. Upon examination of the schematic created from the Test Assistant, the functionality was determined to follow the IEEE Std. 1149.1. The only optional instruction supported was the INTEST. Examining the controller, a few observations were made. The Test Access Port (TAP) controller was generated using a standard 16-state state machine shown in figure 2.

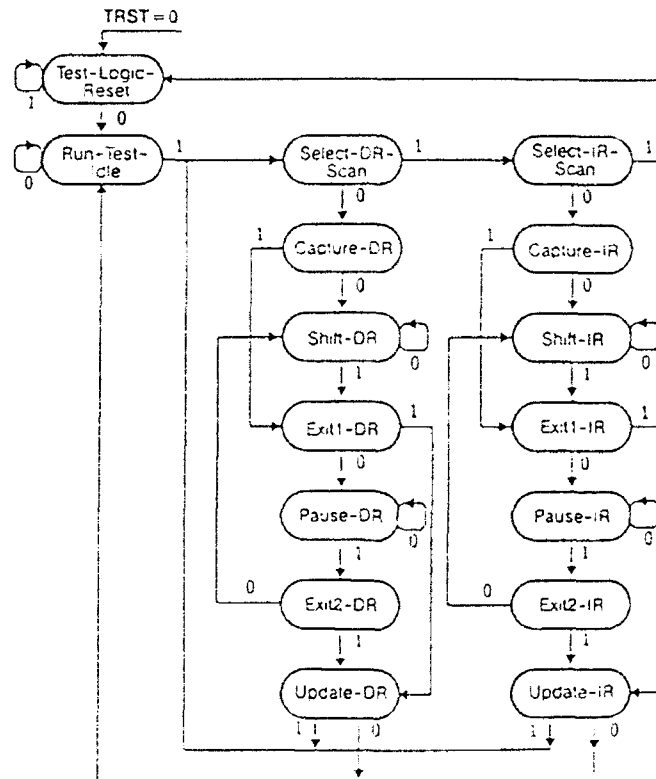


Figure 2. TAP controller state diagram

It is interesting to note that here the state machine was implemented in a conventional manner (combinational gates and Flip-Flops) rather than the PLA implementation used in the state machine compiler. COM-PASS' use of the Single Transport Chain (STC) Concept to reduce silicon overhead is evidenced in the sharing of the instruction and bypass registers. The instruction decoder decoded a 3-bit instruction into five control signals that control the value of the BS register, the update stage of the BS register, the BS test mode, and the EXTEST/INTEST mode.

Test Assistant's MUX Isolation test methodology was also independently implemented into the 8-bit adder/subtractor. The MUX partitioning line is shown dotted in figure 1. Due to confusion in specifying an

instance list, the partition did not occur exactly on the line specified. However, the correct number of multiplexers were inserted and test control logic was generated. Upon preliminary examination of the schematic an interesting observation surfaced. It appears that only ONE MUX can be accessed at any given time. This would seem to prohibit the use of COMPASS' MUX Isolation for any partitioning scheme. Also, it is required that additional test pads be added to a combinational circuit for additional test control, including pads for a test clock, test reset, and test enable. Examination of Test Assistant's Threshold Detection circuit was interrupted due to a software error and Full Scan Design was disabled. Further evaluation is in progress.

The Test Assistant has the capability to add Built-In Self-Test (BIST) architecture to RAMs, ROMs, and multipliers via its BIST compilers. Furthermore, Linear Feedback Shift Registers (LFSR) may be generated automatically from Test Assistant's LFSR Compiler. However, Test Assistant can NOT automatically add ANY custom BIST circuitry to any component. All BIST MUST be done manually. BIST can be incorporated into Boundary Scan Test by supporting the RUNBIST instruction.

3. SUITABILITY OF PIXEL ELEMENTS FOR BUILT-IN SELF-TEST

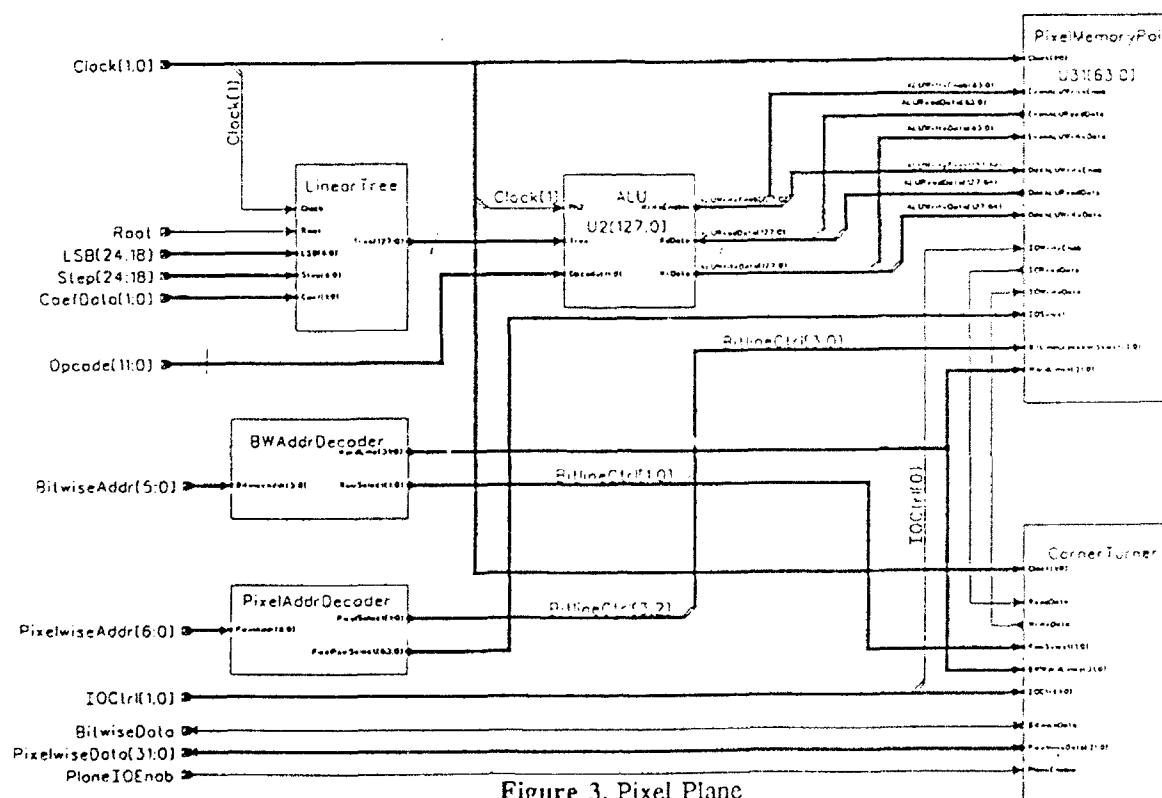


Figure 3. Pixel Plane

Figure 3 shows the PIXEL Plane. As can be seen from figure 3, the PIXEL chip is composed of a number of large blocks each of which can be evaluated to determine its suitability to Built-In Self-Test techniques. Our initial focus was on the 128-bit ALU block and the Linear Tree structure.

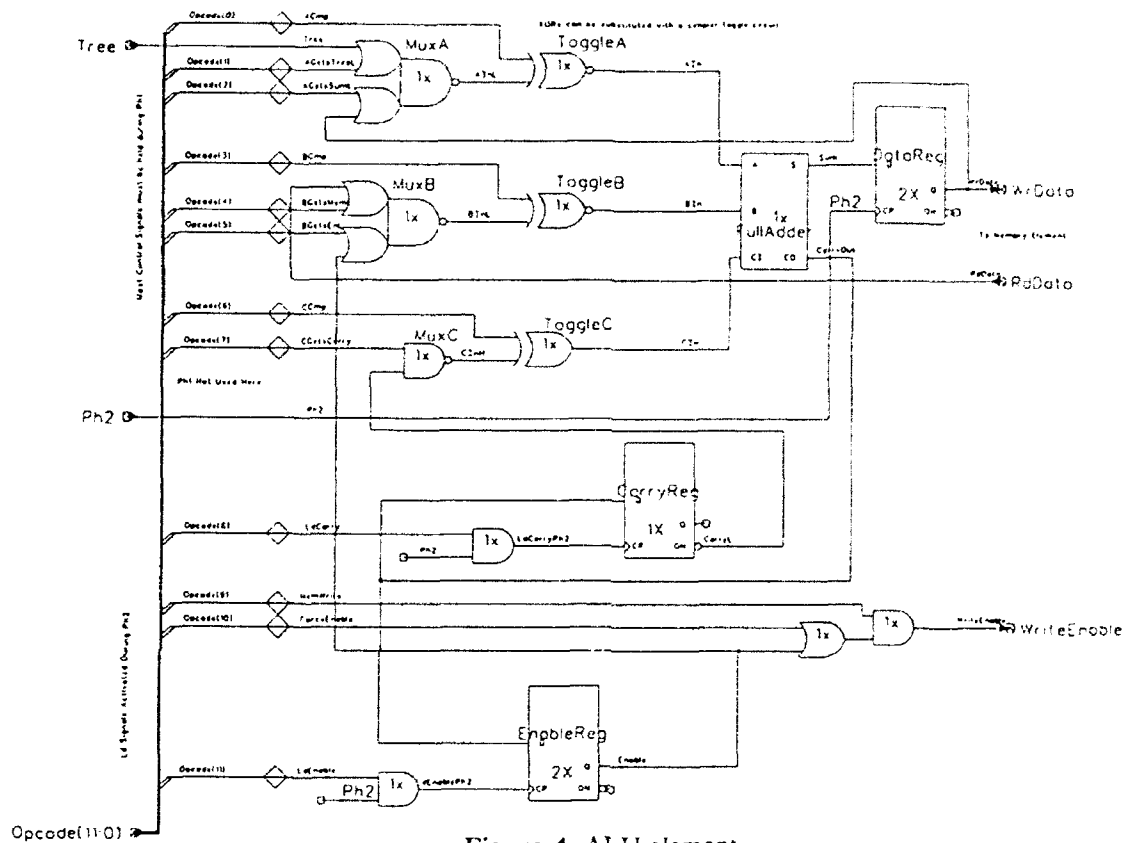


Figure 4. ALU element

A bit-slice of the 128-bit ALU block is shown in figure 4. The ALU element is a sequential circuit employing 14 simple gates, a full adder, and three D-type flipflops. The Linear Tree block is divided as shown in figure 5, with the Tree3Stage shown in figure 6. The basic block of the Linear Tree, shown in figure 7, is also a sequential circuit containing three simple gates, a full adder, and three D-Type flipflops. For analytical purposes, the full adder of figure 8 is assumed for both the full adder of the ALU and for the full adder of the single tree stage.

The first step in our analysis of the elements of the PIXEL chip was to determine if the elements could achieve high fault coverage in a simple combinational model. In order to construct a combinational model of a sequential circuit, the feedback paths must be broken and replaced with extra inputs. A combinational model of the ALU element is shown in figure 9. Note that the feedback paths have been cut and replaced by extra

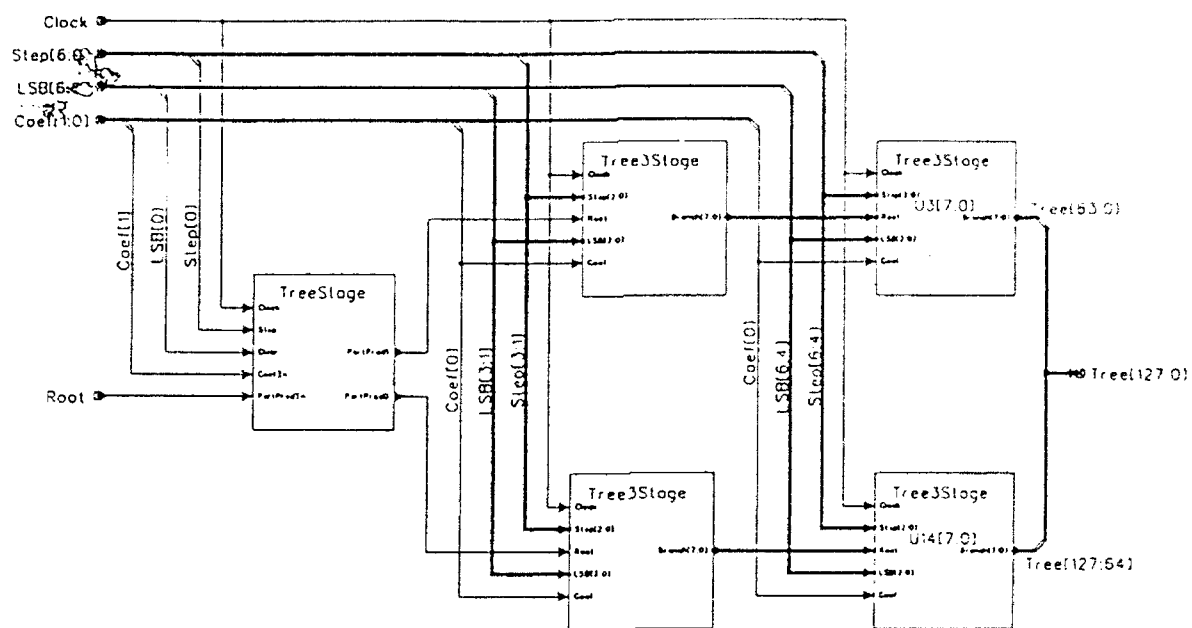


Figure 5. Linear Tree

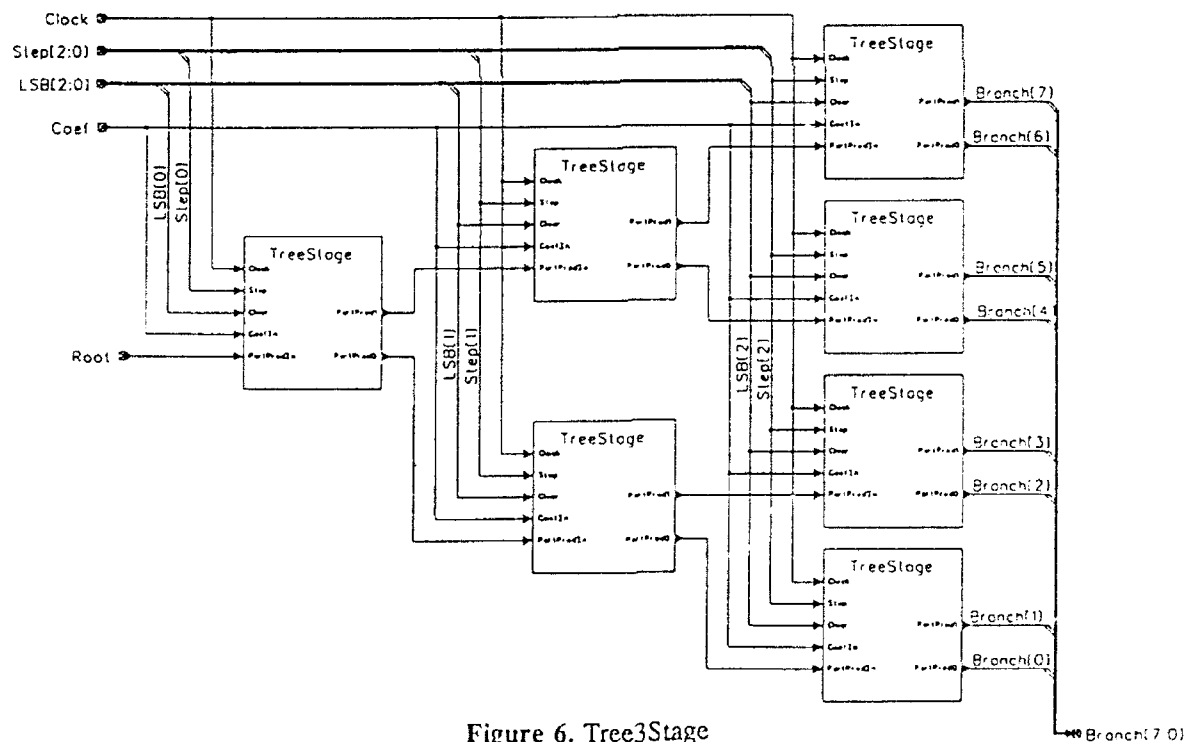


Figure 6. Tree3Stage

inputs (shown in Fig. 9) and that extra outputs now exist. The combinational model of the Tree stage can be constructed similarly. From these models, gate level and transistor level fault simulations were performed. The results are tabulated in table 1. Examination of table 1 reveals that the sub-circuits are 100% testable for

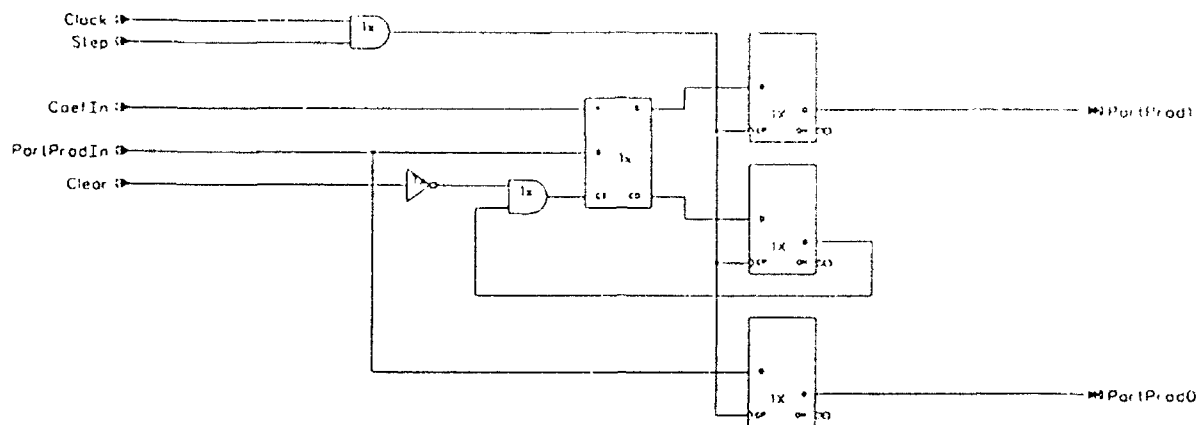


Figure 7. TreeStage

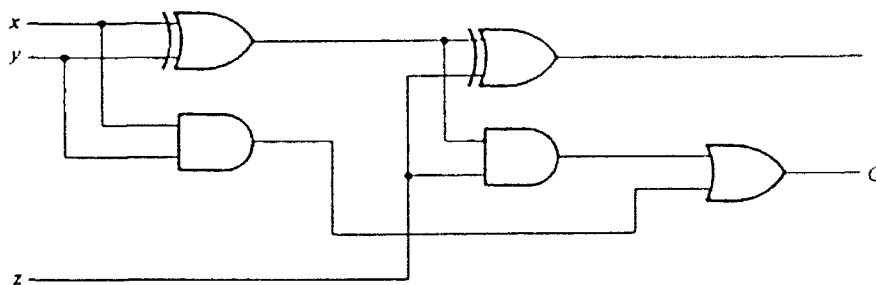


Figure 8. Full Adder

stuck-at-1, stuck-at-0, and stuck-open faults in the combinational model. This does not guarantee that the sequential circuit is 100% testable.

Table 1. Combinational fault simulation of ALU and Linear Tree

Circuit	PI	PO	GATE LEVEL				TRANSISTOR LEVEL			
			Vectors	Faults	Fault Coverage	Test Time	Vectors	Faults	Fault Coverage	Test Time
ALU Slice	18	5	9	70	100.00%	0.133 sec	20	62	100.00%	0.167 sec
Linear Tree	6	3	6	32	100.00%	0.067 sec	12	25	100.00%	0.100 sec

Sequential models of the ALU element and the single tree stage were also constructed, and sequential fault simulation was performed. Note that in sequential testing, gated signals connecting to clock inputs are omitted from testing. Results from the sequential fault simulation are shown in table 2. Examination of table 2 indicates that the sub-circuits are 100% sequentially testable provided the test patterns are applied in the

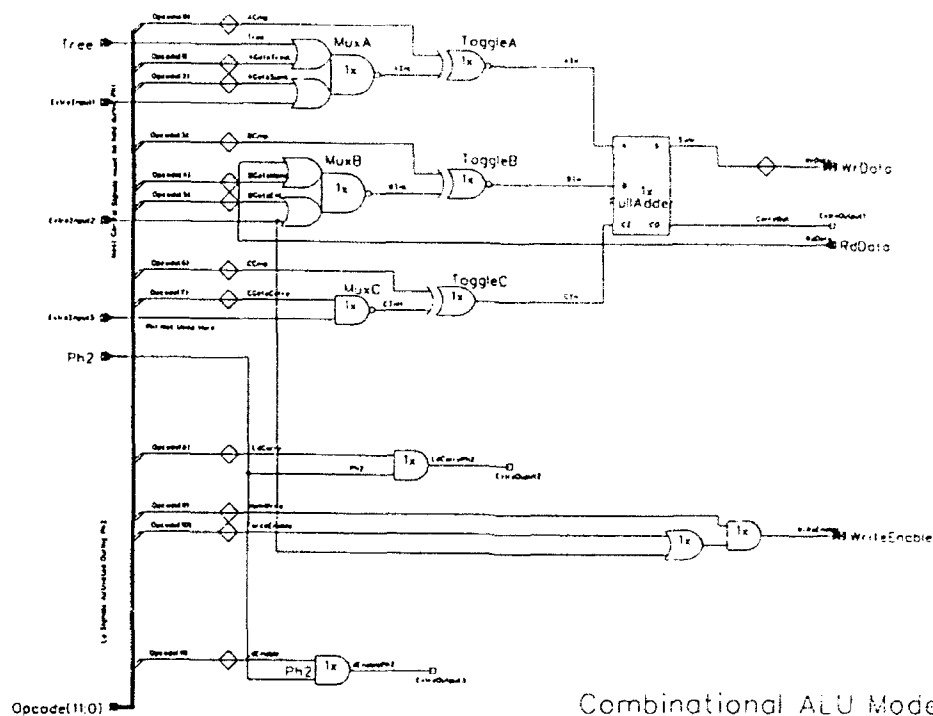


Figure 9. Combinational model of ALU element

proper order. It should be noted that both the combinational and sequential models were tested with random patterns.

Table 2. Random sequential fault simulation

Circuit	PI	PO	Vectors	Faults	Fault Coverage	Test Time
Alu Slice	12	2	17	68	100.00%	0.388 sec
Linear Tree	3	2	15	34	100.00%	0.317 sec

The most efficient method to implement Built-In Self-Test is to use Linear Feedback Shift Registers (LFSR) as test pattern generators. Linear feedback shift registers produce non-zero, non-repetitive patterns which are dependent on the previous state of the register. Given an initial seed, the LFSR cycles through a set of test patterns automatically. However, since the patterns generated from the LFSR are no longer arbitrarily chosen and placed, high fault coverage can not be guaranteed for a given test set.

The patterns generated from LFSR were then used in sequential fault simulation of the ALU slice, the single tree stage, and the DualTree stage of figure 10. The results are tabulated in table 3.

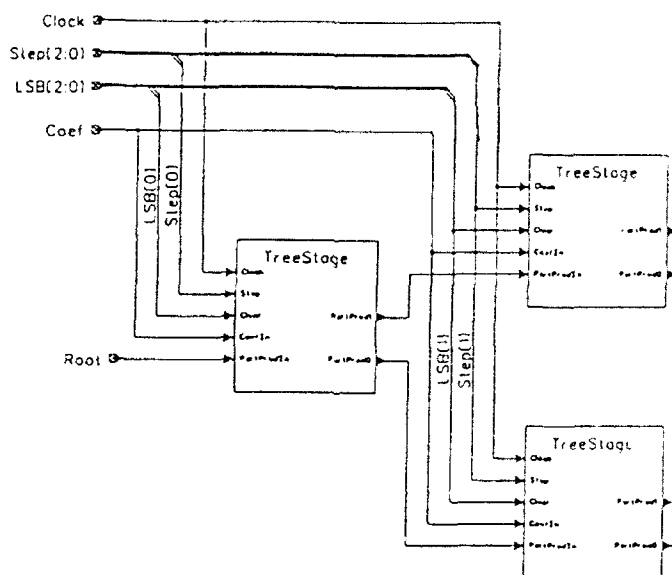


Figure 10. DualTree stage

Table 3. LFSR sequential fault simulation

Circuit	PI	PO	Vectors	Faults	Faults Detected	Faults Undetected	Fault Coverage	Test Time
ALU Slice	12	2	31	68	68	0	100.00%	0.450 sec
TreeStage	3	2	7	34	27	7	79.412%	0.263 sec
DualTree	3	4	7	96	68	28	70.833%	0.433 sec

From table 3 we can see that, even with exhaustive testing (applying all of the non-zero patterns $2^{**}n - 1$), the Tree and DualTree stages have poor fault coverage, 79.412% and 70.833% respectively. It can easily be concluded that the fault coverage will decrease more as the nested-tree structure gets deeper. On the other hand, table 3 shows that by applying only 31 tests, all 128 bit slices of the ALU can be concurrently tested with 100% fault coverage.

As a possible solution to sequentially testing the Linear Tree block, a reciprocal pattern was generated using the reciprocal primitive polynomial of the LFSR. The test set produced from such a LFSR steps through the first $2^{**}n - 1$ patterns and then steps down towards the initial seed. For example, a 3-bit LFSR with an initial seed of 100 would yield a test set of: 100,110,111,...,010,001,010,...,110,100. Using a reciprocal LFSR,

a test set was generated and sequential fault simulations were again performed on the single Tree and Dual-Tree stages. The results are given in table 4.

Table 4. LFSR sequential fault simulation with Reciprocal sequence

Circuit	PI	PO	Vectors	Faults	Faults Detected	Faults Undetected	Fault Coverage	Test Time
TreeStage	3	2	25	34	32	2	94.118%	0.367 sec
DualTree	3	4	25	96	89	7	92.708%	0.467 sec

Although the fault coverage is more encouraging, it is still not 100% and will tend to decay faster as the tree branches out. Sequential testing did not appear a viable option for the Linear Tree structure.

Since sequential testing appeared to be grossly inefficient for the Tree stage, we had to revert back to the previous combinational models. The feedback paths were again broken and replaced by extra inputs. Unlike the previous combinational fault simulations, which used randomly generated patterns to form a test set, further fault simulations were performed, which used test sets generated by LFSRs, in order to determine the Tree stages' suitability to BIST. Neglecting the gated clock signals, the fault simulation was also performed on the DualTree stage. The results are tabulated in table 5. Examining table 5, we note that for a negligible increase in test length, we can completely generate the entire test pattern, given an initial seed, from BIST hardware. While sequential testing did not provide an efficient means to test the Linear Tree structure, combinational testing provides 100% fault coverage at the expense of test length and hardware overhead.

Table 5. LFSR Combinational fault simulation

Circuit	PI	PO	Vectors	Faults	Faults Detected	Faults Undetected	Fault Coverage	Test Time
TreeStage	3	2	10	32	32	0	100.00%	0.283 sec
DualTree	3	4	11	80	80	0	100.00%	0.317 sec

4. INITIAL PLAN FOR BIST TO ALU AND LINEAR TREE ELEMENTS OF PIXEL CHIP

Having determined the suitable BIST methods for the ALU and Linear Tree elements of the PIXEL chip, we proposed the following plan to implement the BIST hardware into the existing design.

As was previously determined, the Linear Tree was to be tested from a combinational model. There are 127 single tree stages in the Linear Tree, each containing a feedback path to be broken. This requires the addition of an extra input and an extra output. Hardware requirements include 127 multiplexers to be inserted to break the feedback paths. The test mode inputs to the MUXs can all be tied to a common signal (the extra input) thus reducing the number of test signals from a possible 130 to four.

When test mode $(T1, T2) = 00$ (phase 1), the nested Linear Tree is tested using the test patterns generated from Test Pattern Generator TPG 1 (see figure 11).

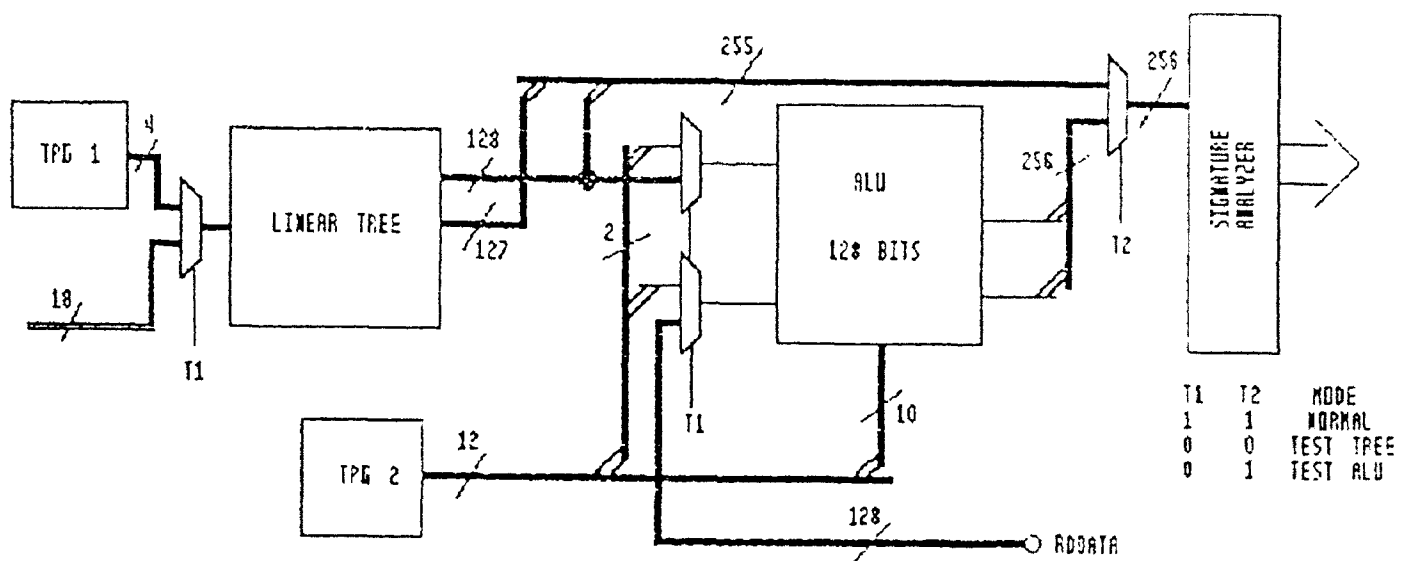


Figure 11. Initial BIST plan for ALU and Linear Tree

Since the Linear Tree is not tested sequentially, the feedback paths are broken and a common signal is simultaneously sent to each of the 127 single tree stages. Furthermore, each of the tree stages produces an extra testable output which must be routed to the signature analyzer, thus resulting in a 255-bit output from the Linear Tree.

As fault simulation has only been performed on a DualTree stage, the test length may increase rapidly as the tree branches out. If the test length exceeds the normal limit $2^{*}25$, partitioning of the Linear Tree structure into smaller subcircuits may have to be considered. However, we are fairly optimistic that this will not be necessary since the Linear Tree has only seven levels.

Again referring to previous discussion, the 128 bits of the ALU may be tested concurrently using a 12-bit test pattern generated from TPG 2 (see figure 11). Since it was determined that the ALU can be efficiently tested sequentially, no feedback paths need to be broken. However, since the ALU element contains two gated clock signals, MUXs must be inserted to break the clock path, thus enabling a single test clock (normally the system clock) to simultaneously shift data in all of the registers on the same clock edge. An example of this is shown in figure 12 for the LdEnablePh2 gated clock signal. The insertion of the multiplexers results in a hardware increase of 256 MUXs.

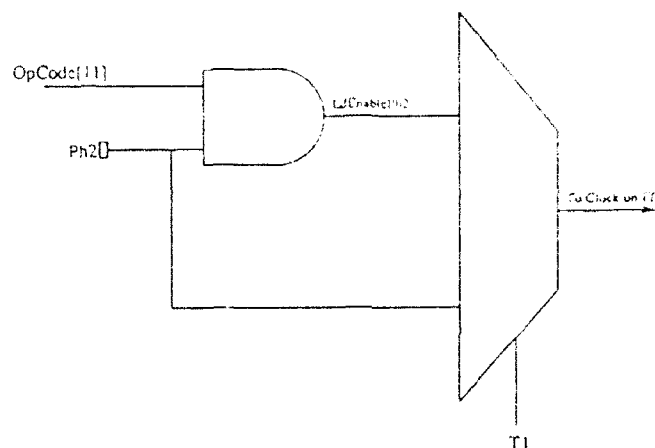


Figure 12. MUX insertion to gated clock

When test mode = 01 (phase 2), TPG 2 provides a 12-bit test pattern: 10 bits to OpCode[0-7,9-10] and one each to Tree and RdData. At the same time, OpCode[8] and OpCode[11], which control the gated clock signals LdCarryPh2 and LdEnablePh2 respectively, are disabled by the internal MUXs and the system clock is wired directly to the Enable and Carry registers. When test mode = 11 (normal operation), TPG 2, which is a Built-In Logic Block Observer (BILBO), is in the parallel register load mode, and provides the 12-bit OpCode, including the gated clock signals from OpCode[8] and OpCode[11].

In either test mode, the output signals are sent through the selection MUX to the Signature Analyzer (BILBO/MISR), modeled in figure 13, where they may be scanned out. The test pattern generators are simple LFSRs generated from the standard polynomial. No reciprocal patterns are necessary.

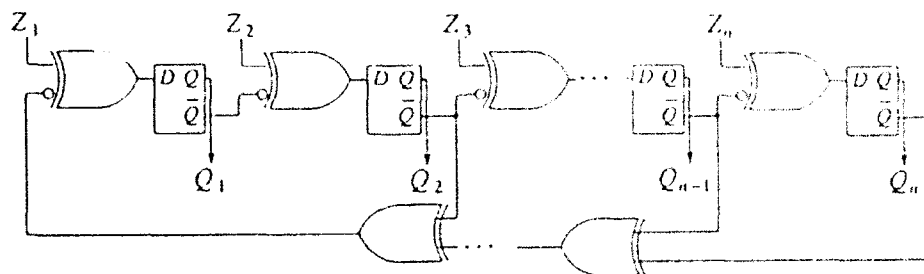


Figure 13. BILBO as a MISR

5. REEVALUATION AND UPDATED BIST PLAN

While the BIST plan for the Linear Tree structure proposed in section 4 had a low hardware overhead and a guaranteed 100% fault coverage, we did not consider routing difficulties. Based on space limitations, it was not possible to route the extra 127 outputs from the broken feedback paths of the Linear Tree to the 256 bit MISR. We concluded that a sequential model must be used to implement BIST in the Linear Tree and thus eliminate the extra output lines.

In Section 3, we showed that the Linear Tree can be tested sequentially while achieving 100% fault coverage (see table 2). However, the patterns generated to perform this simulation were not generated using a LFSR as a test pattern generator. Referring to table 3, we showed that the test patterns generated from the LFSR test pattern generator did not provide a high fault coverage. The lower fault coverage is a result of the LFSR removing the randomness in the test pattern generation. While the LFSR test pattern generator will generate non-repetitive, non-zero test patterns, the same sequence of patterns always develops after every $n-1$ clock cycles. In order to improve the fault coverage while using LFSR generated test patterns, we must scramble the patterns, thereby restoring some randomness to the test set.

A derivative of Circular Built-In Self-Test was implemented on the single tree and DualTree stages. The Circular BIST idea is to chain all of the input, output, and selected internal registers into a Circular Self-Test Path where the next state of the registers depends on the current state of the register as well as the combinational logic. To implement Circular BIST, all of the registers in the circular path must be modified as shown in figure 14. While this method yielded from 98% to 100% fault coverage, the increase in hardware overhead

(albeit a small increase) and the inability to set or reset the registers to an initial state forced us to abandon the Circular BIST approach.

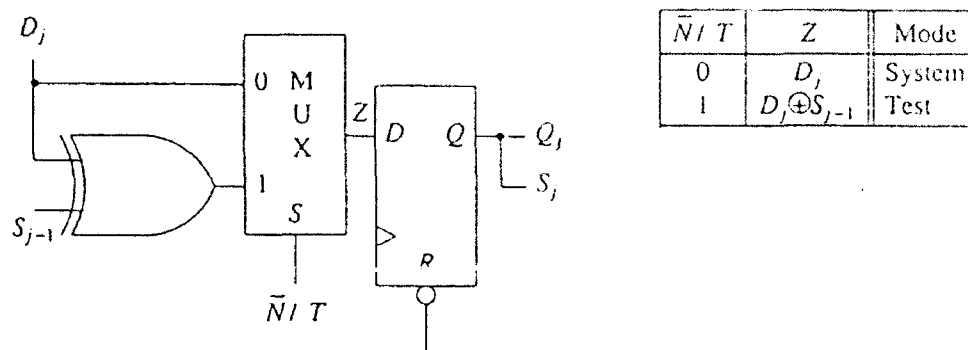


Figure 14. Circular BIST Register

By scrambling the connections between the LFSR test pattern generator and the inputs of the Linear Tree, as shown in figure 15, we were able to increase the fault coverage to 94.118%. Although an improvement, the fault coverage is still too low.

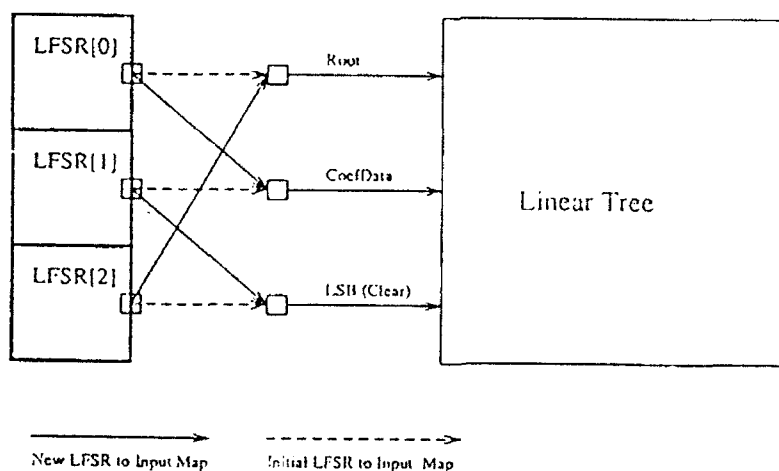


Figure 15. Scrambled test signals

Another method to increase the randomness in the test pattern generation, is to increase the length of the LFSR. Since only the first three bits will be used, the test pattern is no longer non-repetitive and non-zero. While this method provides an advantageous increase in the randomness of the test patterns, the increase in

hardware overhead is its obvious disadvantage. Disturbingly, this method was employed and yielded the same results as table 2.

Combining the previous two methods, we increased the length of the LFSR and we also scrambled the input connections. In addition, an easily testable Full Adder was used (shown in figure 16).

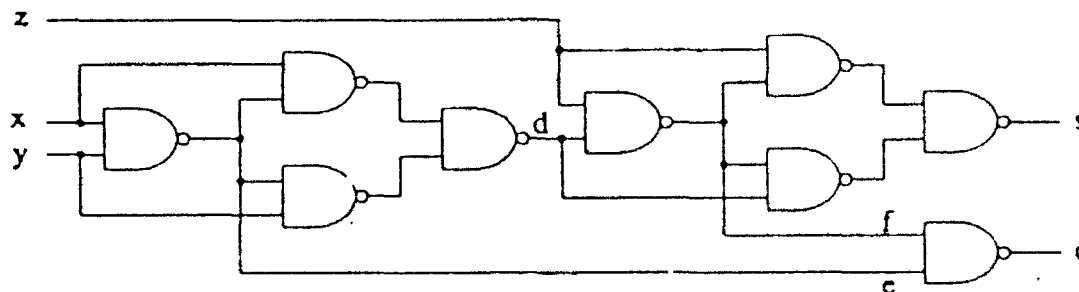


Figure 16. Testable full adder

The sequential fault simulation results for the single tree stage, the DualTree stage, and the entire Linear Tree block using various LFSRs, input sequences, and Full Adder models are tabulated in table 6.

Table 6 shows that by increasing the randomness of the test pattern generation using larger LFSRs and scrambling the input sequence, the Linear Tree can be sequentially tested with 100% fault coverage.

By sequentially testing the Linear Tree, the extra 127 output lines are eliminated, and only the 128 outputs of the Linear Tree must be sent to the MISR. In order to reduce the size of the MISR, we attempted to limit observation of the ALU slice to the WriteEn output only. Since the WriteData output was connected to a feedback path, it was our hope that the ALU could still be sequentially tested and achieve 100% fault coverage by propagating all of the faults through the WriteEn output. Sequential fault simulation results showed that by increasing the test pattern length from 31 vectors to 107 vectors, the ALU could still be tested with 100% fault coverage. Therefore, the size of the MISR can be reduced from 256 bits to 128 bits as both the Linear Tree and the ALU need only 128 outputs to be observed. An updated BIST plan for the ALU and Linear Tree is shown in figures 17 and 18.

Note also that the test pattern generators are no longer local to the Pixel Plane. Since the gated clocks in the ALU were globalized in the Pading, the 12-bit LFSR test pattern generator used to test the ALU has

been moved to the Pading. Also, since the Linear Tree is to be sequentially tested, the 8-bit LFSR TPG has likewise been moved to the EMCTrunk inside the Pading. Using multiplexers, the test signals can be sent along the appropriate buses to the subsystems under test.

Table 6. LFSR sequential fault simulation

Circuit	LFSR	FA	Vectors	Faults	Faults Detected	Faults Undetected	Fault Coverage	Test Time	Input Map
TreeStage	4bit	xor	255	34	27	7	79.412%	0.42 sec	0 0 1 1 2 2
TreeStage	4bit	xor	255	34	32	2	94.118%	0.39 sec	0 2 1 1 2 0
TreeStage	4bit	nand	255	42	33	9	78.571%	0.47 sec	0 0 1 1 2 2
TreeStage	4bit	nand	255	42	41	1	97.619%	0.40 sec	0 2 1 1 2 0
TreeStage	4bit	nand	255	42	41	1	97.619%	0.42 sec	0 2 1 0 2 1
TreeStage	4bit	xor	255	34	32	2	94.118%	0.37 sec	0 2 1 0 2 1
TreeStage	5bit	xor	255	34	33	1	97.059%	0.38 sec	0 2 1 0 2 1
TreeStage	5bit	nand	14	42	42	0	100.000%	0.27 sec	0 2 1 0 2 1
TreeStage	5bit	nand	14	42	42	0	100.000%	0.30 sec	0 2 1 1 2 0
TreeStage	6bit	xor	58	34	34	0	100.000%	0.32 sec	0 2 1 0 2 1
DualTree	4bit	xor	255	96	91	5	94.792%	0.70 sec	0 2 1 0 2 1
DualTree	4bit	xor	255	96	85	11	88.542%	0.62 sec	0 2 1 1 2 0
DualTree	5bit	xor	255	96	88	8	91.667%	0.60 sec	0 2 1 1 2 0
DualTree	5bit	xor	255	96	93	3	96.875%	0.63 sec	0 2 1 0 2 1
DualTree	6bit	xor	255	96	90	6	93.750%	0.63 sec	0 2 1 1 2 0
DualTree	6bit	xor	59	96	96	0	100.000%	0.45 sec	0 2 1 0 2 1
DualTree	5bit	nand	255	120	109	11	90.833%	0.72 sec	0 2 1 1 2 0
DualTree	5bit	nand	255	120	111	9	92.500%	0.72 sec	0 1 1 2 2 0
DualTree	5bit	nand	255	120	119	1	99.167%	0.65 sec	0 2 1 0 2 1
DualTree	6bit	nand	39	120	120	0	100.000%	0.47 sec	0 2 1 0 2 1
Linear Tree	10bit	xor	255	3816	3557	259	93.213%	27.4 sec	0 0 1 1 2 2
Linear Tree	8bit	xor	255	3816	3803	13	99.659%	32.3 sec	0 2 1 1 2 0
Linear Tree	7bit	xor	255	3816	3808	8	99.790%	52.9 sec	0 2 1 0 2 1
Linear Tree	8bit	xor	227	3816	3816	0	100.000%	28.5 sec	0 2 1 0 2 1

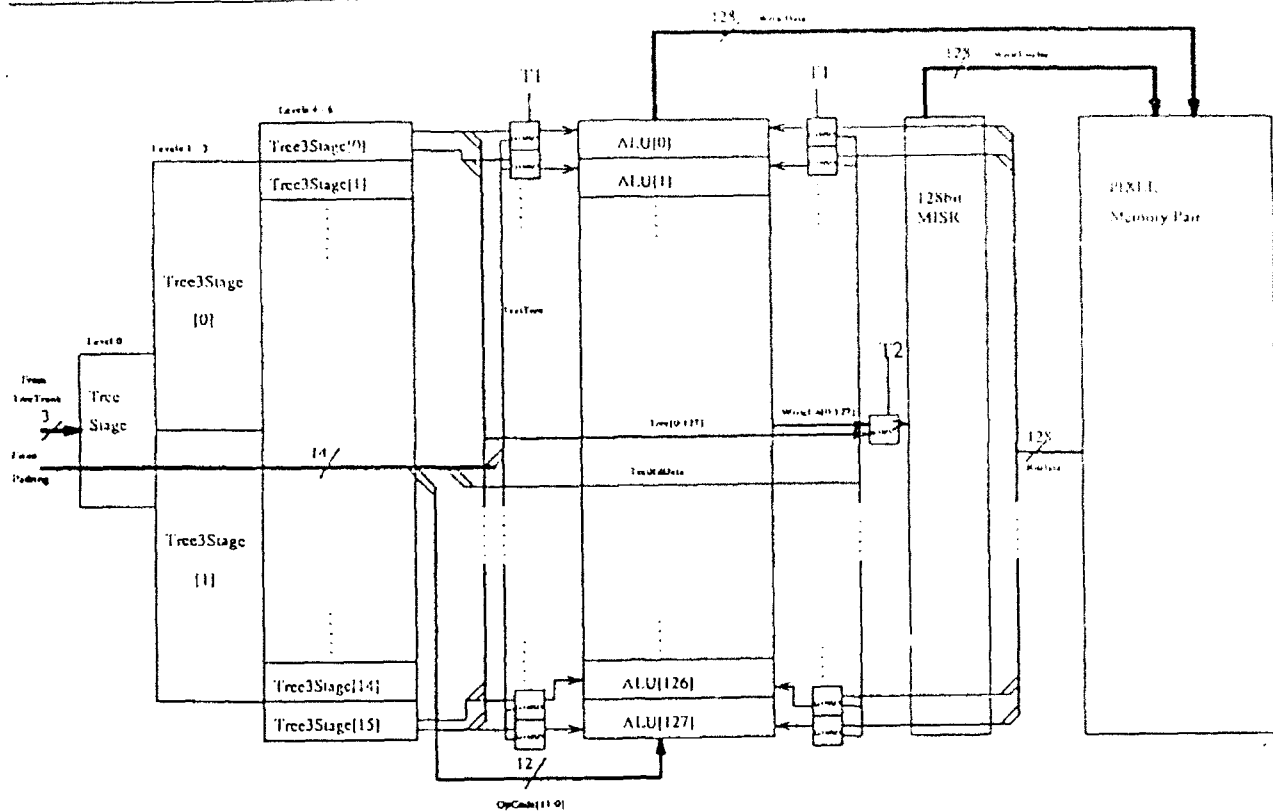


Figure 17. Final BIST plan for the ALU and Linear Tree

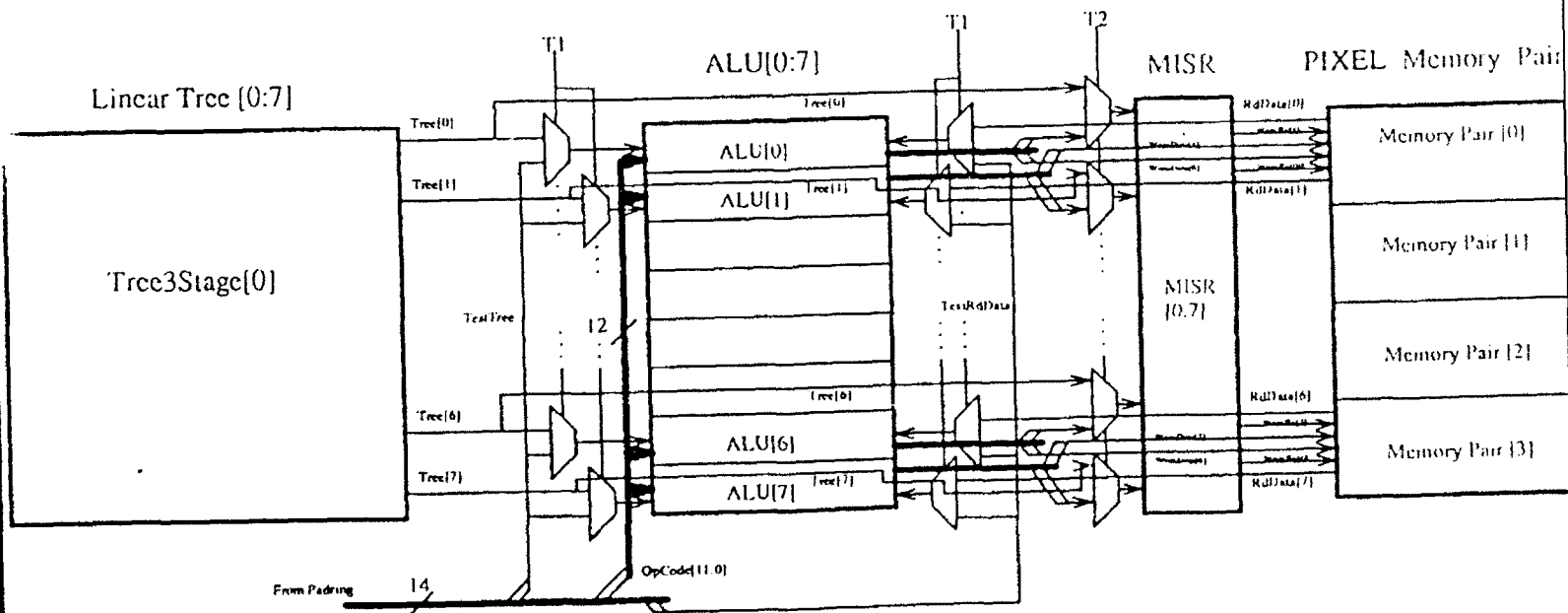


Figure 18. Final BIST plan slice

6. CONCLUSION AND TOPICS FOR FUTURE RESEARCH

As can be seen from the above results, sequential testing offers a much more efficient test methodology than combinational testing, evident most notably in a much lower hardware overhead and ease of signal routing. Although straightforward sequential testing does not always provide a higher fault coverage than its combinational counterpart, new methods of sequential testing have been developed that drastically improve fault coverage.

A new technique for designing self-testing sequential circuits, referred to as Circular BIST, was analyzed in this research. To implement Circular BIST, all of the registers in the sequential circuit must be modified as shown in figure 14. The test patterns are generated from LFSRs and fed into the primary inputs of the sequential circuit. The output response from the primary outputs is fed into a MISR for the signature analysis. The stuck-at fault simulation results from benchmark circuits shown in table 7 indicate that the circular BIST on sequential circuit achieves a higher fault coverage in comparison to several previous work on sequential circuit testing. Employing the partial scan into the Circular BIST design to achieve even higher fault coverage is currently under investigation and also a topic for future research.

Table 7. Circular BIST fault simulation

Circuit	PI	PO	FF	Gates	Faults	Vectors	Faults Detected	Faults Undetected	Fault Coverage	Run Time	Hardware Overhd
s298	3	6	14	150	366	395	366	0	100.000%	2.30 sec	17.07%
s344	9	11	15	199	418	256	417	1	99.761%	1.82 sec	14.02%
s349	9	11	15	200	426	398	423	3	99.296%	2.45 sec	13.95%
s382	3	6	21	203	495	766	495	0	100.000%	5.15 sec	15.55%
s400	3	6	21	207	520	1358	514	6	98.846%	7.95 sec	18.42%
s420	19	2	16	247	494	20500	472	22	95.547%	83.40 sec	12.17%
s444	3	6	21	226	570	478	556	14	97.544%	4.75 sec	17.00%
s510	19	7	6	242	588	20000	562	26	95.578%	109.82 sec	4.84%
s526	3	6	21	238	641	6538	635	6	99.064%	318.73 sec	16.22%
s641	35	24	19	452	581	20000	574	7	98.795%	144.30 sec	8.07%
s713	35	23	19	466	695	20000	650	45	93.525%	167.40 sec	7.54%
s820	18	19	5	317	870	20000	847	23	97.356%	213.30 sec	3.11%
s832	18	19	5	315	890	20000	852	38	95.730%	223.20 sec	3.15%
s838	35	2	32	489	985	20000	885	100	89.848%	199.60 sec	12.28%
s953	16	23	29	469	1241	20000	1211	30	97.583%	155.68 sec	11.65%
s1196	14	14	18	579	1338	20000	1333	5	99.626%	202.50 sec	6.03%
s1238	14	14	18	558	1449	20000	1375	74	94.893%	267.50 sec	6.25%
s1423	17	5	74	822	1835	20000	1819	16	99.128%	331.00 sec	16.52%
s1488	8	19	6	673	1510	512	1338	172	88.609%	14.20 sec	1.75%
s1494	8	19	6	667	1530	2048	1390	140	90.850%	36.50 sec	1.76%
s3378	35	49	179	3172	5397	20000	4742	655	87.864%	1873.01 sec	10.68%
s9234	19	22	228	6072	8001	20000	5514	2487	68.916%	6109.73 sec	7.24%
s13247	31	121	669	9320	13167	20000	8930	4237	67.821%	9727.57 sec	13.39%
s35932	35	320	1728	19556	46094	20000	42020	3984	91.340%	17394.3 sec	16.28%
Average	17	31.4	133	1910	3750.4	15000.3	1246.67	503.79	93.647%	1554.82 sec	10.71%

A STUDY OF DAMAGE IN GRAPHITE EPOXY PANELS SUBJECTED
TO LOW AND HIGH VELOCITY IMPACT

Mohammed A. Samad
Graduate student
Department of Mechanical Engineering

University of New Orleans
Lake Front, New Orleans
Louisiana 70148

Final Report For:
Summer Research Program
Wright Laboratory
Flight Dynamics Directorate
Vehicle Subsystem division
Survivability Branch

sponsored by:
Air Force Office of Scientific Research
Bolling Air Force Base, Washington, D.C.

September 1992

A STUDY OF DAMAGE IN GRAPHITE EPOXY PANELS SUBJECTED
TO LOW AND HIGH VELOCITY IMPACT

Mohammed A. Samad
Graduate Student
Department of Mechanical Engineering
University of New Orleans

Abstract

The Effect of Graphite Epoxy composite panel layup on damage was studied using constant angle change between adjacent plies. Use 4 different type of layup, constant angle change used, A) 0/90 B) 0/45 C) 0/22.5 D) 0/11.25. The composite panels were impacted at a series of velocities from 400 ft/sec to 6000 ft/sec with a 1/2 inch diameter steel sphere. The loss in weight due to impact, initial and residual velocity were measured and latter by deplying method the damage for each ply were analyzed. The total delamination area is correlated with the initial velocities. The ballistic limit (V50) also measured for all 4 layup. The ballistic limit is correlated with constant angle change between the plies, as the angle decreases between the plies and the ballistic limit increases.

A STUDY OF DAMAGE IN GRAPHITE EPOXY PANELS SUBJECTED TO LOW AND HIGH VELOCITY IMPACT

Mohammed A. Samad

INTRODUCTION

Graphite epoxy composite material consist of fibers of graphite embedded in a matrix. Composite materials have more advantages as compare to other materials like light in weight, high strength, and easy to mold into complicated shell structures. Being a new material some of its properties are not yet known fully. Damage due to impact is one of them, research on the damage of graphite epoxy composite panels due to impact has been studied extensively over past many years. However, the constant change in angle between adjacent plies in composite panels has never been done before.

The experiment work of this research includes impacting composite panels with 4 different layup, A) $[(0/90)_8]_s$, B) $[(0/+45/90/-45)_4]_s$ C) $[(0/+22.5/+45/+67.5/90/-67.5/-45/-22.5)_2]_s$, and D) $[(0/+11.25/+22.5/+33.75/+45/+56.25/+67.5/+78.75/90/-78.75/-67.5/-56.25/-45/-33.75/-22.5/-11.25)]_s$, at a series of velocities with a 1/2 inch diameter steel sphere, measurement of the thickness of the panel before impact, measurement of weights of the panel before and after impact,

calculating initial and residual velocities, collecting spall from the front and back of the plate to study the behavior of the damage latter. The main purpose of this research is to analyze the effect of angle change on damage by deplying method and energy absorbed per unit area. The total damage area is found to be decreases with the increase of velocity. The ballistic limit for 4 different layup is also studied.

METHODOLOGY

The composite panels is impacted at a known velocity with a 1/2 inch diameter steel sphere. The figure 1 shows all the test setup of gas gun, fixture, composite panel, coil tube, and break papers. The dimension of the plate is 5"x 5", 32 ply and constant angle change between the adjacent ply. The initial and residual weight of the panel and thickness at the point of impact were measured. Before the panel is impacted a few air shots were done to get exact gas (Helium) pressure for required velocity. Table 1 shows a sample of a few shots with different layup, initial and residual velocities, thickness of plates, weights of plates before and after impact, and weight of spall collected from back and in front of the panel.

To calculate the initial velocity of the projectile a wire or a small piece of graphite is placed across the barrel, also a small piece of graphite is placed in front of the fixture. These graphite are connected to a circuit. As the projectile break the graphite a

signal is sent to a oscilloscope called "Kontron". The Kontron is used to determine the time between the two graphite, by knowing the distance between the two graphite and time, velocity can be calculated.

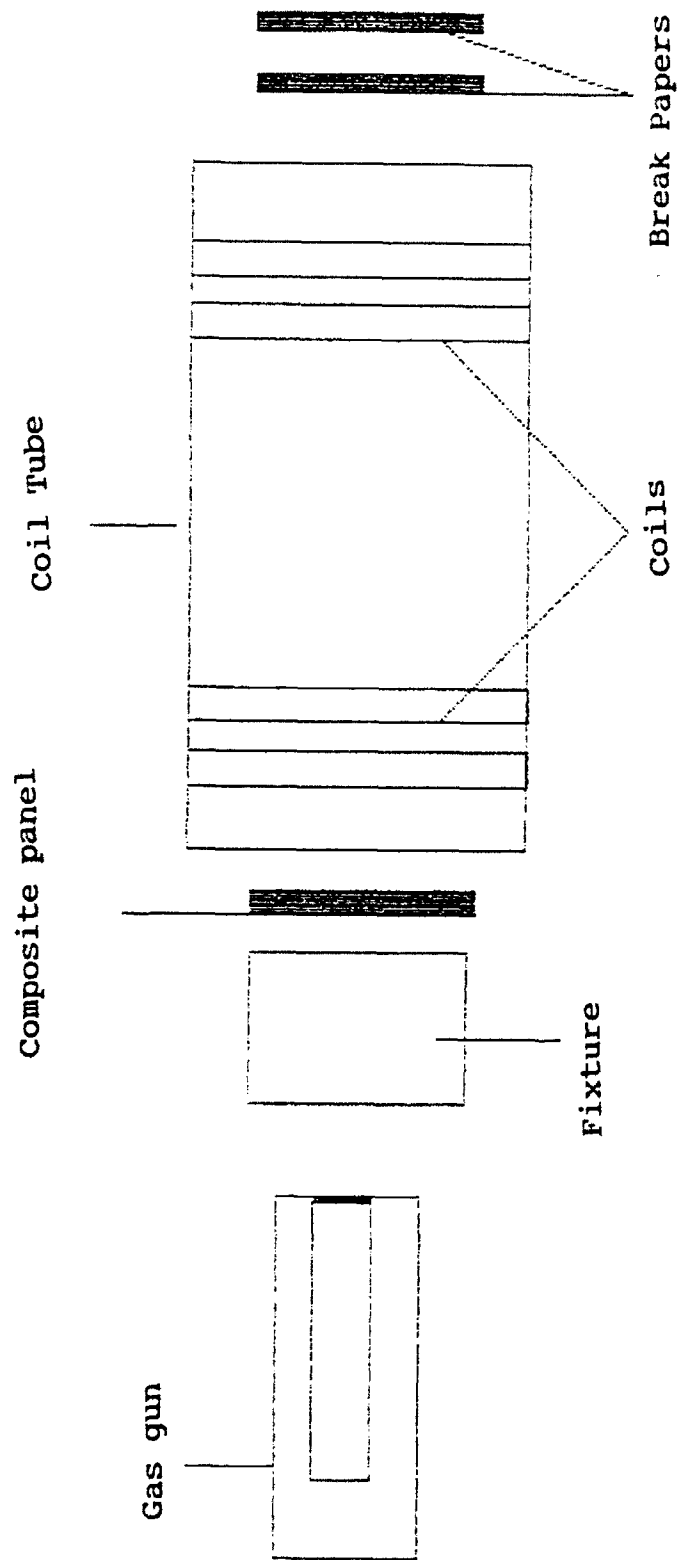
The residual velocity can be calculated by two methods 1) by coils 2) by break papers. Coils are wrapped around coil tube as shown in fig.1, the distance between the coils is 24 inches, these coils produce a two different magnetic field in the tube. when the projectile pass through the coil tube, it disturbed the magnetic field and a signal is sent to Kontron. The time between two coils is calculated from Kontron and the residual velocity can be calculated.

In the second method the break papers are used to calculate the residual velocity. These break papers have continuous lines of silver ink, one is placed at the back of the coil tube and the second paper is placed at a distance of 24 inches from the first paper and these are connected to separate circuits. When the projectile pass through the first and the second paper, a high voltage signal can be read in Kontron when it breaks the circuit and the time between the two papers is calculated from Kontron and the residual velocity.

After the experimental work is done, the panels are deplied and total damage area is calculated. A series of process has to done to deply the panel. A Wax wall is built over the hole of the panel created by the impact and gold chloride solution is poured

into the hole for 3 to 4 times and soak the panel for 2 days. After two days the panel is heated in the oven for approximately about half hour at 800 degree Fahrenheit to burn the matrix of the composite panel and deply each ply. When the panel is heated at high temperature gold chloride turn into gold and appear at the damage of each ply. A white paint line is drawn across the gold, the white line helps in measuring the delamination area quickly under image analyzer. The image analyzer consist of a camera, a RGB monitor, and a computer. Figure 2 shows the setup of image analyzer First a ply is placed under a camera and an image of the ply is displayed on monitor. By using JAVA a software system, damage area under white paint line is calculated. This procedure is repeated for 32 plies to calculate the damage of a single panel impacted at a known velocity.

C-scan is another method to find out the damage of the panel. A few panels were c- scan to compare the damage area calculated by deplying method and c-scan. Deplying method give more accurate result than c-scan because in deplying method it give the damage for each ply were as in c-scan it gives the total damage. A separate test is conducted to study the ballistic limit of the each layup, several panel were shot at each layup and calculate the ballistic limit. It is found that the ballistic limit increase with decrease in angle between ply.



TEST SETUP

Figure 1

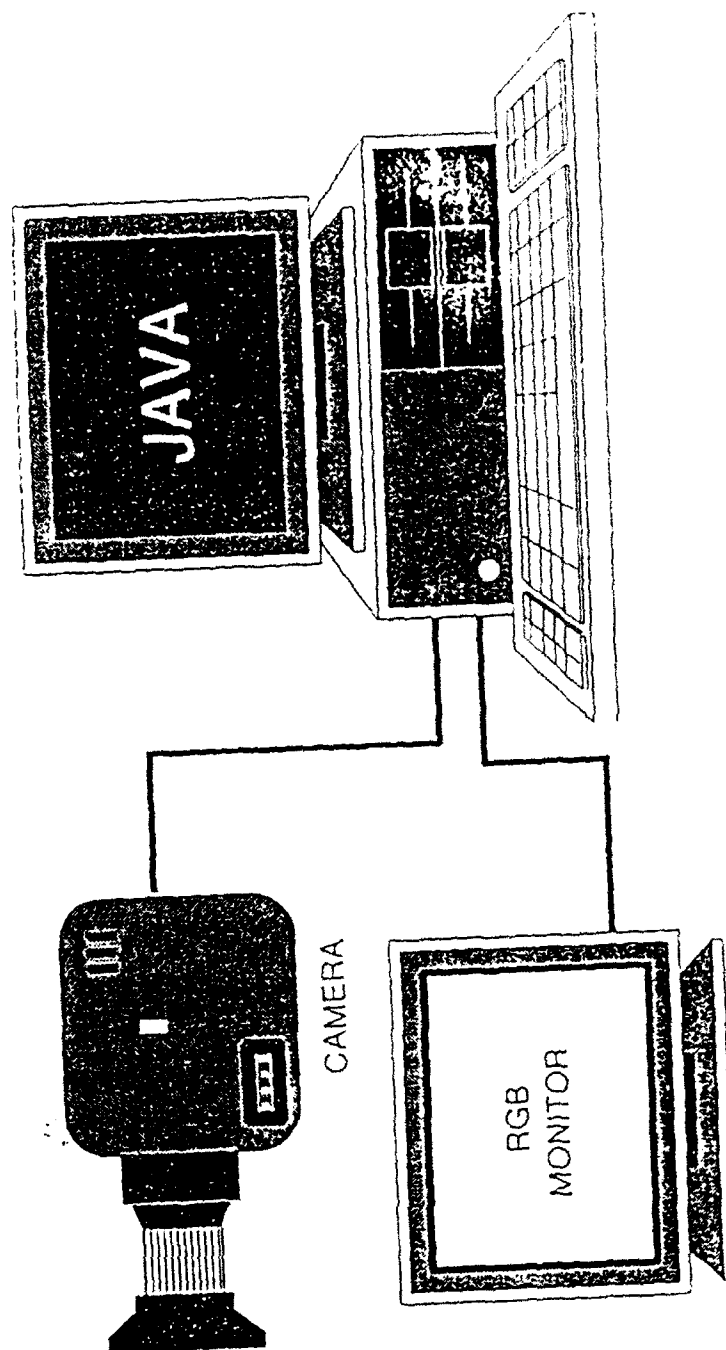


IMAGE ANALYZER

Figure 2

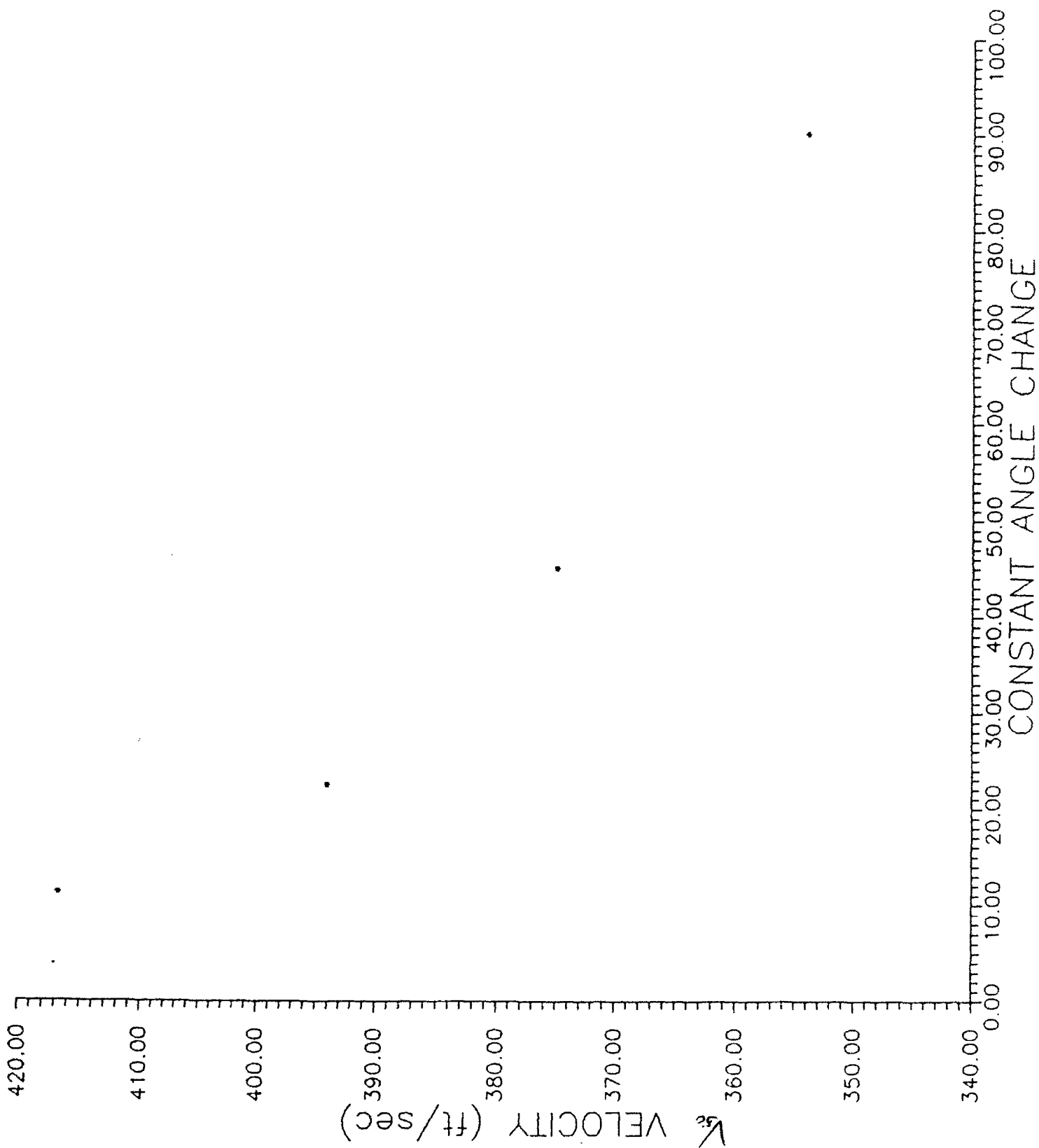


Table 1

PANEL I.D. #	THICKNESS In.	INIT. VEL. ft/sec	RES. VEL. ft/sec	INIT. WT. Grams	RES. WT. Grams
PANEL LAYUP: 0/90					
F 1.2	0.183	1990	1683	115.55	113.83
F 1.3	0.184	1980	1731	114.06	112.54
F 1.4	0.184	397	226	116.54	116.35
F 1.5	0.183	398	213	117.26	117.09
F 1.7	0.184	1003	774	115.70	114.69
F 1.8	0.185	989	798	116.36	115.15
F 1.9	0.183	994	796	114.80	113.65
F 2.1	0.183	1007	817	116.07	114.13
F 2.2	0.186	696	575	116.10	113.88
F 2.3	0.183	694	331	115.32	113.27
F 2.4	0.185	696	557	116.75	115.03
PANEL LAYUP 0/45					
F 9.1	0.175	994	846	111.79	110.65
F 9.2	0.179	696	576	113.26	112.58
F 9.3	0.176	697	581	112.11	111.45
F 9.4	0.185	696	568	114.78	114.01
F 10.2	0.177	2008	1636	114.50	112.98
F 10.3	0.176	1994	1709	111.63	110.23
F 10.4	0.182	1967	1736	114.58	112.99
F 10.5	0.184	392	103	116.78	116.60
F 10.6	0.179	388	----	112.72	112.56

F 10.7 0.183
F 10.8 0.183
F 10.9 0.175

398
993
1091

858
867

112.37
114.01
110.92

112.19
112.67
109.87

PANEL LAYUP: 0/22.5

F 11.1 0.173
F 11.2 0.176
F 11.3 0.171
F 11.4 0.175
F 11.5 0.179
F 11.6 0.176
F 11.7 0.171
F 11.8 0.175
F 11.9 0.171
F 12.1 0.177
F 12.2 0.182
F 12.3 0.178

994
992
994
705
692
697
392
395
396
1983
2015
1983

870
878
862
581
568
565
0
0
0
1777
1778
1798

108.68
110.41
108.05
110.35
112.69
109.52
107.34
109.54
107.73
110.63
113.39
113.37

107.28
108.90
106.66
109.02
111.64
108.37
115.87 Inc. ball wt.
117.56 Inc. ball wt.
107.63
109.11
111.66
109.94

PANEL LAYUP: 0/11.25

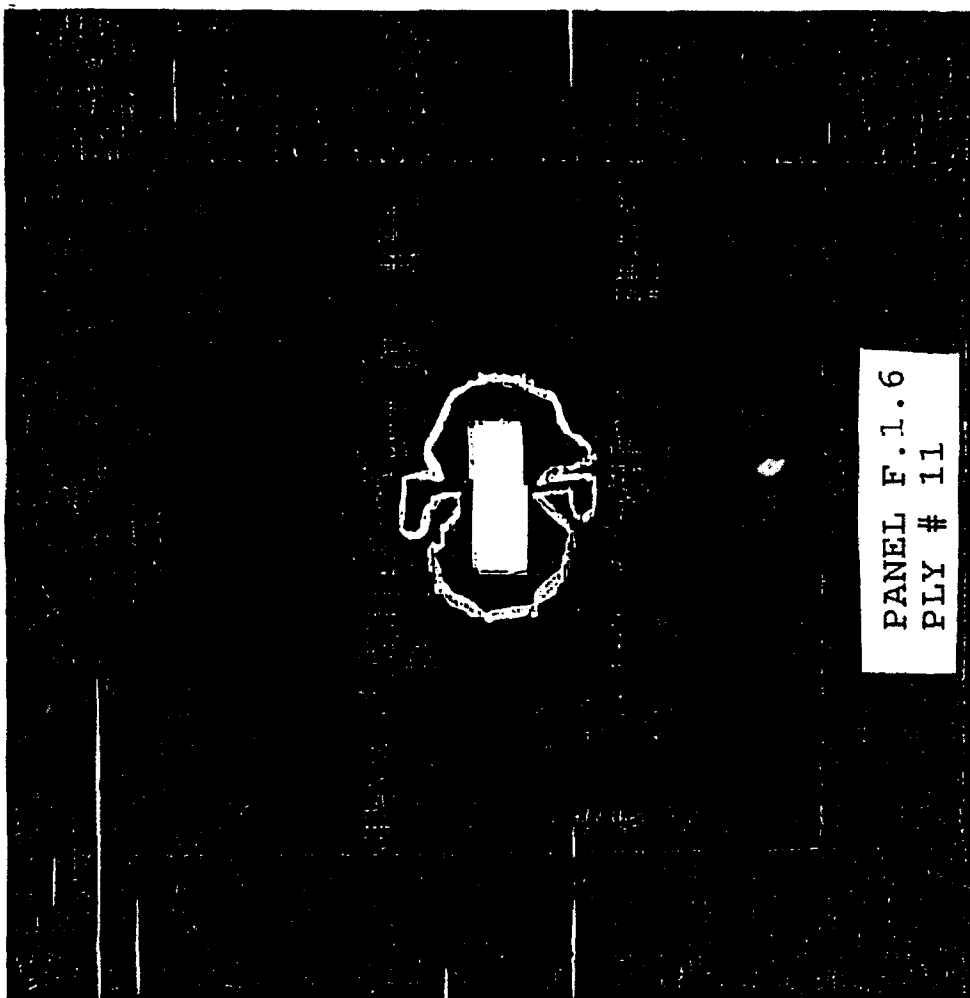
F 16.1 0.176
F 16.2 0.183
F 16.3 0.179
F 16.4 0.183
F 16.5 0.180
F 16.6 0.183
F 16.7 0.160
F 16.8 0.182
F 16.9 0.179
F 17.1 0.178
F 17.2 0.183
F 17.3 0.180

716
709
5112
394
394
398
1997
2046
991
1003
991
1965

577
569
410
264
0
0
1808
1792
850
857
855
1787

109.05
113.83
112.60
114.43
116.47
113.04
117.71
113.67
111.90
111.96
114.82
111.98

107.87
112.68
112.35
114.35
124.78
113.01
109.23
111.33
110.81
110.70
113.38
109.77



PANEL F.1.1.6
LAYUP 0/90
VELOCITY = 397 ft/sec

Deplied ply shot at 397 ft/sec,
The area under White line is delamination area.

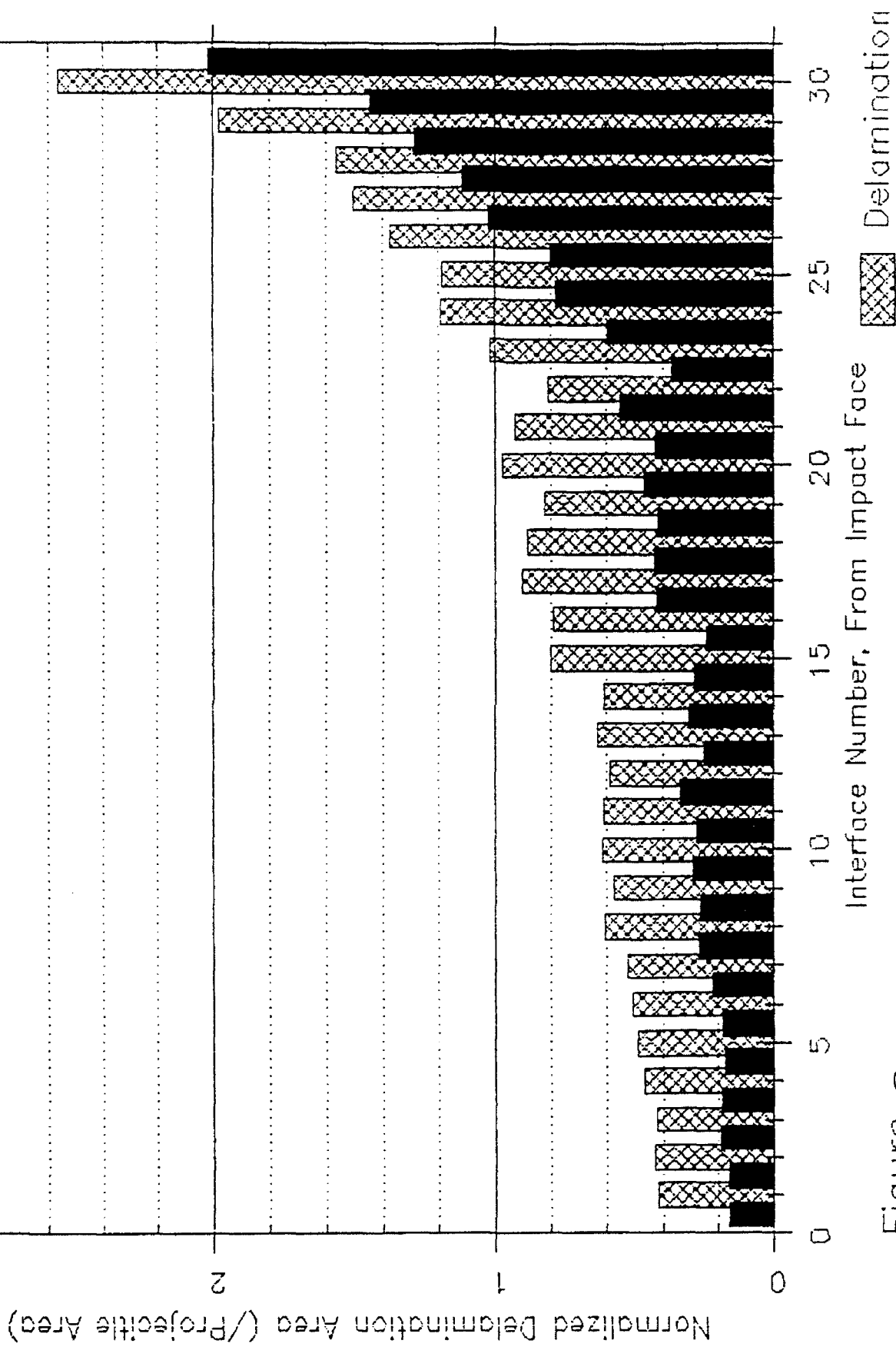


Figure 3
Delamination and Hole Area of Individual Plies
For Panel Impacted at 697.3 Ft/Sec, 0/90 Layup

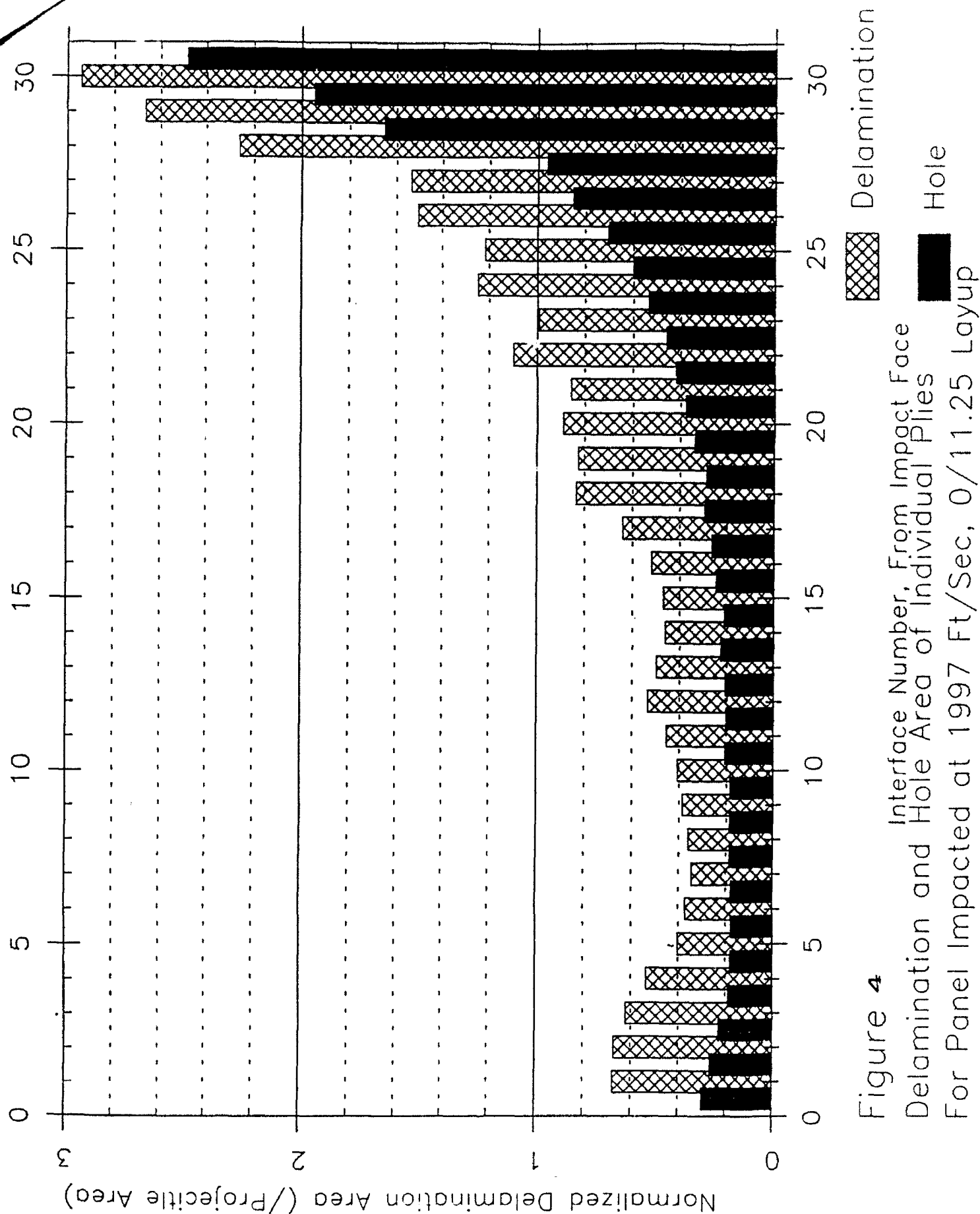


Figure 4

Interface Number, From Impact Face
Delamination and Hole Area of Individual Plies

For Panel Impacted at 1997 Ft/Sec, 0/11.25 Layout

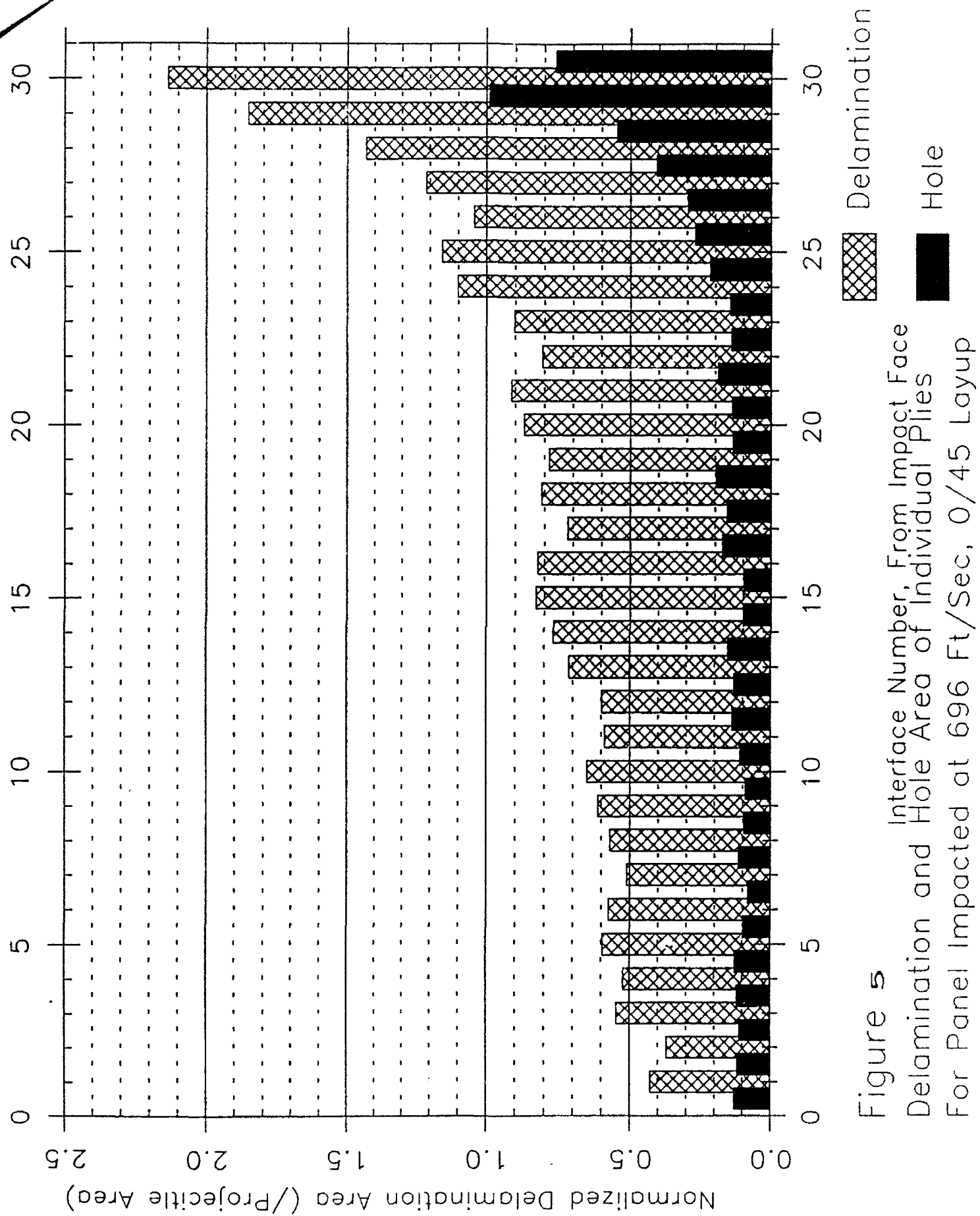
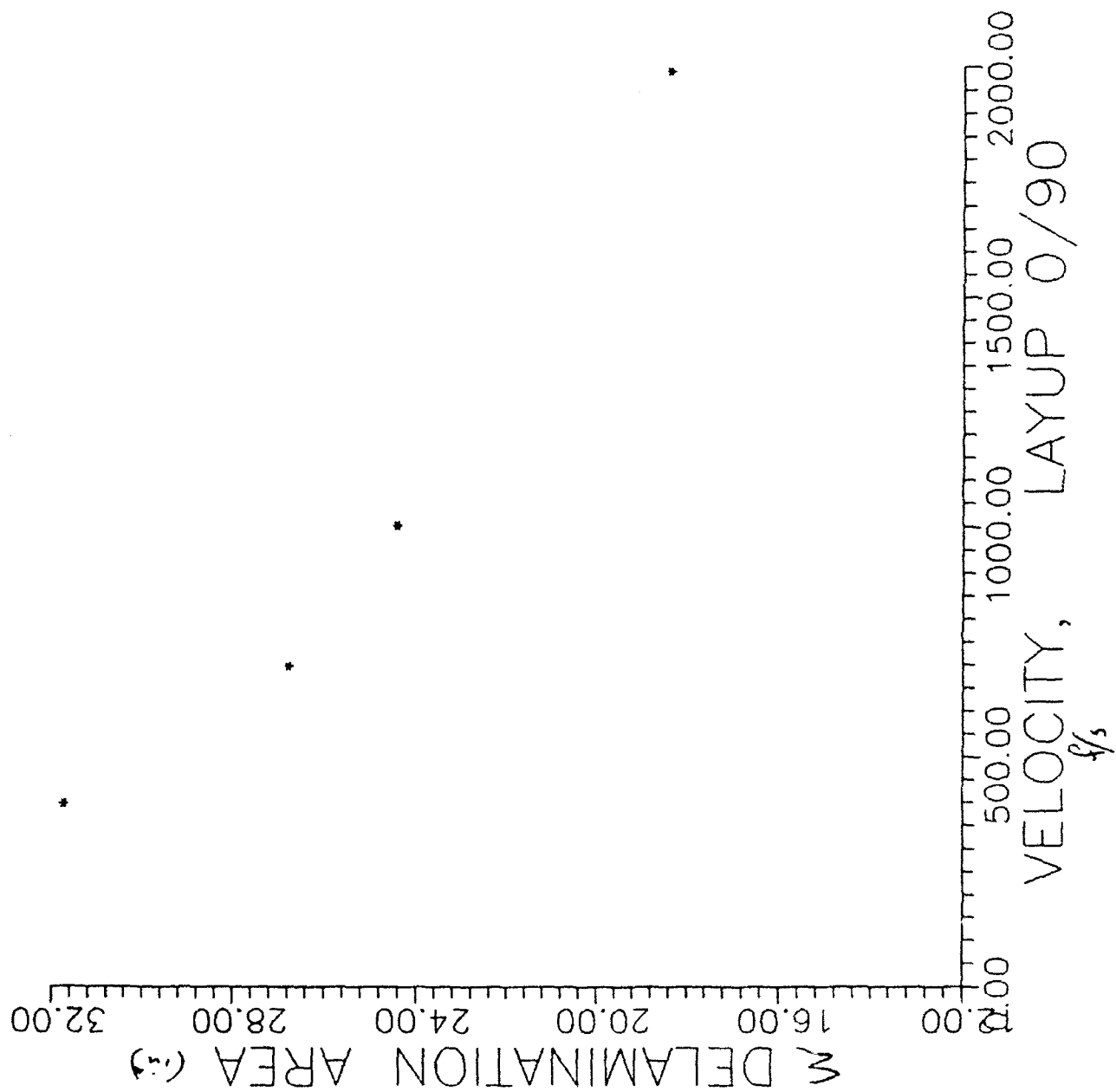
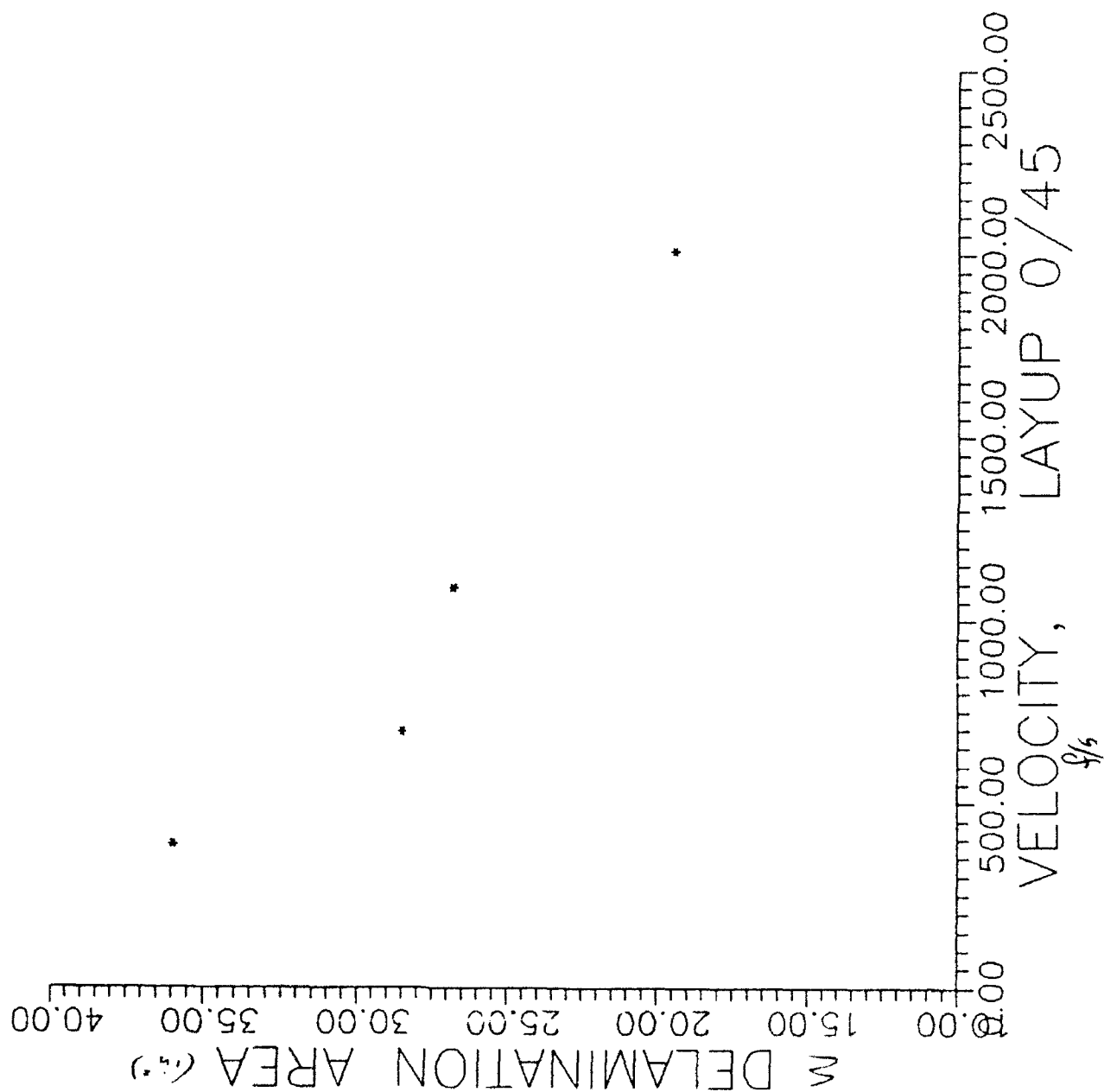
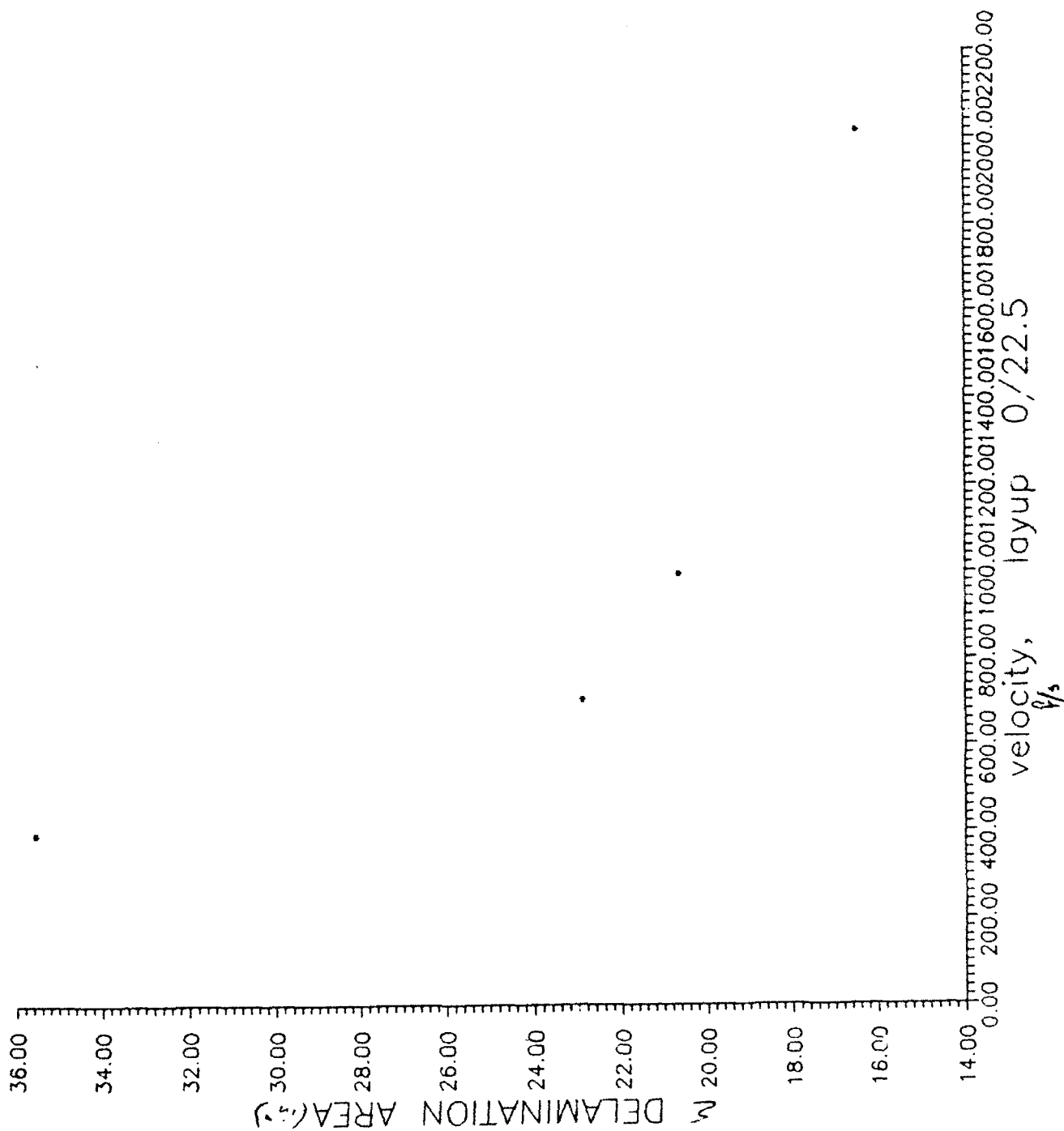
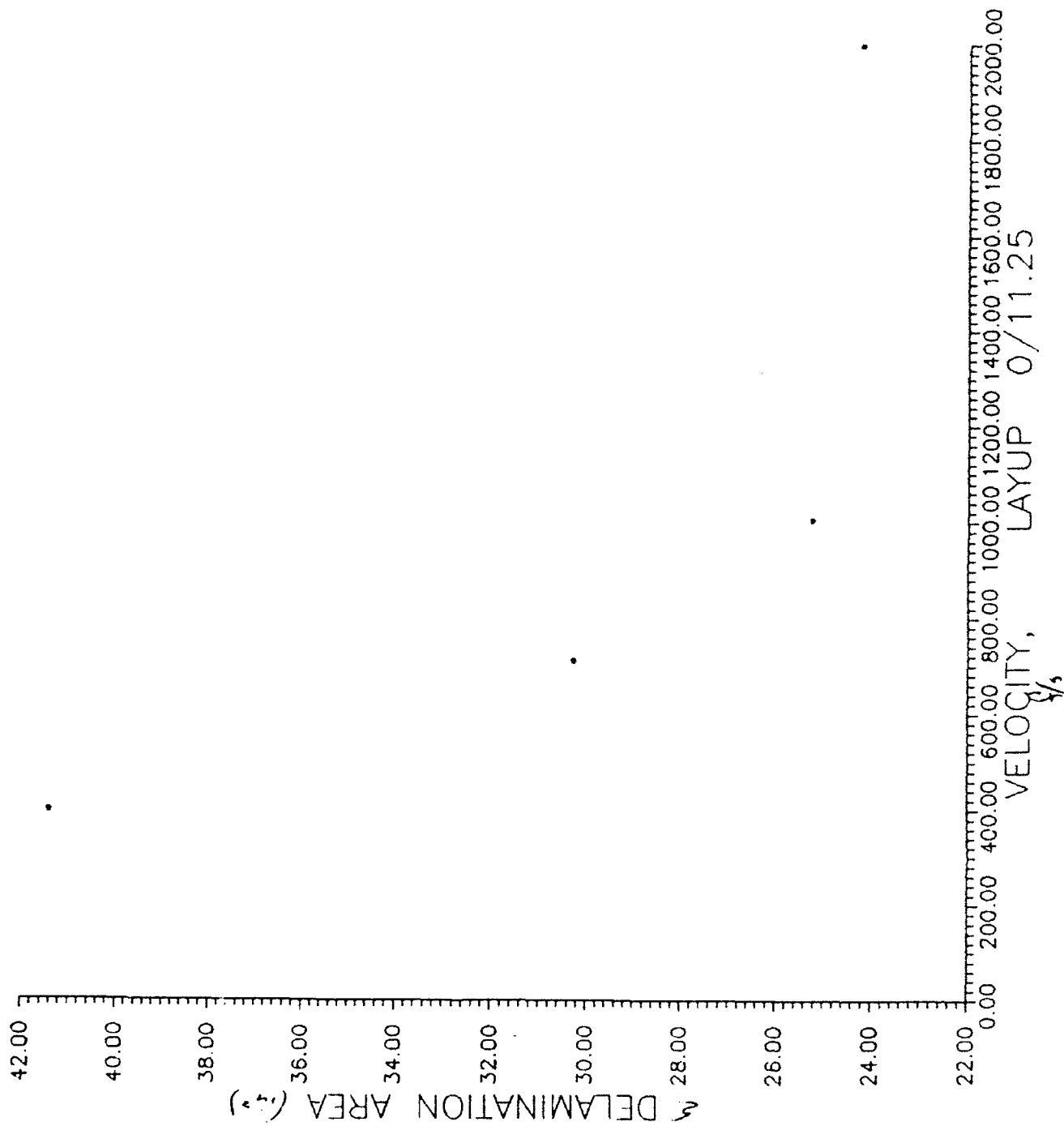


Figure 5 Interface Number, From Impact Face
Delamination and Hole Area of Individual Plies
For Panel Impacted at 696 Ft/Sec, 0/45 Layout









CONCLUSION

Because of very short time 12 weeks, only 120 panels were shot at various velocities. 16 panels with different layup deplied and the delamination area for each ply is analyzed for 16 panels. 35 were ~~g~~-scanned. Some of the results shown in graphs, on page 16 the graph shows the relation between total delamination area and velocity for the layup 0/90. Similarly on page 17, 18, and 19 graph shows the relation between total delamination area and velocity for layup 0/45, 0/22.5, and 0/11.25. Velocity at 400 ft/sec, the layup 0/11.25 absorbed more energy than layup 0/90 and other layups. On page 9 the graph shows the relation between the ballistic limit and the velocity. the ballistic limit for layup 0/11.25 is more than the other layups. The ballistic limit increases as the angle change between the adjacent ply decreases.

ACKNOWLEDGEMENTS

Working in wright laboratory is a learning experience of me specially working at WL/FIVS branch. First of all I would to thank to Air Force office of scientific Research for sponsor of this research. There are many individuals to whom I would like to thank my supervisor Greg Czarnaki, Range officer Ryon Studbeaker, sergeant Humble, and technician Mark Morgan. Also I would like to thank John T. Lair Graduate student of UNO for helping me from beginning to end of this research.

I would like to express my sincere gratitude to my Professor Dr. Hui David, and Dr. Arnold H. Mayer, who provide the technical background and work environment to make my research successful.

LOW VELOCITY IMPACT DAMAGE
OF COMPOSITE MATERIALS

Rob Slater
Graduate Teaching Assistant
Department of Mechanical, Industrial,
and Nuclear Engineering

University of Cincinnati
Cincinnati, OH 45221-0072

Final Report for:
Summer Research Program
Wright Laboratory

Sponsored by:
Air Force Office of Scientific Research
Wright-Patterson Air Force Base

September 1992

LOW VELOCITY IMPACT DAMAGE OF COMPOSITE MATERIALS

Rob Slater
Graduate Teaching Assistant
Department of Mechanical, Industrial,
and Nuclear Engineering
University of Cincinnati

Abstract

Damage of composite materials due to low velocity impact was studied. Tolerance to low velocity impact damage is a critical consideration in design of aircraft structures because the damage is very difficult to detect by visual inspection but may have severe effects on the residual strength and stiffness of the composite laminate. Low velocity impact is fundamentally different than ballistic impact associated with battle damage. A review of published literature has demonstrated that there is no acceptable theory available which can predict damage for a given set of impact parameters. The use of a discrete laminate plate theory has been suggested as a possible area for future research which may make damage prediction possible. Finite element modeling of impact events using commercial codes was also investigated. Several researchers have shown the capability to predict post-impact properties of composites if the damage state is known. If the state of a composite laminate subject to foreseeable impacts such as tool drop can be predicted, a set of design criteria for low velocity impact can be established.

Introduction

Impact damage of composite laminates is a concern as they gain wider use in aerospace structures. Low velocity impacts are particularly critical because they often produce significant internal damage which is difficult to detect visually. This is fundamentally different from high velocity impact as is commonly associated with battle damage.

Low velocity impact is not an absolutely defined term. It generally refers to the type of damage produced which is determined primarily by the characteristics such as geometry, material, and boundary conditions of the impacted laminate. Low velocity impact can also be defined as one in which the duration of impact is long compared to the time a stress wave takes to travel through the thickness of the target. This means the stress distributes through the laminate during the impact event. The impact generally excites the fundamental frequency of vibration of the target structure.

For a low velocity impact the stress waves propagate in-plane and reflect from the boundaries. Damage may occur at relatively large distances from the impact site, and manifests itself as mostly matrix cracking and delamination. There is usually little fiber breakage unless it occurs directly at the surface at the point of impact or in the back plies of thicker laminates due to bending.

As previously mentioned, the damage caused by low velocity impact is particularly insidious because it is very difficult to detect by visual inspection. The surface damage may be very slight, often only a small shallow dent, but the internal and back-ply damage can be severe. Conversely, a high velocity impact produces significant visible damage at the surface which may

extend deep into the target, but that damage tends to be highly localized. The effect of delamination and other damage associated with low velocity impact on strength and stiffness of panels is unpredictable unless some assessment of the damage can be made.

Researchers have proposed numerous models of low velocity impact damage of composite materials, but there is not general agreement as to what are the important phenomena of the damage mechanisms. A large number of parameters have been identified as possibly having an effect on impact behavior, and a great volume of experimentation has been performed. Unfortunately no standard testing regimes have been developed to guide researchers in setting up experiments and the data produced is highly dependent on the particulars of the test configuration. No reliable scaling laws exist to allow specific prediction of the dynamic behavior for an arbitrary panel or its post-impact damage state from a laboratory coupon test.

Little research has been done in the area of impact on composites carrying in-plane loads. This effect is important for aircraft structures because panels will almost always be carrying some load when impacts occur. Even for the commonly mentioned "tool drop" problem, for example, an airplane wing will be loaded by its own weight and that of any fuel inside it. It is important to investigate how pre-loading changes impact response and damage tolerance in order to design properly for foreseeable impact events.

Summary of Summer Research

The work conducted this summer was under the guidance of Dr. V.B. Venkayya, WL/FIBR. Approximately 60% of the 12 week period was spent identifying, locating, reading, and discussing published articles and books on the topics of impact and composite materials. The remainder of the time was spent learning the computer codes ASTROS and NASTRAN and performing finite element analyses to simulate low velocity impacts on composite materials.

A great deal of research on the topic of low velocity impact on composites has been performed, mostly over the last twenty years, and many articles have appeared in the scientific journals. Abrate [1] published a wide-ranging survey of research articles concerning impact on laminated composites. Much of this review is devoted to low velocity impact of plates. Force-indentation relationships, simulation of the impact event, studies of damage, effect of materials, geometry, and stacking sequence, damage prediction, and residual properties are topics which are covered at length and are of interest in the proposed research effort.

Determining an appropriate force-indentation relationship, or contact law, is a primary step in describing a plate's response to impact and predicting damage. The indentation can be described mathematically if the position of the target and impactor are known as functions of time. In experiments where both quantities are measured, a piecewise continuous contact law can be constructed [Sun & Grady,2]. But for analyses when these are not known a priori a contact law must be assumed. The classic Hertzian solution for elastic contact between a sphere and a semi-infinite body has frequently been used [Sun,3] and the results are generally accurate for small force levels. For impacts involving larger forces, Yang and Sun [4] proposed modifications to the Hertzian theory. The initial loading scheme is unchanged, but there is capability to handle

plastic deformation on unloading and subsequent reloading. These equations are based on quasi-static experiments performed by pressing steel balls into plates, so strain rate effects are not included. Further, several of the parameters necessary for application of the contact law must be determined experimentally. Due to the non-linearity of the equations it is difficult to estimate these constants without numerous tests.

Once a force-indentation law has been established, the next task is to determine how the plate responds to the loads applied by the impactor. Classical thin-plate theories such as Kirchhoff's do not apply well because of the transverse loading and the distribution of the transverse shear and normal strains through the thickness. Shear deformation effects must be included as they are significant [Schoeppner.5]. Also, membrane effects should be considered if the boundary conditions found on the aircraft structure being modeled indicate so. Kelkar, Elber, and Raju [6] developed a membrane-bending coupling theory for large deflection of circular plates, which gave more accurate results for impact-type point loads than the classical solutions. Sun [3] and Sun and Grady [2] investigated impact of balls on laminated composite beams. Finite element models were used to find the energy levels that cause permanent indentations and propagation of an internal crack, respectively.

These two finite element analyses modeled beams and neglected width effects which have a very significant effect. Many researchers have used 2-D plate and 3-D solid elements to model targets in impact simulations. Lee, Du, and Liebowitz [7] presented a 3-D dynamic analysis. The loading is ramp-on, ramp-off and the equilibrium equations are solved using the central difference method. This is an explicit method which is conditionally stable and requires a very small time step ($\Delta t < T/\pi$, where T is the smallest period of the system). For a model with many degrees of

freedom the highest natural frequency will be large (and period small) so it is very expensive to carry out the integration for even one full period at the fundamental frequency.

Another 3-D finite element analysis was performed by Wu and Springer [8]. They used 8-noded elements with incompatible displacement functions to provide more accurate information on bending stiffnesses and interlaminar shear stresses. Their goal was to determine stresses that cause damage and did not report how the dynamic response of the plate compared to experimental results. Sun and Chen [9] used 2-D plate elements in their finite element model and the modified Hertzian contact law proposed by Yang and Sun. The Newmark method is used for solving the equilibrium equations. The effects of biaxial pre-load (both tensile and compressive) and mass, size, and velocity of impactor were studied. They found that a tensile pre-stress effectively stiffens the plate, resulting in smaller maximum deflection and higher natural frequencies. This causes a shorter time period between the two target/impactor contact events which have also been observed by Wu and Springer and others. A compressive pre-load has exactly the opposite effect.

Impacts at higher velocity caused deflections of greater magnitude as would be expected, but also excited higher modes of vibration resulting in more rapid fluctuations in both contact force and plate displacement. Change in impactor mass causes a very fundamental change in the nature of the contact event. As previously mentioned, there is frequently a contact sequence in which contact is made, then broken as momentum is transferred from impactor to target, then re-established as the target rebounds back into the impactor. But as impactor mass increases, the initial contact time becomes longer, and at a certain mass level only a single contact is observed. If the contact time exceeds one-half the fundamental period of the target structure, the loading

may be considered to be no longer of the impact variety [Faupeil and Fisher, 10]. Dynamic effects are most important but wave propagation effects are not. For massive or flexible targets or massive impactors, the characteristics of impact noted in this paper may not apply. Sun and Chen noted size the impactor head had no effect on the response of the target. No experimental data was presented to determine the accuracy of these finite element simulations.

Palazotto, Perry, and Sandhu [11] compare impact damage in cylindrical composite panels to a dynamic analysis. Using experimentally determined force-time curves as loading for finite element models, the experiments were analyzed with shell elements incorporating shear deformation. Good results were obtained for deflections. The boundary conditions were ostensibly clamped, but the deflections indicate that the actual boundary conditions behave between simply supported and clamped.

Several papers have been published on the topic of testing methods with low velocity impacts. Sjoblom, Hartness, and Cordell [12] describe in detail an experimental setup using a pendulum-type impactor with a load cell for measuring contact force. They compare their facility to the more common drop tester and list some advantages and disadvantages of each. They also discuss some of the errors inherent in instrumented impact testing such as in velocity measurement and dynamic calibration of load cells.

A great amount of test data has been gathered and published by researchers but there is yet to be an accepted measure of impact performance that is independent the test method details [Robinson and Davies, 13]. The variations in impactor mass and size, impactor velocity, test specimen size and boundary, et cetera are so wide that no coordinating theory has been derived to predict the behavior for any given set of impact parameters.

An interesting analytical theory has been developed by Liu [14] which predicts the oft-seen lemniscular or "peanut" shaped delaminations associated with low velocity impact. This is based on the quantity known as the mismatch angle, which is the difference in fiber orientation between adjacent lamina in a laid-up composite structure. A mismatch coefficient is derived from the difference in bending stiffnesses of the upper and lower lamina at an interface which predicts the relative area and orientation of the delamination. For an interface where the adjoining layers are oriented in the same direction, no delamination is predicted, which is consistent with experimental evidence that delamination only exists where the fibers change orientation. The effects of material properties, laminate thickness, and impact energy are discussed.

The purpose of damage prediction is to be able to design aircraft structures which are tolerant to foreseeable impact events. Low velocity impacts an aircraft might experience include tool drop, hail, footsteps, and runway debris. If damage can be predicted, then the properties of the composite structure post-impact can be characterized also in order to estimate the residual strength and stiffness. Laminates may suffer a significant degradation of their properties due to low velocity impact, and to design structures that can survive such impacts, one must be able to predict the damage which is likely to occur.

The determination of residual properties has been undertaken by many investigators. The strength of damaged laminates, particularly in compression, is an important field of study. Since low velocity impact damage can extend such a relatively long distance from the impact site, strength can be reduced much more than for the case of a penetrating impact. Similarly stiffness can be change significantly due to back-ply damage.

Frequently there is a loss of symmetry in a laminate due to impact damage. This

introduces bending-stretching coupling. The vast majority of composite laminates are laid up symmetrically about the mid-plane in order to eliminate coupling of the bending and extensional strains. This greatly simplifies the analysis of such laminates. But when damage occurs the symmetry is lost and the behavior of the laminate may change drastically.

The effects of impact damage on tensile strength has been investigated by El-Zein and Reifsnider [15]. They hypothesized that residual strength is controlled by stress concentration effects in the immediate vicinity of the damaged region. A complex variable solution based on Lekhnitskii's problem of an anisotropic plate with an elliptical inclusion is derived, and the results agree fairly well with experiment.

The change in compression strength after impact is a more widely reported phenomenon. Dost, Ilcewicz, and Gosse presented a sublaminar stability based approach to predict damage tolerance (i.e. post-impact strength) [16]. This approach does require an accurate description of the state of the damage inside the plate. Further it was noted that there were significant differences in residual strength between experimental coupons and actual composite structures due to finite width effects in the test specimens.

Buckling of delaminated composites is a third failure mode. Global buckling is the same mode as occurs in undamaged panels, but may occur at markedly smaller load levels when damage is present due to loss of stiffness and the previously mentioned bending-stretching coupling. The phenomenon of local or delamination buckling is strictly related to delaminations near the free surface of the laminate. Under load the delaminated region may buckle while the rest of the laminate remains stable. Chai and Babcock [17] modeled an anisotropic layer separated from a thick isotropic base laminate. The delamination is elliptic in shape and the

material axes coincide with the ellipse's axes. The buckling for the damaged region by the Rayleigh-Ritz method and propagation of the delamination area is predicted via a fracture mechanics approach. Results show stable, unstable, or unstable growth with crack arrest depending on material properties and orientation, loading, and fracture energy.

Davidson [18] considered failure of a damaged laminate by all three failure modes: compression, global buckling, and delamination buckling. Buckling loads are calculated by applying the Trefftz criterion to governing equations found from the Rayleigh-Ritz method and compression failure by a modified maximum strain criterion. If the initial failure is delamination buckling, that layer is removed (it carries no load), the laminate properties are recalculated, and the loading sequence is continued until catastrophic failure (compressive or global buckling) is reached.

Davidson compared five analyses (two performed by himself and three reported by other researchers including Chai and Babcock) to experimental results. He found only one gave conservative predictions of buckling loads. The remaining four over-estimated the failure loads. The model which gave conservative results employed the reduced bending stiffness approximation. The $[D]$ matrix is replaced by a matrix $[D^*]$ defined as $[D^*] = [D] - [B][A]^{-1}[B]$ for cases where the coupling matrix $[B]$ is non-zero. The analysis is then performed as though the laminate was symmetric.

It was discovered that under certain conditions delamination buckling can occur under tensile loading. The fibers in a lamina normally have a small Poisson's ratio, approximately one order of magnitude smaller than the matrix material or a quasi-isotropic laminate. When the fibers in a delaminated layer are oriented normal to a tensile load, both the delamination and the

base undergo lateral contraction. Due to the mismatching of the Poisson's ratios, the delaminated region experiences compression and the base tension. Thus buckling may occur under loading cases when it is not expected.

Modeling of impact events using finite elements was another major thrust of the summer work effort. Many researchers have published data on the techniques they used to simulate low velocity impact. It is felt that finite element modeling and other numerical techniques are promising methods for finding solutions to impact problems. The commercial code NASTRAN was chosen as the analysis tool for this investigation, as it has extensive capabilities and the lab personnel are well-experienced in using it. The preliminary findings show qualitatively that the basic characteristics of the impact of a spherical impactor striking a simply-supported composite plate can be reproduced. But more work with the model is needed before the quantitative values derived agree with experiment.

Current design criteria for low velocity impact requires composite panels to be able to withstand an incident impact energy of 100 ft-lbf or a surface indentation of 0.1 inch. These are rather arbitrary guidelines which are too simple to govern a complex issue such as low velocity impact. An object dropped on an aircraft with 100 ft-lbf of energy could possibly cause damage over a broad spectrum. As an example, a 50 lbf tool box dropped from a height of two feet and landing flat on its bottom in the middle of a large panel would probably cause little damage. The impact energy would be spread over a large area, absorbed elastically, and dissipated by the internal damping. But if the tool box were to land on its corner or near an underlying stiffener, damage could be quit severe. The vagueness of the present design criteria means that composites may be designed with large factors of safety (and high weights) in order to guard against impacts.

Results of Summer Research

The work completed under the Summer Faculty/Graduate Student Program has demonstrated that low velocity impact damage of composite materials is an area which warrants further investigation. There are many issues pertaining to aircraft structures which remain unresolved, and there is no cohesive theory available to explain phenomena such as the dynamic response of a plate to impact loading, stresses which cause delamination and other internal damage, and residual properties.

A proposal to continue the work started in the summer program is to be submitted. The research that is proposed can be divided into three areas of concentration:

1. Development of an analytic model to describe behavior of impacted plates
2. Development of finite element modeling techniques for impact problems
3. Experimental studies of low velocity impact on pre-loaded plates

The first area has been researched heavily in the recent literature. Many theories have been brought forth. Most involve two-dimensional plate theories of varying complexity. Because of the transverse loading in impact problems, the transverse normal and shear stresses vary significantly through the thickness. Shear deformation is also important, so higher order shear deformation theories are necessary. It is believed that a discrete laminate theory is needed in order to determine the interlaminar stresses which result in delamination, so the research will concentrate on developing such a theory. These theories are already capable of handling in-plane loads, so adding the effects of pre-loading should be fairly straightforward.

Finite element methods represent an analysis tool capable of speeding modeling times and reducing the amount of expensive experimental testing on composite laminates. Developing techniques to simulate impact events using general-purpose finite element codes such as NASTRAN will make the complicated analyses more easily performed by designers. Other numerical solutions, such as the boundary element method, will be investigated. The laboratory has considerable computer resources and experience in finite element methods. Work done this summer shows promise of gaining valuable knowledge in this area.

The experimental portion of the research effort will serve as verification for the analytical and finite element models. A Dynatup drop tester and a 55 kip MTS testing machine are available for the experimental apparatus, and specimens will be prepared at the FIBC facilities. In addition, the tests will provide new information on the effect of pre-loading on damage. Some general knowledge on the dynamic behavior of pre-loaded plates is presented in the literature, but no experimental results concerning damage has been presented.

Conclusions

A review of published research on impact damage to composite materials has revealed several areas which would benefit from further study. Existing theories are limited in their ability to predict damage. Thus residual properties are very difficult to estimate, placing severe limitations on engineers' ability to design the most efficient composite laminates. Further investigation is currently proposed to determine a theory that will allow for accurately computing the stresses causing delamination and other damage due to low-velocity impact. Pre-loaded plates, which have received little attention previously, will be considered. Finite element modeling techniques will be developed in order to provide a readily available design and analysis tool for engineers. Experimental testing will be done to verify the analytical theory and finite element results and also to provide data on the effect of pre-loading on impact damage.

References

1. Abrate, S., "Impact on Laminated Composite Materials", Applied Mechanics Reviews, Vol. 44, No. 4, 1991, pp. 155-190
2. Sun, C.T., and Grady, J.E., "Dynamic Delamination Fracture Toughness of a Graphite/Epoxy Laminate Under Impact", Composites Science and Technology, Vol. 31, 1988, pp. 55-72
3. Sun, C.T., "An Analytical Method for Evaluation of Impact Damage Energy of Laminated Composites", Composite Materials: Testing and Design (Fourth Conference), ASTM 617, American Society for Testing and Materials, 1977, pp. 427-440
4. Yang, S.H., and Sun, C.T., "Indentation Law for Composite Laminates", ASTM STP 787, I.M. Daniel, Ed., 1982, pp. 425-449
5. Schoeppner, G.A., "A Stress-Based Theory Describing the Dynamic Behavior of Laminated Plates", Ph.D. Dissertation, The Ohio State University, 1991
6. Kelkar, A., Elber, W., and Raju, I.S., "Large Deflection Behavior of Quasi-Isotropic Laminates Under Low Velocity Impact-Type Point Loading", 26th Structures, Structural Dynamics, and Materials Conference, Orlando, FL, Part 1, 1985, pp. 432-441
7. Lee, J.D., Du, S., and Liebowitz, H., "Three-Dimensional Finite Element and Dynamic Analysis of Composite Laminate Subjected To Impact", Computers and Structures, Vol. 19, No. 5/6, 1984, pp. 807-813
8. Wu, H.T., and Springer, G.S., "Impact Induced Stresses, Strains, and Delaminations in Composite Plates", Journal of Composite Materials, Vol. 22, June 1988, pp. 533-560

9. Sun, C.T. and Chen, J.K., "On the Impact of Initially Stressed Composite Laminates", *Journal of Composite Materials*, Vol. 19, November 1985, pp. 490-504
10. Faupel, J.H., and Fisher, F.E., Engineering Design, John Wiley and Sons, 1981
11. Palazotto, A., Perry, R., and Sandhu, R., "Impact Response of Graphite/Epoxy Cylindrical Panels", *AIAA Journal*, Vol. 30, No. 7, July 1992, pp. 1827-1832
12. Sjoblom, P.O., Hartness, J.T., and Cordell, T.M., "On Low-Velocity Impact Testing of Composite Materials", *Journal of Composite Materials*, Vol. 22, January 1988, pp. 30-52
13. Robinson, P., and Davies, G.A.O., "Impactor Mass and Specimen Geometry Effects in Low Velocity Impact of Laminated Composites", *International Journal of Impact Engineering*, Vol. 12, No. 2, 1992, pp. 189-207
14. Liu, D., "Impact-Induced Delamination—A View of Bending Stiffness Mismatching", *Journal of Composite Materials*, Vol. 22, July 1988, pp. 674-692
15. El-Zein, M.S., and Reifsnider, K.L., "On the Prediction of Tensile Strength after Impact of Composite Laminates", *Journal of Composites Technology and Research*, Vol. 12, No. 3, 1990, pp. 147-154
16. Dost, E.F., Ilcewicz, L.B., and Gosse, J.H., "Sublaminar Stability Based Modeling of Impact-Damaged Composite Laminates", *Proceedings of the American Society for Composites, 3rd Technical Conference, Seattle, WA, 1988*, pp. 354-363
17. Chai, H., and Babcock, C.D., "Two-Dimensional Modeling of Compressive Failure in Delaminated Laminates", *Journal of Composite Materials*, Vol. 19, January 1985, pp. 67-98
18. Davidson, B.D., "A Determination of the Strength and Mode of Failure of Compression Loaded Laminates Containing Multiple Delaminations", *JPL Document D6447*, September 1989

**MONITORING OF DAMAGE ACCUMULATION FOR THE PREDICTION OF FATIGUE
LIFETIME OF CORD-RUBBER COMPOSITES**

**Jeffrey A. Smith
Graduate Student
Department of Engineering Science and Mechanics**

**The Pennsylvania State University
227 Hammond Building
University Park, PA 16802**

**Final Report for:
AFOSR Summer Research Program
Wright-Patterson Air Force Base
Flight Dynamics Laboratory
Vehicle Subsystems Division
Aircraft Launch and Recovery Branch**

**Sponsored by:
Air Force Office of Scientific Research**

September, 1992

MONITORING OF DAMAGE ACCUMULATION FOR THE PREDICTION OF FATIGUE LIFETIME OF CORD-RUBBER COMPOSITES

Jeffrey A. Smith, Graduate Student
Department of Engineering Science and Mechanics
The Pennsylvania State University

Abstract

This study attempted to monitor fatigue damage in angle plied cord-rubber composites by measuring dynamic creep, temperature changes, and acoustic emission (AE). In addition to real-time monitoring damage accumulation, non-destructive evaluation of damaged specimens was performed using ultrasonic C-scan and x-ray techniques. Two types of composites were used: steel cord-reinforced model composites and nylon cord-reinforced composites representing the actual aircraft tire carcass. Results showed that the rate of AE changes when debonding and delamination appear. However, the location of the damage with respect to the sensor seemed to effect the intensity and the accumulation rate of AE. Dynamic creep was found to be a good indicator of the remaining fatigue lifetime of composite coupons. X-ray testing clearly showed the areas of delamination in model composites. In the case of nylon cord composites, the effect of frequency was examined in detail with a special attention to temperature rise characteristics. The relationship between temperature and fatigue life was not straightforward. On the other hand, a power law relationship between frequency and fatigue life could be established. Hysteresis curves show that energy dissipation per cycle did not change for specimens tested at different frequencies, meaning that energy dissipation per unit time is proportional to frequency.

MONITORING OF DAMAGE ACCUMULATION FOR THE PREDICTION OF FATIGUE LIFETIME OF CORD-RUBBER COMPOSITES

Jeffrey A. Smith

I. INTRODUCTION

Currently, the operational life of bias-ply aircraft tires used by the U.S. Air Force is determined by costly dynamometer tests which are designed to determine the number of takeoff and landing cycles a tire may endure before replacement (1). This method of accelerated testing provides valuable information on the durability and life expectancy of aircraft tires. However, dynamometer testing often reflects merely the sensitivity of a particular tire design to a given set of test conditions, unless the mechanisms of degradation and structural failure are identified.

An initiative is underway to develop more cost-effective methods to predict aircraft tire lifetimes. The initial goal is to define the fatigue fracture mechanisms of the tire carcass by identifying the parameters which control the process of damage accumulation of cord-rubber composite elements. These efforts could extend tire replacement intervals which could prove more cost-effective than existing criteria that now govern tire replacement.

The focus of this study addresses the above-stated goal of identifying the parameters which control the fatigue life of angle-ply tire carcass composites, with a special provision of establishing a viable experimental technique for real-time monitoring of the damage accumulation process. Real-time monitoring of the damage accumulation process in composite specimens is necessary for the fatigue fracture mechanism study. Additionally, when established it will also lay a technical basis for the future research activities of predicting the imminent failure of an actual tire carcass in the field.

Past studies (2-7) have shown that angle-ply cord-rubber composites exhibit an unusually high level of interply shear strain under uniaxial tension. Above a critical value of interply shear strain, localized damage is induced in the form of cord-matrix debonding. Further

increase of strain results in the initiation of matrix cracking. Under static loading, debonding and matrix cracking become more extensive and eventually develop into the delamination leading to the gross fracture of the composites.

The same sequence of failure modes was observed in cord-rubber composites under cyclic loading as long as the minimum stress remains in the tensile regime (2). It was also observed that damage accumulation in the form of debonding and delamination is accompanied by a steady increase in temperature and local strain. The first phase of the present study involves using these values to monitor and predict damage in cord-rubber composites. The damaged samples were subsequently studied by visual inspection, ultrasonic C-scan and an X-ray method.

For real-time monitoring of damage accumulation, acoustic emission (AE) analysis was utilized in addition to the measurement of changes in strain and temperature. The damage accumulation under cyclic as well as static tensile loading is normally accompanied by AE events that are associated with the sudden release of energy. Past research has attempted to use this method to monitor damage and associate an acoustic signature with a particular damage mode (8). Since the catastrophic failure of the tire carcass is often caused by delamination, there is great potential for the use of AE as a means of sensing imminent tire failures if the acoustic signature of the delamination mode is identified. However, little work has been done with cord-rubber composites using AE analysis, even though it is a technique which has been used for years with other structural materials such as metals and advanced fiber composites (9-11). This research could provide a NDE tool which could detect the failure mechanisms responsible for controlling the fatigue lifetime of cord-rubber composites and ultimately aircraft tires.

Previous research work (2) has defined S-N curves for cord-rubber composites representing aircraft tire carcass. These tests have shown a clear relation between the stress amplitude and the number of cycles to failure, but only study the effect of relatively low frequencies, between 1 and 10 Hz, on fatigue lifetime. The current study will assess the effect of higher frequencies of up to 30 Hz on these composites. Load, displacement, and temperature

were recorded versus time. From this data, hysteresis curves were produced to quantify the amount of heat generated by the coupon during testing, which may affect the fatigue life. Displacement time history curves may also be used to predict fatigue damage.

II. Objectives

The present investigation was undertaken in accordance with continuing efforts for the laboratory prediction of the durability of aircraft tire carcass. Two objectives were set forth. The first goal was to identify the parameters controlling damage accumulation and fatigue lifetime of steel-reinforced rubber model composites, and the second was to determine the effects of frequency on the fatigue lifetime of tire carcass composites.

III. Experiments

The first phase of this study was to monitor and predict damage in cord-rubber composites. For practical considerations, model composites consisting of 2+2x0.25 steel wire cable reinforcement (four filament) and a proprietary rubber matrix were used instead of nylon reinforced composites representing actual aircraft tire carcass. The model composite was included to allow better determination of the failure modes, since the chosen cord angle (+/- 23°) maximizes interply shear strain which is a major contributing factor in composite failure. Laminate panels (4" gage length x 3" wide x .25" thick) were supplied by the Goodyear Tire & Rubber Company. To avoid tension-bending coupling, the laminates were constructed with a symmetric ply lay-up. End tabs were added to the specimens to prevent grip failures during testing. Coupon specimens of 0.75" width were machined from these panels and edge polished on a grinding wheel.

The load amplitude and mean load for these tests remained constant at 200lb and 400lb respectively, and the load function was applied uniaxially as a sine wave at a 5 Hz frequency. During testing, strain, temperature, and acoustic emissions were recorded versus time. Temperature was measured with a thermocouple attached by a metal clip to the middle of the surface of the coupon, and strain was measured by recording the position of the actuator during testing.

Acoustic emissions were measured with a LOCAN-AT system (Physical Acoustics Corporation). Signal Processing parameters were based on a past study (8). One Physical Acoustics R50 transducer with a 1.25cm diameter was connected to a wide bandpass preamplifier (Physical Acoustics model 1220A) to amplify the signal. The following internal signal processing parameters were used:

Gain	20 dB
Amplitude Threshold	45 dB
Peak Definition Time (PDT)	20 ms
Hit Definition Time (HDT)	150 ms
Hit Lockout Time (HLT)	300 ms

A thin layer of vacuum grease was applied between the transducer and the coupon to insure good contact and minimize wave reflection at the interface, and the transducer was held in place near the center of the sample with adhesive tape and rubber bands.

The fatigue life of the composites under these conditions was first determined by testing 5 samples until failure. Subsequent tests were performed for various fractions of the fatigue lifetime to produce systematic damage in the coupons. These specimens were then examined using C-scan and X-ray non-destructive examination. The residual strength and stiffness was determined by monotonic loading at a rate of 0.002 in/sec.

During the second phase of testing, the effect of higher frequency on aircraft tire cord composites was examined. load amplitude and mean load were held constant at 60 lbs and 90 lbs, respectively. The loading frequency was varied between 2 and 30 Hz, the minimum and maximum loads were 30 and 150 lbs respectively, and all samples were tested to failure. Load

and displacement were measured at small time intervals during testing; typical data acquisition consisted of measuring load and displacement at 20 points per cycle for 4 cycles, once every 1000 cycles. This data was recorded by the computer software which controlled the testing system (Materials Testing System model 810).

In addition to load and displacement data, the temperature was monitored using thermography equipment (Inframetrics model 525). The temperature was taken across the width of the specimen, and the image of the profile was recorded on a video cassette recorder. All subsequent temperature readings will represent the peak value of the temperature profile at a given time.

IV. Results and Discussion

The first phase of this study involves monitoring damage accumulation of model composites during fatigue testing. Testing conditions and measured values relating to the strain and temperature are given in table 1. Creep rate and temperature rise rate are measured in the steady state region of the testing period. Figure 1 shows typical displacement and temperature versus time plots, and the steady state regions associated with testing. Dynamic creep (figure 1) was defined as the difference between the final displacement prior to failure and an initial extrapolated strain value. The extrapolated value was determined by the intersection of the extrapolation of the steady-state region line and the onset of testing (see figure 1).

Past research (8) has shown that the rate and extent of dynamic creep and temperature rise rate are good indicators of fatigue life. Figure 2 shows that dynamic creep may be a good indicator of damage accumulation in coupons tested at one frequency. This graph plots dynamic creep versus percentage of life, which is defined as the number of cycles divided by the average fatigue life for unfailed samples, and provides a curve to predict the number of cycles to failure for an unfailed specimen.

indicator of damage accumulation in coupons tested at one frequency. This graph plots dynamic creep versus percentage of life, which is defined as the number of cycles divided by the average fatigue life for unfailed samples, and provides a curve to predict the number of cycles to failure for an unfailed specimen.

Acoustic emission was used for in-situ monitoring of the damage accumulation process. Under cyclic tension, AE occurs as soon as the loading starts. Past research has attempted to assign particular rates of accumulation of AE hits, counts, or energy to specific failure modes such as debonding and delamination. The present study found that the rate of AE activity changes when debonding, cracking, and delamination appear. However, there was difficulty assigning a particular rate to a specific damage mode. This study found that the location of the damage with respect to the sensor affected the strength of the AE signal and the rate of accumulation. Therefore, the rate of accumulation was affected by both location of the damage and the damage mode. Future work may focus on the use of two or more sensors to determine the location of the damage and its relation to AE rate. This would assist in developing a more comprehensive method of determining the failure mode from acoustic emission data.

Acoustic emission was also monitored during monotonic testing of damaged and undamaged steel cord rubber model composites. The tests showed that Kaiser's rule applies to this system. Kaiser's rule states that a sample which is subsequently loaded and unloaded will not produce emissions upon reloading until the maximum stress from the previous loading is exceeded. Undamaged coupons subjected to a constant strain rate until failure show emissions uniformly throughout the test. However, samples which underwent cyclic testing showed acoustic emissions only after the maximum cyclic stress was exceeded. Examples of AE signals from tests of both damaged and undamaged samples are given in figure 3.

Coupons were examined before monotonic testing using C-scan and x-ray methods to determine the amount of damage caused by cyclic loading. Examples of results are given in figure 4. C-scan results were inconsistent and difficult to characterize. However, the extent of

coupon damage was clearly evident in x-ray images of the specimens. Although a limited number of samples were tested, images of coupons cycled for increasingly longer times clearly showed successively larger areas of delamination. Future work may attempt to correlate the observed delamination area with residual strength or remaining cycles to failure.

The second phase of the study involved testing nylon cord-reinforced composite coupons at different frequencies under identical load conditions. Table 2 provides information on the testing conditions and results for these composites.

Force and displacement data recorded during testing allowed for the generation of hysteresis loops for individual cycles. The area within the loops was used to determine the amount of energy loss in a sample during a particular cycle and the value of $\tan \delta$, where δ is the phase angle difference between stress and strain. $\tan \delta$, also known as the loss tangent, is important since it is a fundamental parameter for expressing energy losses relative to energy stored in a system. This value may be expressed as:

$$\tan \delta = E''/E' \quad (1)$$

where E'' and E' are the imaginary and real components of the elastic modulus. This may also be calculated from the equation:

$$\tan \delta = \tan (\sin^{-1}(A/(\pi\sigma_o\epsilon_o))) \quad (2)$$

where A is the area inside the hysteresis loop, and σ_o and ϵ_o are the stress and strain amplitudes, respectively. Values of hysteresis loop areas and $\tan \delta$ are given in table 3 for samples with different frequencies, at regular intervals in the fatigue life. The numbers are relatively constant, with slightly higher values of $\tan \delta$ at lower frequencies. According to Gehman (12), $\tan \delta$ should be found in the range of 0.1 to 0.2. Dodge and Clark also report (13) that $\tan \delta$ should decrease slightly with increasing frequency. The values obtained from this experiment correlated well with past research and theory.

The area within the hysteresis loops was also relatively constant, with values ranging between 3.9 and 5.5 lb•in. The loop area generally increased with time, however some samples

did not support this trend. From this data, it was determined that since the energy dissipation per cycle was approximately constant and higher frequencies cause more dissipation per unit time, the energy dissipated per unit time was proportional to the frequency. Most of the energy lost is assumed to be produced as heat energy within the sample. Since thermal degradation is thought to contribute to tire failure, fatigue lifetime was plotted versus frequency on a logarithmic scale in figure 4. The least squares curve fit assumes a power law relation with an exponent of nearly -1, which indicates that fatigue life and frequency, and thus energy loss per unit time, are inversely proportional.

To determine the frequency effect on temperature, the peak value of the temperature profile across the coupon was plotted versus time. A typical graph is given in figure 5. Generally, the temperature increased during the initial part of the test until it reached a steady state, and rose again at the onset of severe damage near the end of the test until failure occurred. As the frequency decreased, the length of the steady state region increased. Several temperature profiles are given in figure 6. The temperature of each coupon is plotted versus the normalized time for each test to show the variation of temperature over fatigue life for several samples. The normalized time is defined as the time at a given point divided by the failure time for a particular sample. This graph shows that the average temperature increases with frequency. However, the difference between samples at 20 and 30 Hz is very small. The coupons at these frequencies appeared to undergo a phase change which limited the maximum temperature of the specimen.

The surface temperature of the specimen at a particular frequency can be estimated from the data generated in this study. Although temperature seems to have an effect on the fatigue life of cord rubber composites, there is a much clearer correlation between fatigue life and frequency. It should be noted that surface temperature is a dependent variable which may not be directly proportional to internal temperature or heat dissipation rate which may be frequency related.

V. Conclusions

This study continues an ongoing effort to establish a viable technique for in-situ monitoring of the damage accumulation process of cord-rubber composites. Previous work has examined the effect of a broad range of loading parameters on the fatigue life of these materials (8). The first part of the present study maintains constant fatigue loading conditions in order to identify key parameters which may predict the life of the composite. Among them, dynamic creep was found to provide the most accurate measure of the fatigue life. Acoustic emission may also provide a method of determining the specific damage mode during fatigue. However, further study with the use of two or more transducers should be performed to determine the effect of location on signal rate of accumulation. Additionally, x-ray techniques may prove to be a reliable non-destructive method of characterizing damage in cord-rubber composites.

The second part of this study attempted to determine the effect of frequency on nylon cord tire carcass composites. Temperature, force and displacement were measured versus time, and samples tested at different frequencies were compared. Hysteresis loops were generated, and the loss tangent was calculated to be consistent with past research. Loop area was found to be similar for different frequencies which means that the energy loss rate, thought to be mostly heat energy, is directly proportional to frequency. Since thermal degradation is thought to be a very important factor in the fatigue life of cord-rubber composites, lifetime was plotted versus frequency and an inverse relationship was demonstrated.

Coupon surface temperature was compared for various frequencies. Higher frequencies produced higher steady state and failure temperatures. However, the correlation between temperature and fatigue life was not clear. Additional samples are currently being tested, and should provide a more statistically significant database to explore trends between these values.

TABLE 1: Fatigue-Induced Changes in Strain and Temperature for Steel Cord-Rubber Composites (5 Hz Frequency)

Number of Cycles (N_f)	Percent of Lifetime ($\%N_f$)	Creep (in)	Creep Rate (in/Kcycle)	Final Temperature ($^{\circ}\text{F}$)	Temperature Rise Rate ($^{\circ}\text{F/Kcycle}$)
64450	100	.291	2.07×10^3	88	-----
53564	100	.250	2.05×10^3	93	.010
73127	100	.250	1.40×10^3	94	.050
53710	100	.250	1.95×10^3	94	.085
43856	100	.220	2.05×10^3	95	.055
68149	95	.160	1.70×10^3	98	.045
50100	87	.115	2.00×10^3	96	.088
50100	87	.095	1.70×10^3	95	.059
50100	87	.118	2.00×10^3	94	.045
38000	66	.071	1.87×10^3	93	.045
38000	66	.075	1.58×10^3	96	.010
38000	66	.060	1.58×10^3	93	.032
19000	33	.035	1.84×10^3	94	0.00
19000	33	.040	2.11×10^3	93	0.00
19000	33	.035	1.84×10^3	94	.107

TABLE 2: Testing Conditions and Results for Nylon Cord Rubber Composites

Sample #	Frequency (Hz)	Cycles to Failure (N_f)	Final Temperature ($^{\circ}$ F)	Dynamic Creep (in)	Dynamic Creep Rate (in/kcycle)
G-1	30	13242	-----	.22	.011
E-1	30	13603	-----		
F-2	30	14550	346.1		
G-3	30	14929	360.1		
G-2	20	17658	-----		
E-3	20	13762	323.4	.27	.013
G-6	20	28892	-----		
F-3	10	28573	311.2		
E-5	10	44896	-----		
F-5	10	37772	-----	.35	.0025
G-5	5	81107	-----	.41	.0013
E-6	5	81832	200.2	.38	.0016
F-6	2	160140	179.0	.30	.0010

TABLE 3: Hysteresis Loop Area and Tan δ Values for Samples at Different Frequencies. Loop area units are in lb•in.

Sample #	Frequency	25% area, tan δ	50% area, tan δ	75% area, tan δ	99% area, tan δ
F-6	2 Hz	4.09, 0.128	4.19, 0.128	4.68, 0.133	5.24, 0.141
E-6	5 Hz	3.88, 0.119	4.02, 0.120	4.05, 0.121	5.49, 0.118
F-5	10 Hz	4.13, 0.117	4.12, 0.117	4.08, 0.113	4.49, 0.109
G-6	20 Hz	4.02, 0.117	3.91, 0.112	4.04, 0.114	4.57, 0.114
G-1	30 Hz	4.89, 0.118	5.29, 0.120	5.17, 0.120	4.81, 0.104

FIGURE 1: Typical Displacement and Temperature versus Time Graphs

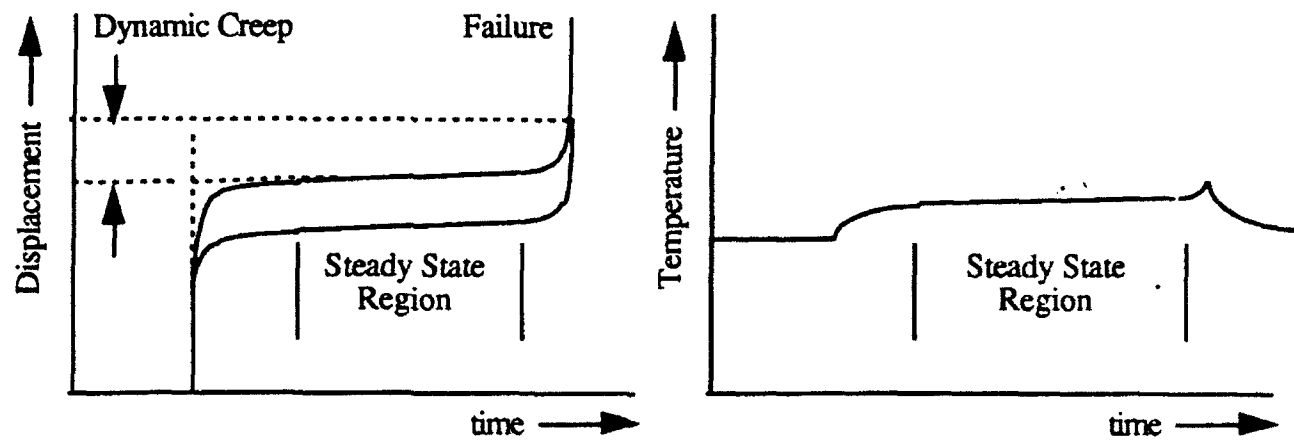


FIGURE 2: Dynamic Creep vs. Percent of Lifetime for Steel Cord Composites

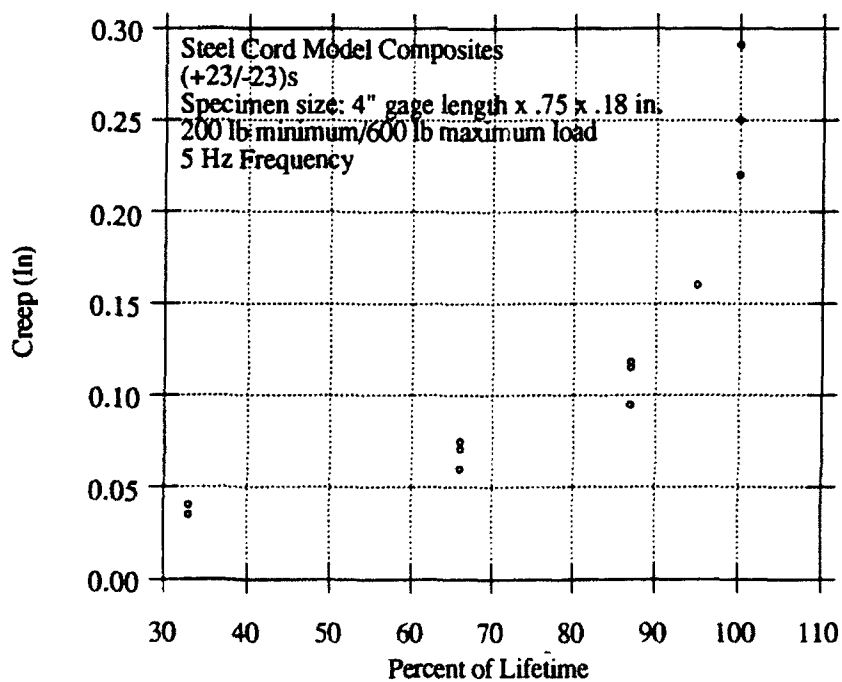


FIGURE 3.1: Cumulative Events Versus Time for Cyclic Testing to Failure. The onset of different failure mechanisms is demonstrated by a change in the slope of the graph.

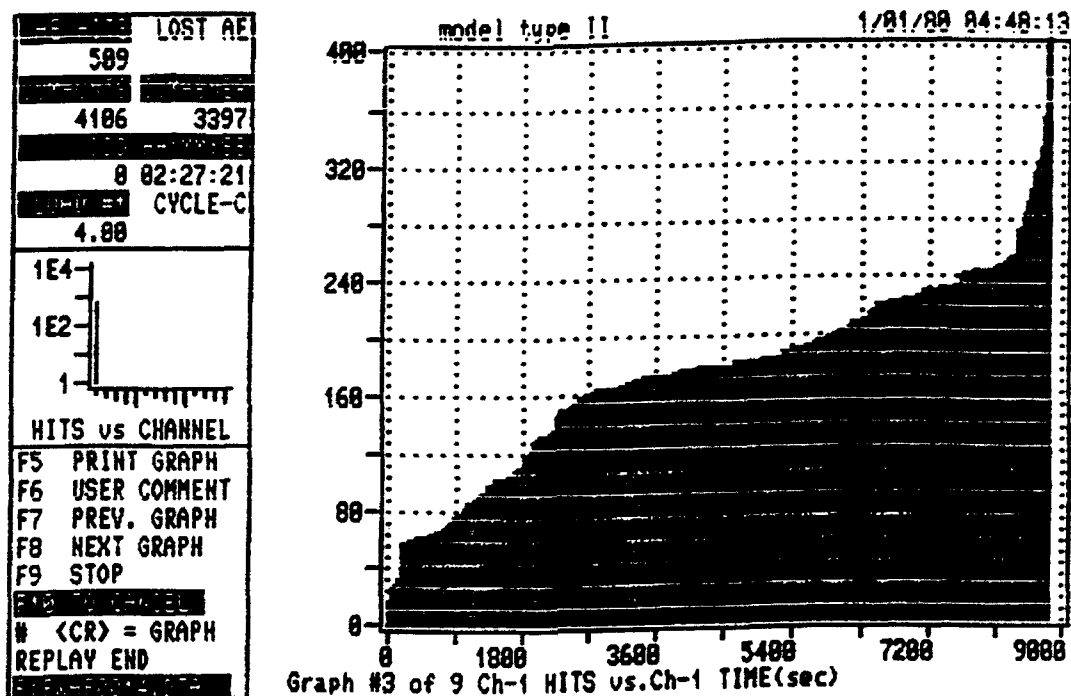


FIGURE 3.2: Cumulative Events Versus Time for Undamaged Sample. Monotonic testing performed to failure on virgin sample at a constant strain rate of 0.002 in/sec.

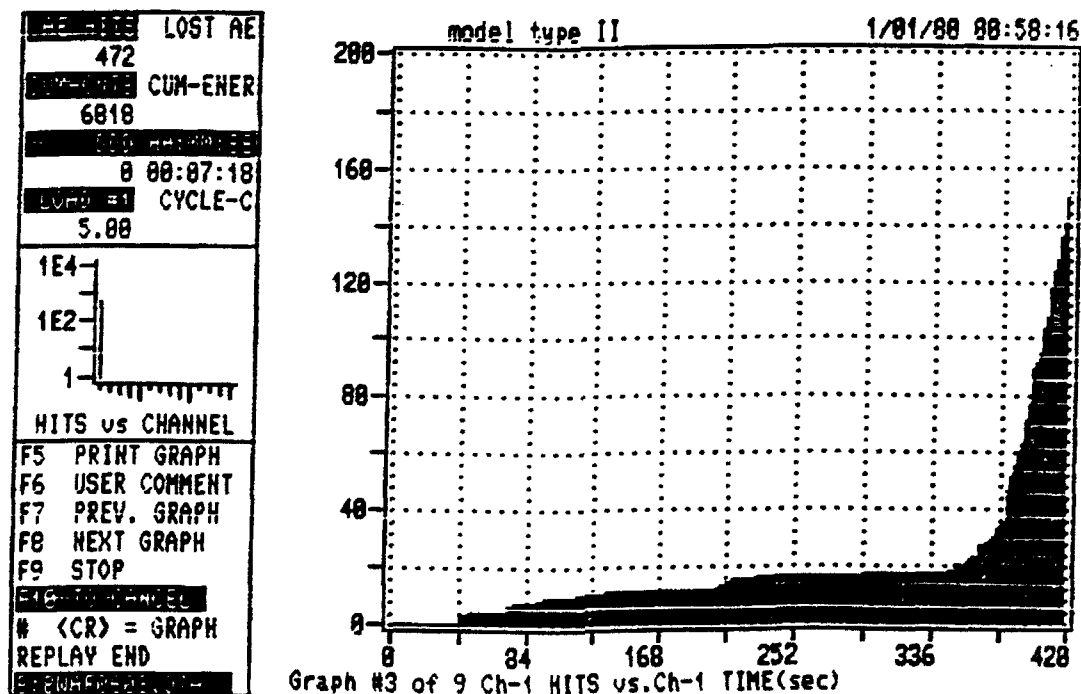


FIGURE 3.3: Cumulative Events Versus Time for Damaged Sample. Monotonic testing performed on sample tested for 38,000 cycles ($N_f = 66\%$). Acoustic emissions are detected only above the maximum cyclic testing load of 600 lb.

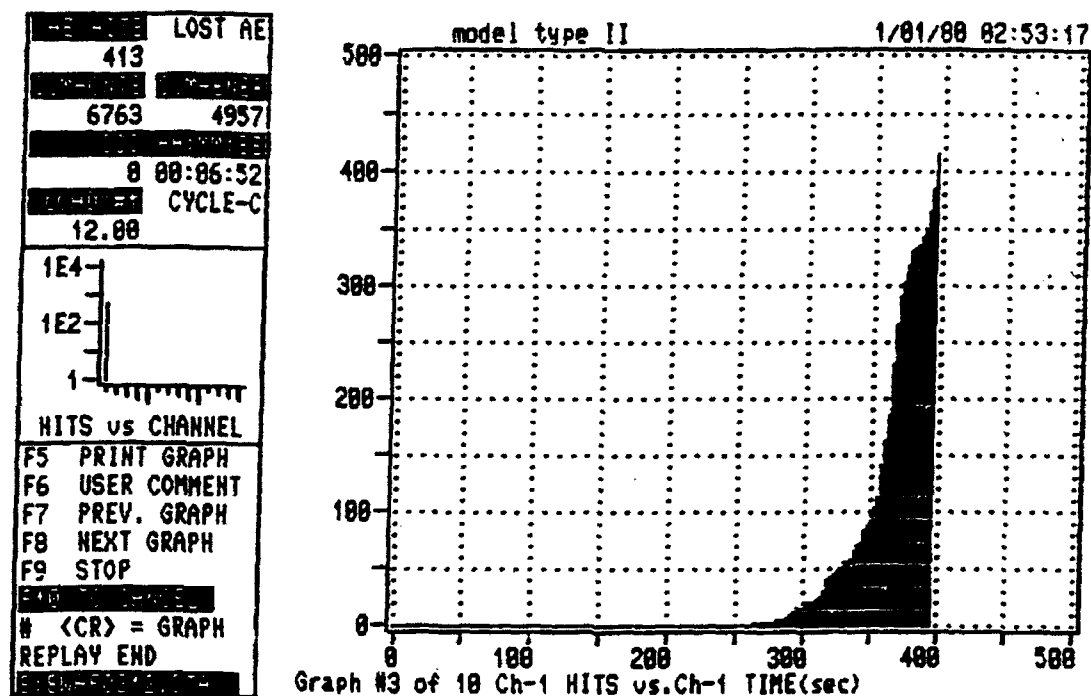
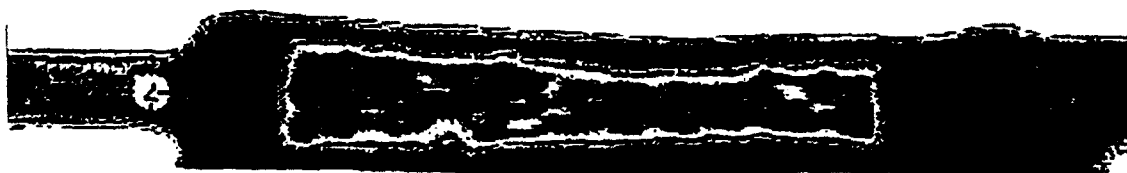


FIGURE 4a: Ultrasonic C-scan images of undamaged and damaged coupons

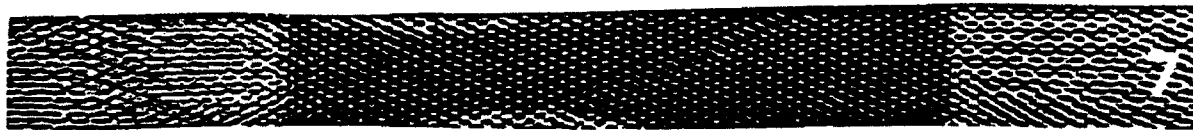


C-scan of damaged coupon (51100 cycles)



C-scan of undamaged coupon

FIGURE 4b: X-ray images of undamaged and damaged coupons



X-ray of damaged coupon (51100 cycles)



X-ray of undamaged coupon

FIGURE 4: Cycles to Failure vs. Frequency

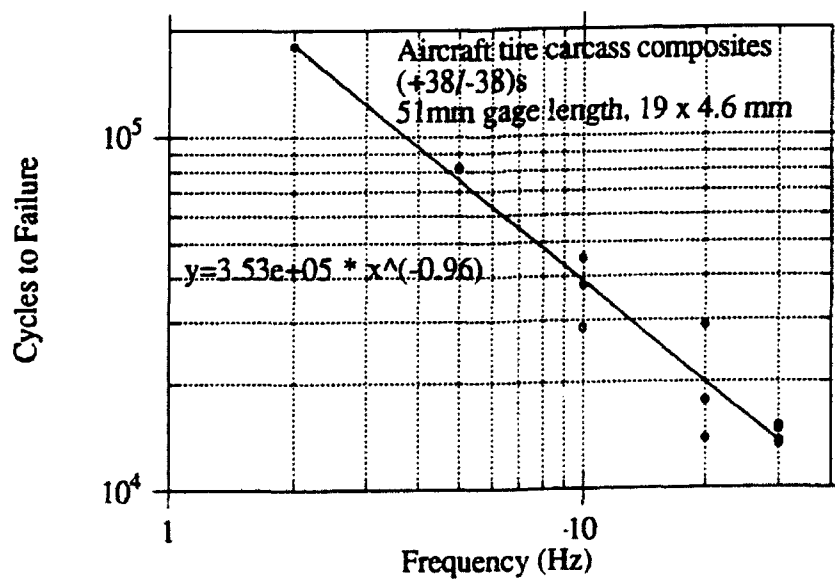


FIGURE 6: Temperature vs. Time For a Typical Nylon Cord Sample

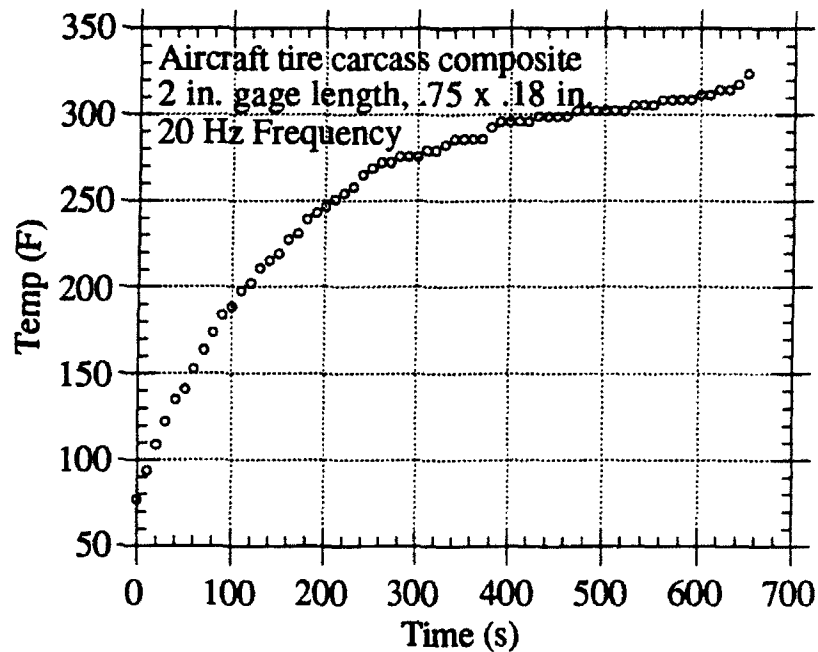
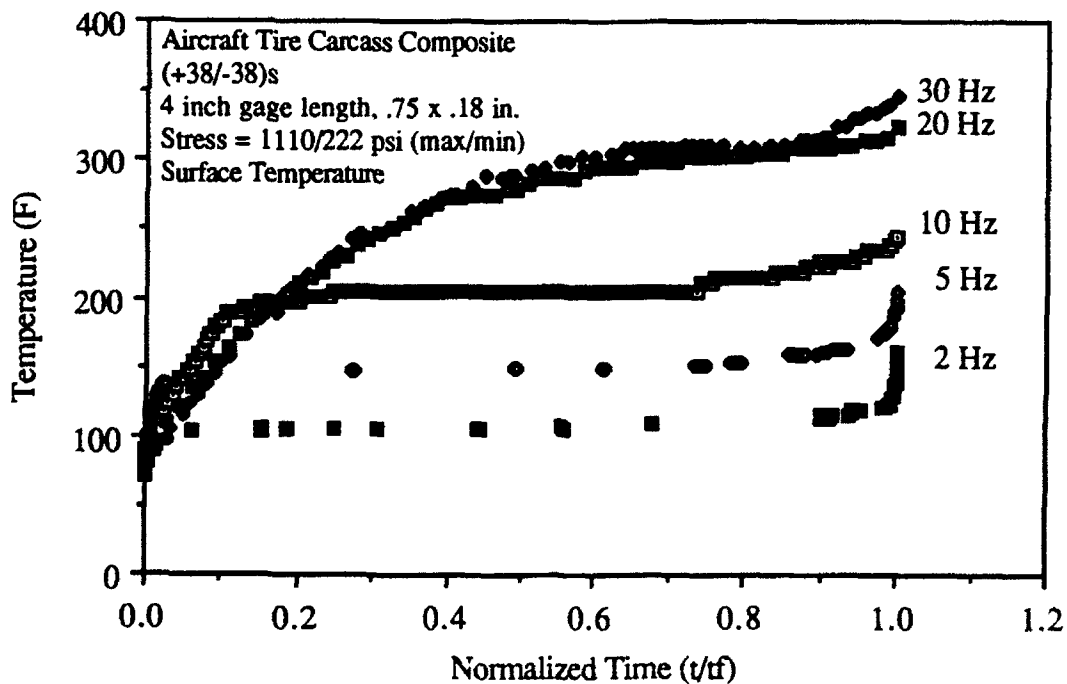


FIGURE 7: Temperature vs. Normalized Time for Various Frequencies



REFERENCES

1. S. N. Bobo, "Fatigue Life of Aircraft Tires," *Tire Science and Technology*, Vol. 16, No. 4, p. 208 (1988).
2. B. L. Lee, J. P. Medzorian, P. M. Fourspring, G. J. Migut, M. H. Champion, P. M. Wagner and P. C. Ulrich, "Study of Fracture Behavior of Cord-Rubber Composites for Lab Prediction of Aircraft Tire Durability," Soc. of Automotive Engineers Paper #901907, Warrendale, PA (1990).
3. R. F. Breidenbach and G. J. Lake, "Mechanics of Fracture in Two-Ply Laminates," *Rubber Chemistry and Technology*, Vol. 52, p. 96 (1979).
4. R. F. Breidenbach and G. J. Lake, "Application of Fracture Mechanics to Rubber Articles Including Tyres," *Philosophical Trans. Royal Soc. London*, Vol A299, p. 189 (1981).
5. J. D. Walter, "Cord-Reinforced Rubber" in Mechanics of Pneumatic Tires edited by S. K. Clark, U.S. Department of Transportation, Washington D.C. (1982).
6. J. L. Ford, H. P. Patel and J. L. Turner, "Interlaminar Shear Effects in Cord-Rubber Composites," *Fiber Science and Technology*, Vol. 17, p. 255 (1982).
7. R. J. Cembrola and T. J. Dudek, "Cord/Rubber Material Properties," *Rubber Chemistry and Technology*, Vol. 58, p. 830 (1985).
8. Byung-Lip Lee, "Study of Fracture Behavior of Cord-Rubber Composites for Lab Prediction of Structural Durability of Aircraft Tires," 1990 USAF-RDL Summer Faculty Research Program, September 30, 1991.
9. A. Rotem, "Effect of Strain Rate on Acoustic Emission from Fibre Composites," *Composites*, Vol. 9, No. 1, p. 33 (1978).
10. B. Harris, F. J. Guild and C. R. Brown, "Accumulation of Damage in GRP Laminate," *J. Physics D:Applied Physics*, Vol. 12, p. 1385 (1979).
11. L. Carlsson and B. Norrbom, "Acoustic Emission from Graphite/Epoxy Composite Laminates with Special Reference to Delamination," *J. Materials Science*, Vol. 18 p. 2503 (1983).
12. S. D. Gehman, "Chapter 1 Rubber Structures and Properties," pp. 1-36 in *Mechanics of Pneumatic Tires*. Samuel K. Clark, Ed. United States Department of Transportation, National Highway Traffic Safety Administration.
13. Richard N. Dodge and Samuel K. Clark, "Properties of Aircraft Tire Materials," SAE Technical Paper Series. Aerospace Technology Conference and Exposition, Anaheim, California, October 3-6, 1988.

EFFECTS OF INTERMOLECULAR INTERACTIONS
IN A
CYCLIC SILOXANE BASED LIQUID CRYSTAL

Edward Peter Socci
Graduate Student
Department of Materials Science and Engineering

University of Virginia
Charlottesville, VA 22903-2442

Final Report for:
AFOSR Summer Research Program
Wright Laboratory

Sponsored by:
Air Force Office of Scientific Research
Bolling Air Force Base, Washington, D.C.

September 1992

Effects of Intermolecular Interactions
in a
Cyclic Siloxane Based Liquid Crystal

Edward Peter Socci
Graduate Student
Department of Materials Science and Engineering
University of Virginia

Abstract

A substituted cyclic penta(methylsiloxane) liquid crystal containing combinations of pendant biphenyl-4'-allyloxybenzoate (B) and cholesterol-4'-allyloxybenzoate (C) mesogens was examined using computer molecular modeling. These compounds are of interest for possible use in the fabrication of laser resistant optical devices. Molecular mechanics (MM) and dynamics (MD) calculations were undertaken to assess the conformation and mesophase structure of this material.

The Gibbs free energy of likely interactions of (B) and (C) mesogens was calculated. Dissociation temperatures for pairs of (B) and (C) mesogens were also calculated. Results suggest possible models for the structural ordering of the mesophase based upon strong interactions between mesogens in certain preferred orientations. These interactions could lead to the formation of a supramolecular, pseudo main chain polymer from this low molecular weight liquid crystal, which would account for the unusually good fiber forming characteristics of this material.

Effects of Intermolecular Interactions in a Cyclic Siloxane Based Liquid Crystal

Edward Peter Socci

Section 1

CYCLIC SILOXANE-BASED LIQUID CRYSTALS

The focus of this work is the cyclic penta(methylsiloxane) based liquid crystal (LC) with combinations of cholesteryl-4'-allyloxybenzoate and biphenyl-4'-allyloxybenzoate mesogens pendant on the siloxane ring (Figure 1.1). These low molecular weight liquid crystals can be drawn into fibers tens of meters in length, but their fiber drawing characteristics are highly dependent upon the mole fraction of cholesterol mesogen ($X_{\text{cholesterol}} = P/(P+Q)$) present in the material [1]. In this chapter the results of a computer molecular modeling study of this LC are presented in which we attempt to understand the intermolecular packing order of these LC molecules in the mesophase.

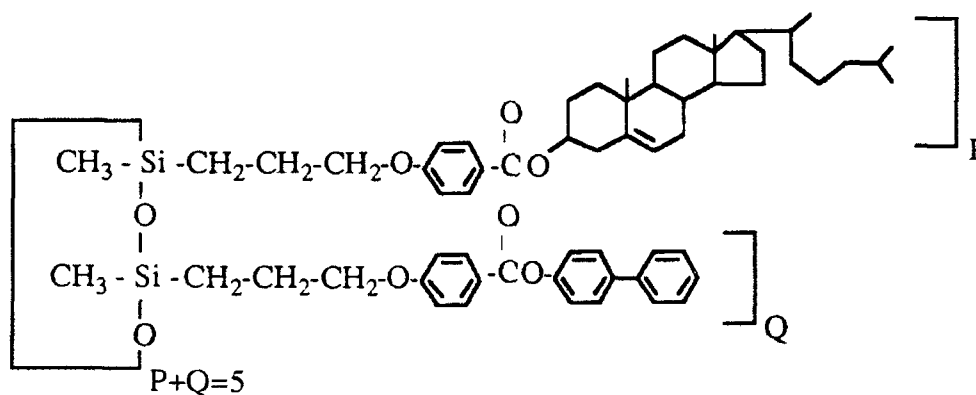


Figure 1.1: General Chemical Structure of the cyclic siloxane-based liquid crystal

The work of Bunning and co-workers [1] involving characterization of cyclic siloxane-based LC's by X-ray diffraction has suggested a variety of possible conformations and packing arrangements for the mesogens pendant on the siloxane ring. Each of these packing schemes is dependent upon the conformation of the LC and the mole fraction of individual mesogens present in the system. One possible conformer as suggested by X-ray diffraction is cylindrically shaped

(herein named the cylinder model) and is illustrated in Figure 1.2. Molecules in this conformation could form strong intermolecular interactions with mesogens of neighboring molecules causing the formation of an interdigitated structure (i.e. like the meshing of fingers on two hands).

When drawn into fibers, the mesogens pendant on the siloxane ring are oriented parallel to the draw direction (orienting the lamellae at 90° to the draw direction) and form what amounts to a S_A mesophase. It has been shown experimentally [1] that materials with a low mole fraction of cholesterol mesogen can be drawn into fibers over 10 meters long. However, fibers cannot be drawn in material with a high mole fraction of cholesterol mesogen. It is apparent that the cholesterol mesogen (C) and the biphenyl mesogen (B) each play important roles in 1) determining the structure of the mesophase, and 2) determining the macroscopic properties of the material. Specifically, the addition of (C) mesogen has a drastic effect on the fiber forming properties of the system.

The first topic in this study addresses the possible orientations that (B) and (C) mesogens may adopt in an effort to find their lowest energy packing relationships. Results of this calculation are useful in predicting possible structural ordering in the mesophase. The method used is molecular mechanics which calculates (in this case) non-bonded interactions between rigid molecules at different orientations in space.

The second topic is the dynamics of mesogen interactions. An analysis of the interactions between mesogen pairs at temperatures between 0 K (minimum energy structures) and the calculated dissociation temperature is made. The ordering of the dissociation temperatures for each mesogen pair is a good indication of the relative stability of each interaction.

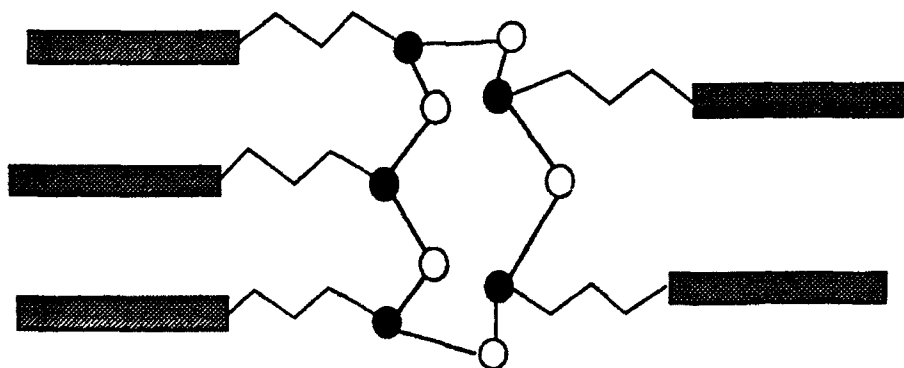


Figure 1.2: Cylindrical conformer of cyclic siloxane based liquid crystal

Section 1.1

METHODS

1.1.1. Molecular Mechanics Calculations

MM techniques were used to study the packing behavior of the (B) and (C) mesogens independent of the siloxane ring, i.e. calculations were made only on the mesogens, (the allyloxybenzoate leader group and the (B) or (C) moiety: see Figure 1.1). Mesogens were constructed and optimized (energy minimized) using the SYBYL molecular modeling package [2].

By quantifying the number and magnitude of favorable mesogen interactions, insight into the structure of the mesophase may be made. MM is used to study six possible interactions likely to occur in a mesophase of cylindrical conformers. These six interactions are illustrated in Figure 1.3. Shown in Figure 1.3(a) are the two orientations for a biphenyl mesogen pair, (B)-(B). In the parallel orientation, the leader groups of each mesogen are side by side, while in the anti-parallel orientation, the leaders are located next to the biphenyl group. Each of these orientations is possible in postulated models for the structure of the mesophase. The two starting orientations considered for the mixed mesogen pair (B)-(C) are shown in Figure 1.3(b). Likewise, the starting orientations for the cholesterol mesogen pair (C)-(C) are shown in Figure 1.3(c).

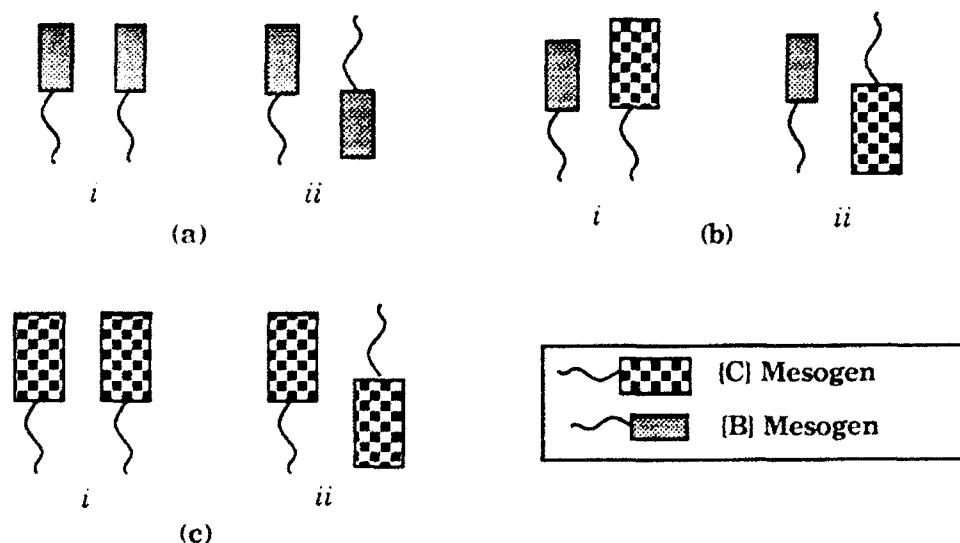


Figure 1.3: Six interactions likely to occur between mesogens in a mesophase of cylindrical siloxane conformers: (B)-(B) parallel (i) and anti-parallel(ii) (a), (B)-(C) parallel (i) and anti-parallel (ii) (b), and (C)-(C) parallel (i) and anti-parallel (ii) (c)

The initial orientation of the mesogens was optimized by allowing a z-translation (along the molecular axis) for each mesogen in the pair. Once an optimal translation was found (minimum energy translation), it was used throughout the calculations. Orientations between mesogens in each pair were then systematically surveyed to locate most favorable relationships by calculation of the intermolecular energy between mesogens. In each case, the angular orientations θ_1 and θ_2 were varied in 40 degree intervals for a given inter-chain distance, d (see Figure 1.4 for an explanation of these variables). This process was repeated for inter-chain distances between 3 Å and 10 Å in 0.5 Å increments. The contour maps resulting from these calculations on {B}-{B}, {B}-{C} or {C}-{C} pairs indicate the most favorable orientations for the molecules as a function of θ_1 , θ_2 and d , i.e. each data point on the contour map corresponds to the value of d which gave the lowest intermolecular energy for the two mesogens at that particular set of orientation angles θ_1 and θ_2 . Only those orientations with intermolecular energies within 50% of the minimum energy orientation are shown.

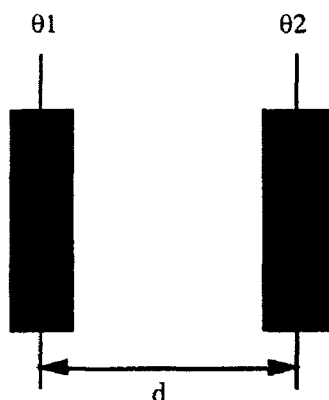


Figure 1.4: Illustration of orientation angles θ_1 and θ_2 and separation distance d

MD simulations (for calculation of mesogen dissociation temperatures) were initiated by slowly heating the molecules to a starting temperature (between 100 K and 125 K) at a rate of 25 K/500 fs. The time step used in all calculations was 1 fs, non-bonded interaction lists were updated every 25 fs and molecular trajectories were saved every 100 fs. The non-bonded cut-off distance was 12 Å. The mesogens were equilibrated at the starting temperature for a period of 10^4 fs. From this point, the temperature of the system was increased in either 5 K or 10 K increments over a period of 500 fs (heating step) and then simulated at this new temperature for 10^4 fs. Upon completion of the simulation at a specific temperature, the trajectories were analyzed for dissociation. The variations in intermolecular energy and distance between mesogens were calculated as a function of simulation time. If the mesogens remained close (i.e. the intermolecular energy between mesogens was non-zero) the simulation was continued at the next highest temperature. This procedure was repeated until the calculated dissociation temperature.

Section 1.2

RESULTS

1.2.1. Molecular Mechanics Calculations

1.2.1.1. Biphenyl Mesogen Pair

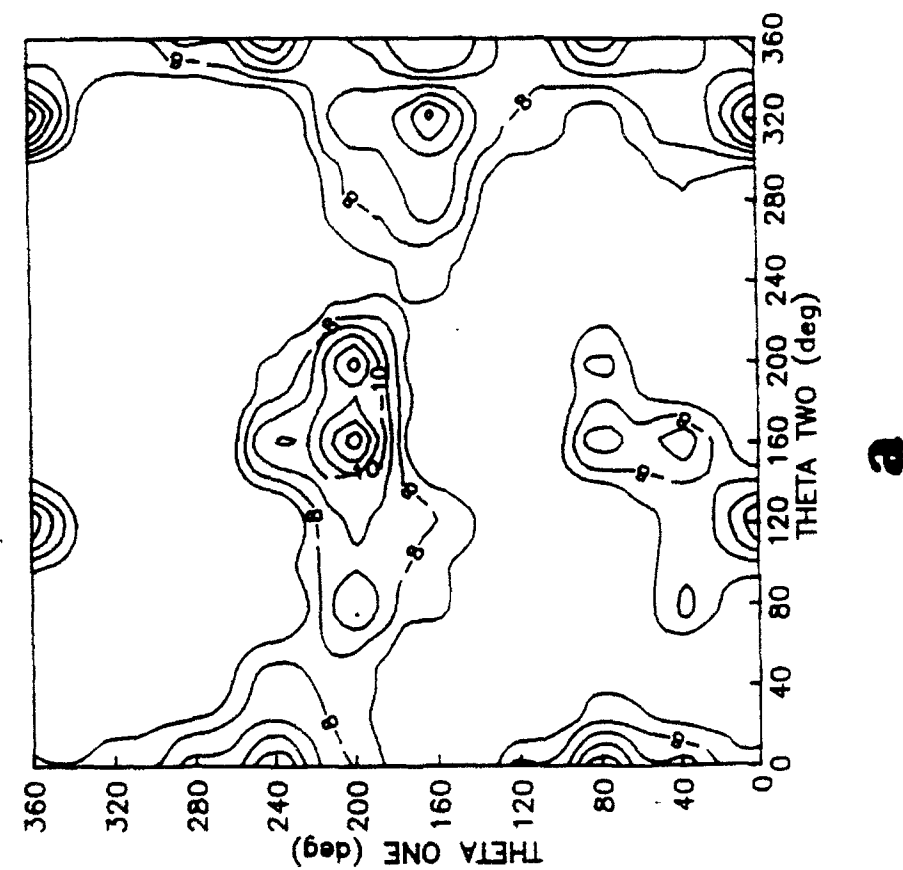
The contour maps for the {B}-{B} mesogen pair (Figure 1.5) illustrate the low energy packing relationships available to these mesogens. The number of favorable orientations (those within 50% of the energy minimum) is approximately equal for the parallel and the anti-parallel orientations. Although the number of favorable orientations is comparable, the magnitudes of the interaction energies are different. The calculated intermolecular energies between mesogens are generally more negative (for most values of θ_1 and θ_2) when the {B} mesogens are in the anti-parallel orientation. Table 1.1 gives a quantitative comparison of the magnitudes of the calculated intermolecular energies between mesogens for several of the low energy orientations in both the parallel and the anti-parallel positions.

Table 1.1. Calculated Intermolecular Energies for {B} mesogen pairs in favorable orientations

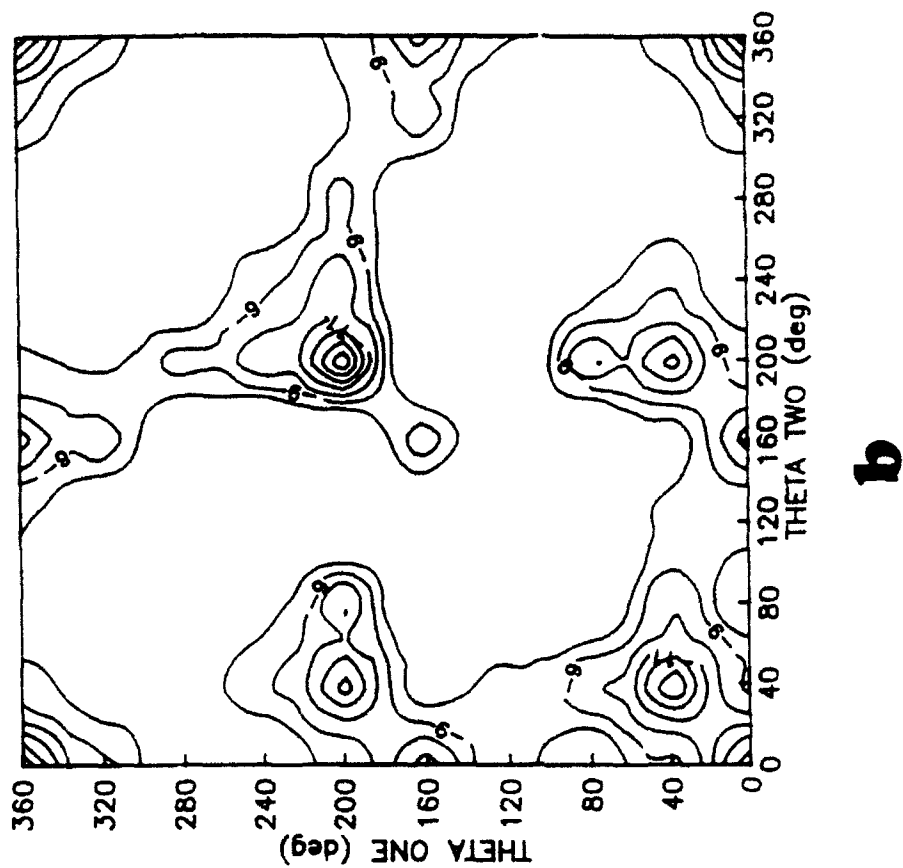
Parallel				Anti-parallel			
θ_1 (deg)	θ_2 (deg)	d (Å)	Inter E (kcal/mol)	θ_1 (deg)	θ_2 (deg)	d (Å)	Inter E (kcal/mol)
160	200	4.0	-13.8	200	200	4.0	-15.0
320	0	4.0	-13.1	40	40	4.5	-14.2
320	160	4.5	-12.4	0	0	4.0	-13.8
0	240	4.5	-12.2	40	200	4.5	-12.6

1.2.1.2. Biphenyl-Cholesterol Mesogen Pair

Contour maps of the interaction search for {B}-{C} mesogen pairs (in the parallel and anti-parallel orientations) as a function of θ_1, θ_2 and d are shown in Figure 1.6. In each map, only those intermolecular energies within 50% of the lowest energy are shown. By visual inspection of the contour maps, there are more favorable packing arrangements for anti-parallel {B}-{C} pairs. The intermolecular energies, itemized in Table 1.2, are more negative in the parallel orientation.

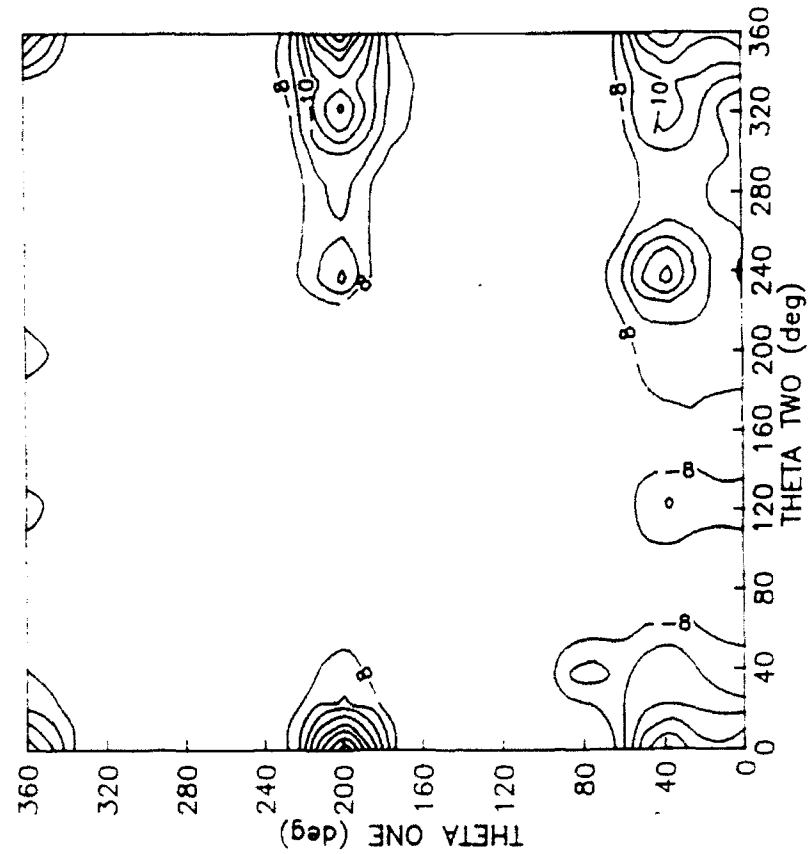


a

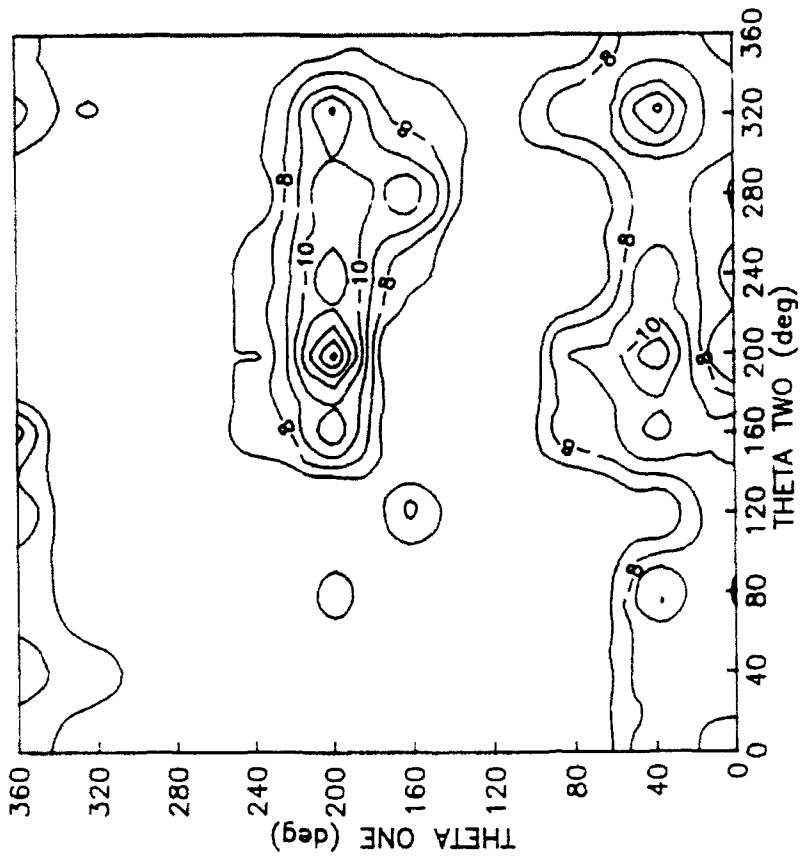


b

Figure 1.5 Contour map describing intermolecular energy as a function of orientation angle for (B) pairs in the parallel (a) and anti-parallel (b) orientations



a



b

Figure 1.6 Contour map describing intermolecular energy as a function of orientation angle for (B)-(C) pairs in the parallel (a) and anti-parallel (b) orientations

Table 1.2 Calculated Intermolecular Energies for {B}-{C} mesogen pairs in favorable orientations

Parallel				Anti-parallel			
θ_1 (deg)	θ_2 (deg)	d (Å)	Inter E (kcal/mol)	θ_1 (deg)	θ_2 (deg)	d (Å)	Inter E (kcal/mol)
0	200	4.5	-15.8	200	200	4.5	-14.4
320	200	4.5	-13.4	320	40	5.0	-12.5
0	40	5.0	-13.2	240	200	5.0	-12.1
240	40	5.0	-12.9	200	40	5.0	-12.0

1.2.1.3. Cholesterol Mesogen Pair

The intermolecular energies calculated for the {C}-{C} mesogens (in both the parallel and anti-parallel orientations) as a function of θ_1, θ_2 and d are shown in contour maps in Figure 1.7. The populations of these two maps (in terms of orientations with 50% of the minimum energy orientation) reveal a dramatic difference in the packing behavior of {C} mesogens. There are substantially fewer favorable orientations available to mesogens in the anti-parallel orientation than there are available to mesogens in the parallel orientation. This indicates that {C} mesogens cannot pack as efficiently in the anti-parallel position, possibly because of steric effects (axial methyl groups in the {C} moiety), for many values of θ_1, θ_2 and d.

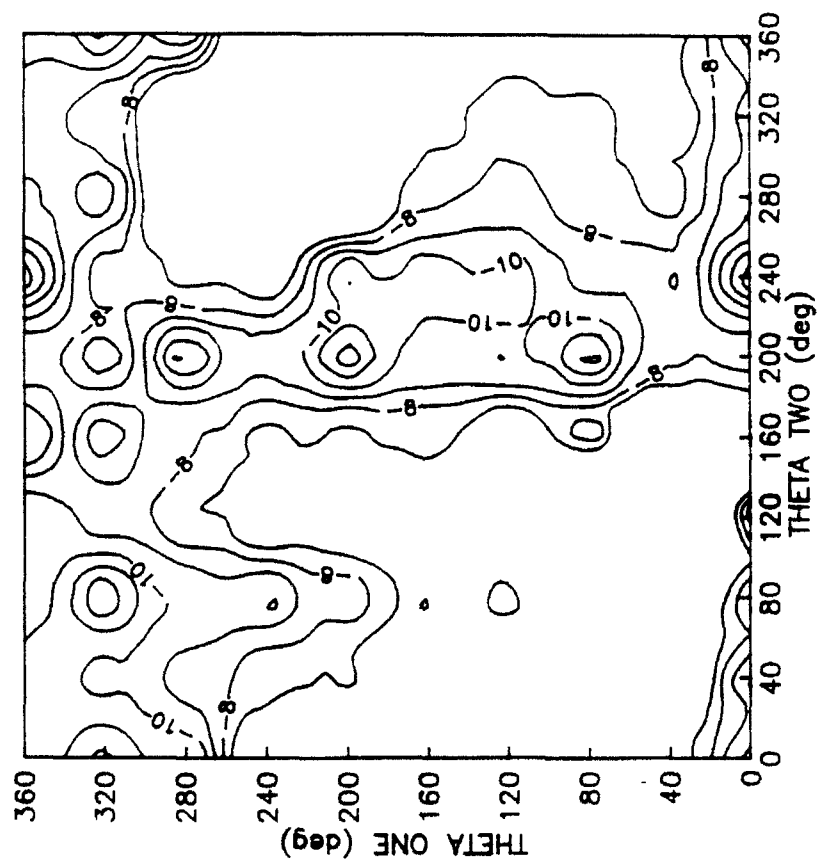
Although the favorable interactions between {C} mesogens are fewer in the anti-parallel orientation, the calculated intermolecular energies are generally more negative in this orientation. The intermolecular energies for several of the most favorable orientations of {C} mesogens in both the parallel and the anti-parallel positions are given in Table 1.3.

Table 1.3. Calculated Intermolecular Energies for {C} mesogen pairs in favorable orientations.

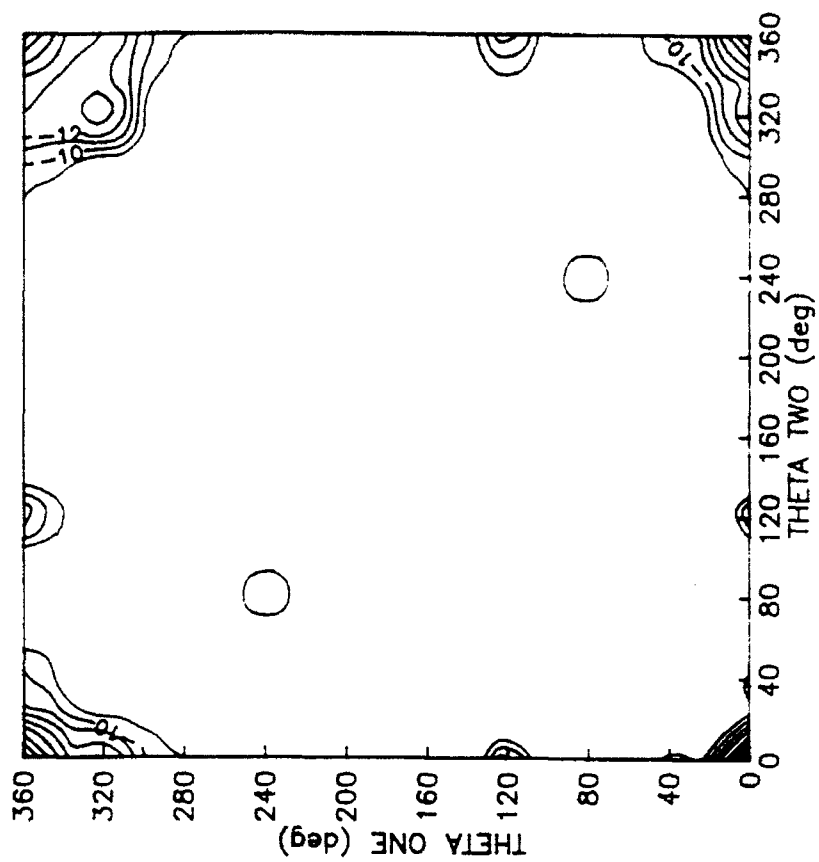
Parallel				Anti-parallel			
θ_1 (deg)	θ_2 (deg)	d (Å)	Inter E (kcal/mol)	θ_1 (deg)	θ_2 (deg)	d (Å)	Inter E (kcal/mol)
240	0	5.5	-13.9	0	0	5.5	-17.4
80	320	6.0	-13.1	320	320	5.0	-14.1
200	200	5.5	-12.8	320	0	5.5	-13.1
200	80	6.5	-12.7				

1.2.2. Molecular Dynamics Simulations

MD simulations were completed in order to evaluate the behavior of {B} and {C} mesogens at elevated temperatures. Emphasis was placed on determining a "dissociation temperature" for the {B}-{B}, {B}-{C} and {C}-{C} mesogen pairs in order to quantify the relative stability of each interaction.



a



b

Figure 1.7 Contour map describing intermolecular energy as a function of orientation angle for [C] pairs in the parallel (a) and anti-parallel (b) orientations

The dissociation temperature was defined as the temperature where the intermolecular energy between mesogens equalled zero. At this temperature, the attractive force between mesogens is overcome by thermal (kinetic) energy input into the system. As a result, the molecules become separated by a distance greater than 12 Å (the non-bonded cut-off distance) and no longer interact.

The graphs in Figure 1.8 indicate the variation of temperature and intermolecular distance for the (B) mesogen pair during the MD simulation (behavior of the (B)-(C) and (C)-(C) mesogen pair was similar). Intermolecular distances were measured between neighboring carbon atoms in each of the mesogens. The observed dissociation temperatures are: $T_d(\text{B})-(\text{B}) = 245 \text{ K}$, $T_d(\text{B})-(\text{C}) = 370 \text{ K}$ and $T_d(\text{C})-(\text{C}) = 445 \text{ K}$.

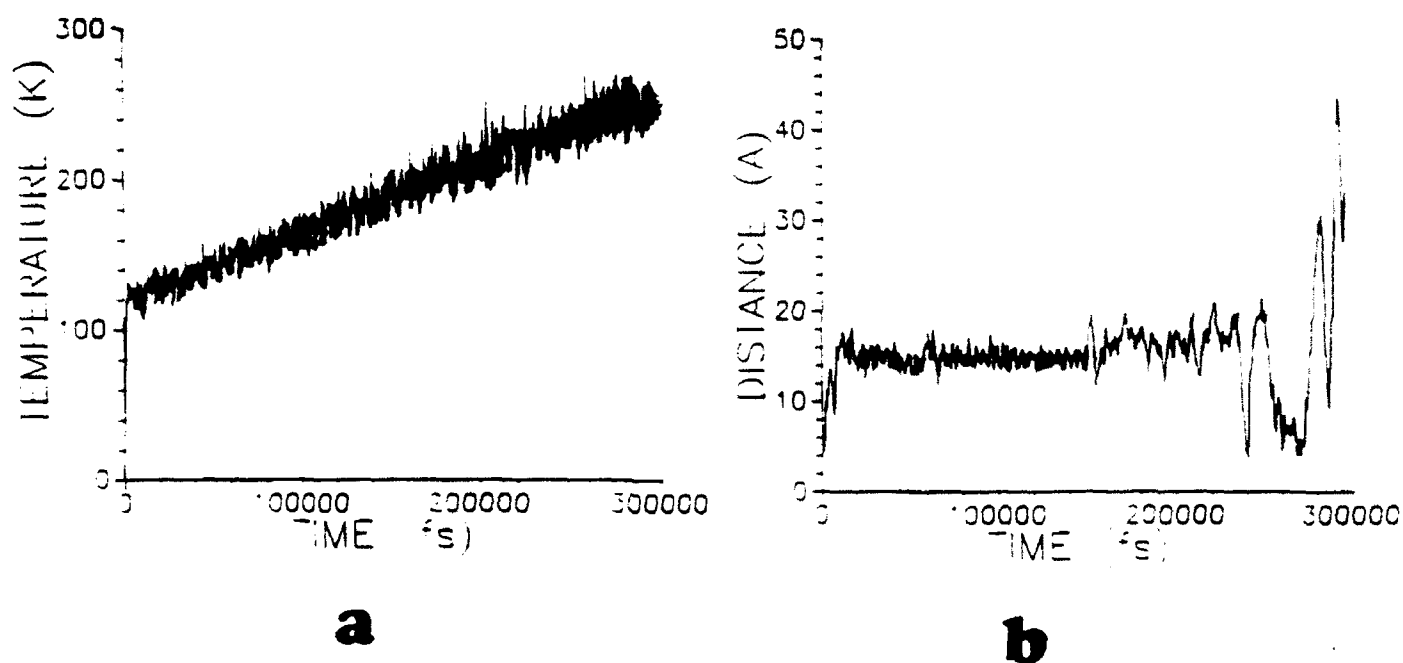


Figure 1.8 Plot of temperature (a) and intermolecular distance (b) versus simulation time for (B) pairs during MD simulation.

Section 1.3

DISCUSSION

As previously indicated, the structure and macroscopic properties of this cyclic siloxane-based LC are affected by the mole fraction of [C] mesogen present in the system. The effect on mesophase structure is revealed by X-ray diffraction measurements. Shown in Figure 1.9 are the measured primary and low angle d-spacings for this LC as a function of mole fraction [C] mesogen [1]. For low mole fraction [C] mesogen ($X_C < 0.30$), there are no measured low angle reflections. The measured primary d-spacings are approximately 25 Å and are consistent with intermolecular ordering of mesogens based upon interdigitation. The 22 Å primary d-spacing ($X_C = 0$) is consistent with a packing pattern based upon complete interdigitation of [B] mesogens. A primary d-spacing of 28 Å ($X_C = 1.0$) is consistent with complete interdigitation of [C] mesogens.

At higher mole fraction [C] mesogen ($X_C > 0.30$) a low angle reflection is consistently observed. This low angle reflection, corresponding to a d-spacing of approximately 50 Å, may describe regions of partially interdigitated [C] mesogens (i.e. [C] mesogens which cannot completely interdigitate because of steric hindrances) which form a secondary packing structure in the mesophase. It is clear from this graph that two or more distinct packing morphologies exist: one dominant at low [C] mesogen concentration and the other at high [C] mesogen concentration.

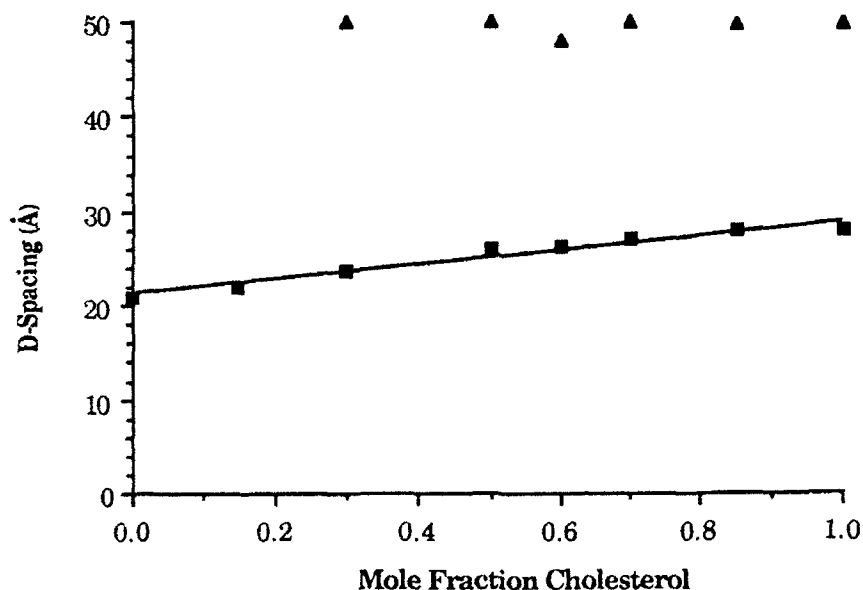


Figure 1.9: Measured primary (-o-) and low angle (-Δ-) d-spacings as a function of mole fraction cholesterol mesogen

1.3.1. Analysis of Molecular Mechanics Results

MM calculations on (B)-(B), (B)-(C) and (C)-(C) mesogen pairs illustrate the effect each can have on the degree of interdigitation in the mesophase. The calculated favorable packing orientations (i.e. parallel or anti-parallel) for each mesogen pair aid in interpreting the structure of the mesophase. A visual inspection of the contour maps in Figures 1.5, 1.6 and 1.7 indicates the different degrees of conformational freedom available to pairs of (B) and (C) mesogens. To quantify this visual impression, the conformational entropy for each pair was computed. The overall partition function and the probability, P_i for a given pair-wise interaction, were computed at 298 K. The energies from which the contour maps were based upon were used to calculate the conformational entropy for each interaction. An entropy of mixing contribution to the total entropy was included for the (B)-(C) mesogen pairs. The conformational entropy is given by:

$$S = -R \sum_i P_i \ln P_i \quad (1)$$

where R is the ideal gas constant and:

$$P_i = E_i / \sum_i \exp(-E_i/RT) \quad (2)$$

The computed conformational entropies for each mesogen pair interaction are listed in Table 1.4.

Table 1.4. Conformational entropies of mesogen interactions (at 298 K)

Interaction Type	Conformational Entropy (kcal/mol.K)
(B)-(B) parallel	3.5
(B)-(B) anti-parallel	3.0
(B)-(C) parallel	3.0
(B)-(C) anti-parallel	2.3
(C)-(C) parallel	3.6
(C)-(C) anti-parallel	2.8

Consistent with the similarity of the contour maps in Figure 1.5, the conformational entropies for (B) mesogen pairs in the parallel and anti-parallel orientation are comparable. However, the conformational entropy is slightly higher for the case of (B)-(B) parallel packing. This result indicates that the parallel orientation for (B) mesogens is slightly less restrictive (i.e. strong interactions can form easier) than the anti-parallel orientation. For the case of (C) mesogen pairs, the conformational entropy of the parallel orientation is significantly higher than for the anti-parallel orientation, indicating that the parallel orientation would be favored.

Similarly, the mixed pair of {B} and {C} mesogens have greater conformational freedom (higher conformational entropy) in the parallel orientation .

The Gibbs free energy of each mesogen pair interaction can also be computed. The thermal average energy, $\langle E \rangle$, for each mesogen interaction ({B}-{B}), {B}-{C}, {C}-{C}) was calculated from the intermolecular energies of the contour maps in Figures 1.5, 1.6 and 1.7 using the relation:

$$\langle E \rangle = \frac{\sum E_i \exp(-E_i/RT)}{\sum \exp(-E_i/RT)} \quad (3)$$

The Gibbs free energy for the interaction (ignoring the $p\Delta V$ term in the enthalpy) is then given by:

$$G = \langle E \rangle - TS \quad (4)$$

The computed values for the Gibbs free energy at 298 K are listed in Table 1.5 and are presented graphically in a "phase diagram" in Figure 1.10. Results from the thermodynamic calculations indicate that the {B}-{B} and {C}-{C} pairs have the lowest free energy in the anti-parallel orientation, while the {B}-{C} mixed pairs are most favorably oriented in the parallel orientation.

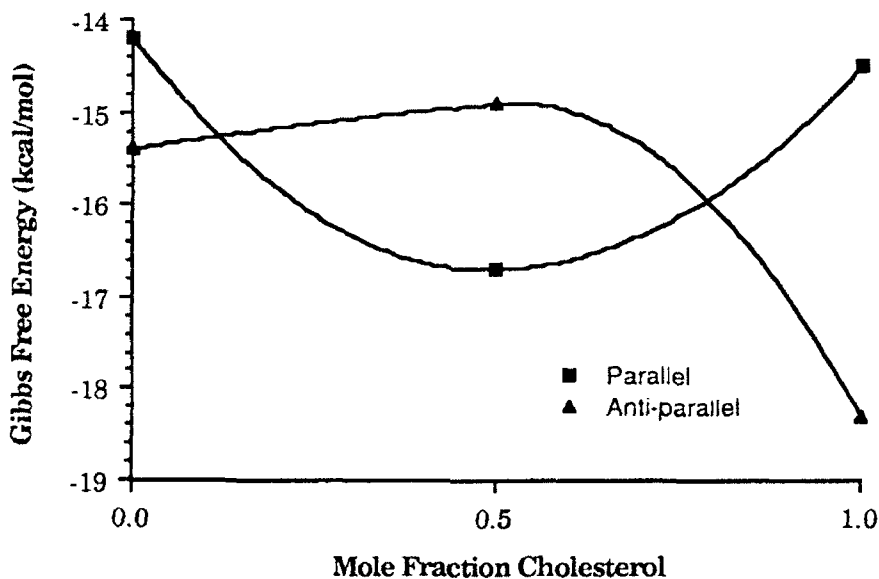


Figure 1.10: Free energy of mesogen interactions as a function of mole fraction cholesterol

Table 1.5. Gibbs free energy of mesogen interactions (at 298 K)

Interaction Type	Free Energy (kcal/mol)
(B)-(B) parallel	-14.2
(B)-(B) anti-parallel	-15.4
(B)-(C) parallel	-16.7
(B)-(C) anti-parallel	-14.9
(C)-(C) parallel	-14.5
(C)-(C) anti-parallel	-18.3

In the context of actual molecular interactions of siloxane conformers, parallel packing relationships correspond to intramolecular interactions and anti-parallel packing relationships correspond to intermolecular interactions (see Figure 1.11: note that mesogen attachment to the ring is via the leader group).

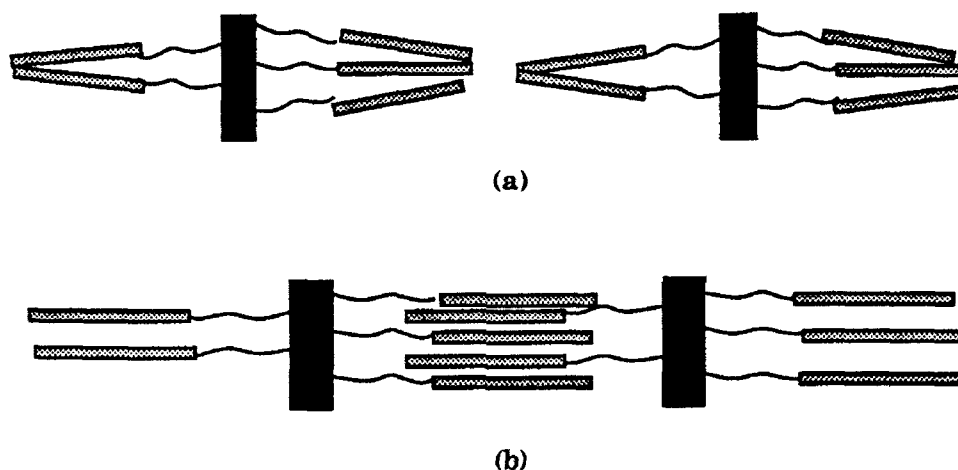
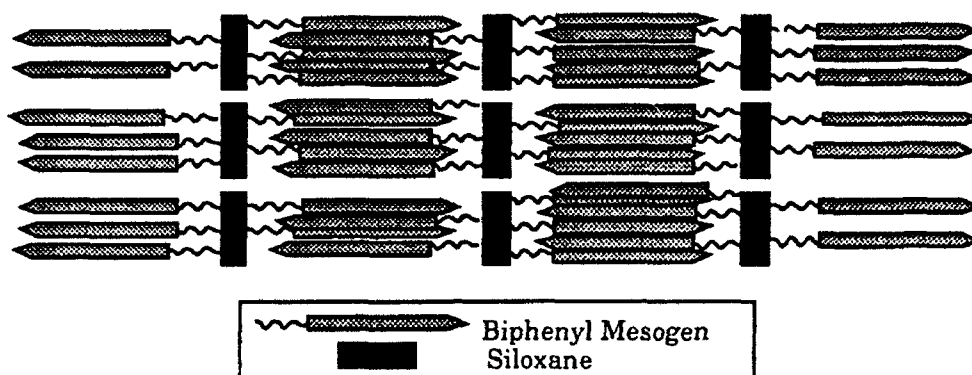
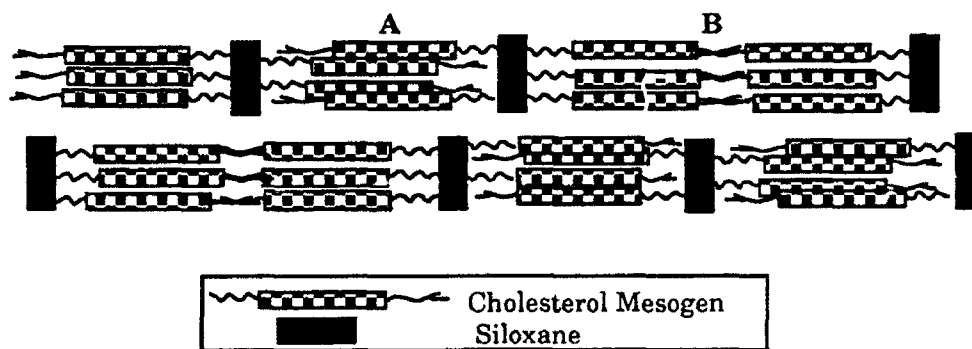


Figure 1.11: Relationship between parallel packing and intramolecular interactions (a) and anti-parallel packing and intermolecular interactions (b)

Illustrated in Figure 1.12 are two possible packing schemes supported by the MM results. In Figure 1.12(a) the pendant mesogens on the siloxane ring are all of type (B) ($X_C=0$). The predicted structure, based upon calculated free energies is one in which the mesogens are fully interdigitated (i.e. intermolecular, anti-parallel interactions between (B) mesogens are favored). This structure of fully interdigitated (B) mesogens corresponds to the experimentally observed d-spacings.



(a)



(b)

Figure 1.12: Schematic of packing in: (B) mesogen substituted siloxane molecules (a) and (C) mesogen substituted siloxane molecules (b)

The influence of (C) mesogen content is illustrated in Figure 1.12(b) which depicts a structure based upon an LC with all (C) mesogens pendant on the siloxane ring ($X_C=1.0$). The packing shown in the figure accounts for both the primary and low angle d-spacings present in the X-ray pattern. Region A of Figure 1.12(b) corresponds to the structural order predicted by the free energy calculations. The four (C) mesogens in region A are fully interdigitated (i.e. *intermolecular*, anti-parallel packing is dominant). In region B, the packing of (C) mesogens is not that predicted by the free energy calculations. The packing in this region is of the parallel (*intramolecular*) type. It is postulated here that in these regions of the lamellae (where six (C) mesogens must pack together), the bulky (C) mesogens cannot achieve the favorable low energy *intermolecular* interactions of the anti-parallel orientation. Interdigitation is disallowed in these regions due to steric hindrances and excluded volume effects.

In the case of (B)-(C) interactions, one must compare several possible packing arrangements of mesogens (illustrated in Figure 1.13). In Figure 1.13(a), the (B)-(B) and (C)-(C) *intramolecular* (parallel) interactions must "compete" with (B)-(C) *intermolecular* interactions (anti-parallel orientation). The average free energies for each of these packing patterns are:

-28.72 kcal/mol for the interdigitated (I) structure and -29.72 kcal/mol for the non-interdigitated (NI) structure. These results indicate that interdigitation would be favored in this environment. A second packing scenario is illustrated in Figure 1.13(b), in which (B)-(B) and (B)-(C) *intramolecular* interactions compete with (B)-(B) and (B)-(C) *intermolecular* interactions. Based upon the calculated free energies for these interactions (I=-30.25 kcal/mol, NI=-30.9 kcal/mol), a non-interdigitated structure is expected. The third (B)-(C) mesogen ordering pattern is illustrated in Figure 1.13(c). In this packing pattern, (B)-(C) and (C)-(C) *intramolecular* interactions compete with (B)-(C) and (C)-(C) *intermolecular* interactions. Free energy calculations predict an interdigitated structure (I=-33.12 kcal/mol, NI=-31.16 kcal/mol) in this environment of mesogens. Illustrated in Figure 1.13(d) is a fourth environment for (B) and (C) mesogens. In this packing pattern, (B)-(C) *intramolecular* interactions compete with (B)-(B) and (C)-(C) *intermolecular* interactions. The lower free energy is associated with an interdigitated structure (I=-33.65 kcal/mol, NI=-33.34 kcal/mol).

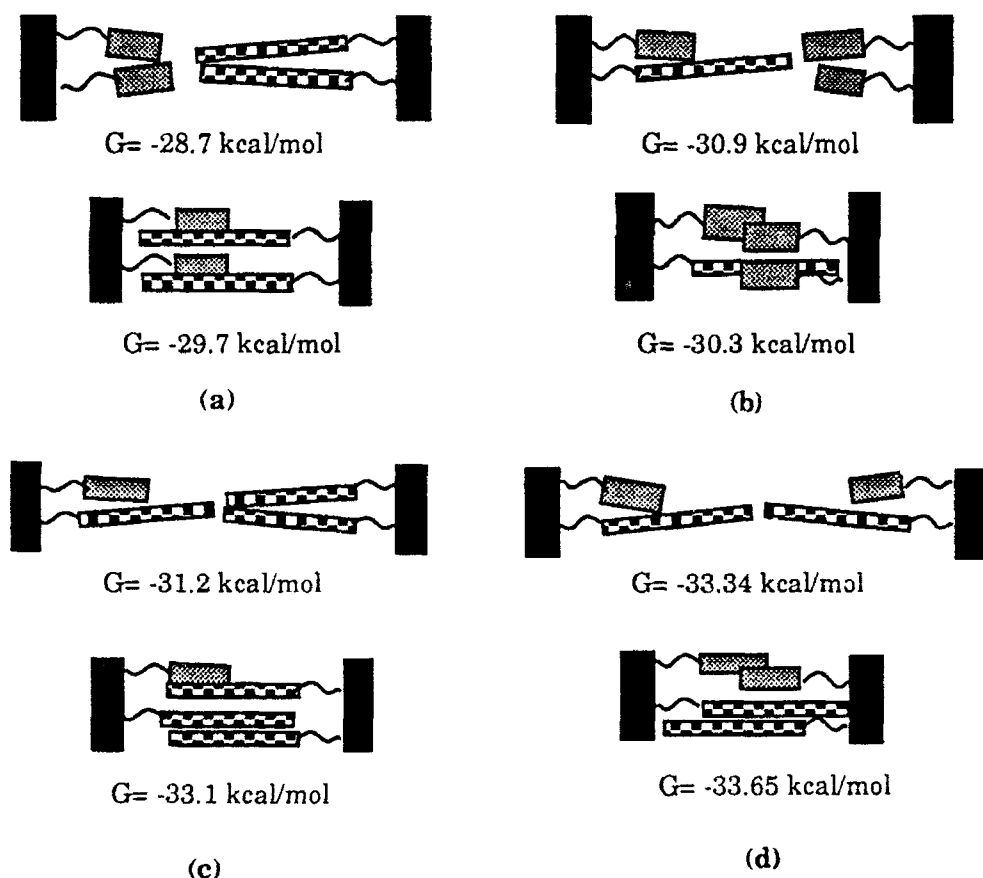


Figure 1.13: Schematic diagrams of packing in a (B)-(C) substituted siloxane. The free energy for parallel (non-interdigitated) and anti-parallel (interdigitated) structures is given.

The results of these calculations indicate that a mesophase of a (B)-(C) substituted LC should contain two general packing patterns. The first is based upon interdigitation of mesogens, in which *intermolecular* interactions are dominant. The second packing pattern is based upon non-interdigitation, in which *intramolecular* interactions are dominant. Indeed, the experimentally observed d-spacings (both low and high angle reflections) indicate two distinct types of structural ordering in a (B)-(C) mesophase.

1.3.2. Analysis of Molecular Dynamics Simulations

Molecular dynamics calculations of the dissociation temperature of the (B)-(B), (B)-(C), and (C)-(C) pairs are also useful in describing the stability of this siloxane based LC. The ordering of the calculated dissociation temperatures for each of the mesogen pairs indicates the relative stability of each interaction. Results from the MD simulations indicate that (C)-(C) mesogen interactions are the most stable, followed by (B)-(C) and (B)-(B) mesogen interactions, respectively. These results are in agreement with the free energies calculated for the individual interactions ($|G_{BB}| < |G_{BC}| < |G_{CC}|$). This suggests that (C)-(C) interactions, once formed, will be difficult to overcome.

The relative positions of mesogens during the MD simulation is also of interest. Each mesogen pair underwent some intermolecular reorganization during the initial heating stage of the simulation. Each of the mesogen pairs changed their relative orientations during the first 10^4 fs of the simulations (See Figure 1.9(b) as an example). These positional fluctuations were accompanied by a significant increase (and a subsequent decrease) in the intermolecular energy between mesogens, caused by the introduction of kinetic energy to the system.

From the graph in Figure 1.9(b), the positional behavior of (B) mesogens during the dynamics simulation can be assessed. At the onset of the MD simulations, the mesogens were anti-parallel. A re-orientation of the mesogens occurred (as a result of heating) after 20000 fs to the parallel position. This parallel orientation was maintained through 240000 fs of the simulation, at which time the mesogens returned to anti-parallel. These results suggest that both the parallel and the anti-parallel orientations can be populated by (B) mesogens. Favorable relationships are found in both orientations of mesogens, as each are within $k_B T$ of one another.

Section 1.4

RECOMMENDATIONS FOR FUTURE WORK AND CONCLUSIONS

1.4.1. Recommendations for Future Work

Further study on the interactions between mesogens should be undertaken. Calculations on the packing of three (B) and three (C) mesogens may help describe the structural ordering of the mesophase. Identification of low energy packing relationships between three mesogens and the types of packing which would result (parallel or anti-parallel) could yield information on the d-spacings and relative intensities of their corresponding reflections that would result from such a structure.

Additional MM calculations should be undertaken to examine the lowest energy conformation of a substituted cyclic penta(methylsiloxane) liquid crystal with various combinations of (B) and (C) mesogens. A systematic study of the conformation of this molecule as a function of mole fraction (C) mesogen should be undertaken to determine the lowest energy conformation for each possible substitution ratio of (B) and (C) mesogens.

1.4.2. Conclusion

The use of computer molecular modeling has been useful in interpreting X-ray diffraction data of a cyclic siloxane-based LC. MM calculations support structural differences in the mesophase as the mole fraction of (C) mesogen is varied. Intermolecular ordering of mesogens ranging from partial to complete interdigitation is predicted. Calculated dissociation temperatures for (B) and (C) mesogen pairs indicate the stability of the various interactions. Results of this computational study lend support to proposed structural models for the mesophase ordering in this LC.

References

- (1) Bunning, T. J., Klej, H. E., Samulski, E. T., Crane, R. L., Linville, R. J., *Liquid Crystals*, 1991, 10, 445.
- (2) SYBYL program and manual, release 5.4, Tripos Associates, 1991.

THIS PAGE INTENTIONALLY LEFT BLANK

A LITERATURE REVIEW OF THE REACTION KINETICS OF THE
PYROLYSIS OF PHENOLIC/GRAPHITE COMPOSITE MATERIAL

Kimberly A. Trick
Ph.D. Candidate
Department of Materials Engineering

University of Dayton
300 College Park
Dayton, Ohio 45469

Final Report for:
Summer Research Program
Wright Laboratory at Wright-Patterson AFB

Sponsored by:
Air Force Office of Scientific Research
Bolling Air Force Base, Washington, D.C.

September 1992

A LITERATURE REVIEW OF THE REACTION KINETICS OF THE PYROLYSIS OF PHENOLIC/GRAPHITE COMPOSITE MATERIALS

Kimberly A. Trick
Ph.D. Candidate
Department of Materials Engineering
University of Dayton

Abstract

The material property/process parameter relationship knowledge base required for expert process control of the carbonization of phenolic/graphite composite materials has been found to be very limited. Specifically, a basic fundamental knowledge of the pyrolysis reaction kinetics does not exist. A review of three incomplete and/or contradictory proposed mechanisms has been made and a research plan for studying the reaction kinetics of the pyrolysis of phenolic/graphite composite materials using thermogravimetric analysis, gas chromatography/mass spectrometry, and infrared spectrometry has been proposed.

A LITERATURE REVIEW OF THE REACTION KINETICS OF THE PYROLYSIS OF PHENOLIC/GRAPHITE COMPOSITE MATERIAL

Kimberly A. Trick

1 Introduction

Carbon-carbon materials are a class of composites which consist of carbon fiber in a carbonaceous matrix. The main advantages of carbon-carbon composites are their low density and excellent mechanical properties to temperatures of 3000°F.

Carbon-carbon was discovered by accident in 1958 in a laboratory at Chance Bought Aircraft company. An analysis was being performed on an organic matrix composite which required the composite to be exposed to air. The crucible lid was accidentally left on and the material underwent pyrolysis instead of oxidation. The resulting charred composite was noted to exhibit structural characteristics. The first intensive research done on carbon-carbon was for use in thermal protection systems for the space shuttle. The space shuttle required a lightweight material which could be fabricated into structural shapes, withstand temperatures as high as 3000 °F, and be reusable for 50 - 100 thermal cycles. The answer to this requirement was carbon-carbon.

Applications of carbon-carbon have broadened to include rocket nozzles, aircraft brakes, turbine engine parts, and even clutches for race cars. The major disadvantages of carbon-carbon are lack of oxidation resistance and expense. The expense is primarily due to the immature state of carbon-carbon processing technology. Carbon-carbon component processing cycles are extremely long and labor intensive and developed largely by trial and error. Because of the trial and error approach to processing, a large amount of reject material is produced.

Carbon-carbon is manufactured in one of three ways. The most common is the impregnation a fibrous preform with carbonizable precursor, such as a phenolic resin, followed by heat treatment to carbonize the precursor. The resulting carbonaceous solid is usually impregnated with additional matrix and re-carbonized to achieve the desired density. The second manufacturing method involves impregnation of the preform with liquid pitch and heat treatment to carbonization or graphitization temperatures. Finally, the third method involves densifying the preform with vapor-deposited carbon.

The longest, most critical step in the processing of carbon-carbon composites from phenolic resin precursors is first carbonization. Damage is frequent, sometimes difficult to detect, and expensive to correct. Studies have shown that matrix cracking is a consequence of "normal" carbonization and is desirable to relieve stress and produce open porosity. However, while micro-cracking is seen as desirable, delaminations cannot be tolerated and are normally thought to be caused by excessively fast heating which results in stresses capable of causing failure. To minimize the chance of delaminations, manufacturers typically carbonize at very slow rates. This method of delamination avoidance is at the cost of long, expensive processing times. An increased understanding of the carbonization of phenolic matrix would allow optimization of the process cycle. This would lead to reduced processing time and less product damage, and thus reduce the cost of carbon-carbon components.

Development of an expert model process control system for the first carbonization step of carbon-carbon processing would significantly advance the state of carbon-carbon processing. An expert model process control system requires a material property/process parameter relationship knowledge base, models describing the processes, and in-situ sensors. The existing property/parameter knowledge base for carbonization of phenolic/graphite composites is very limited. Specifically, a basic fundamental knowledge of the pyrolysis reaction kinetics does not exist.

1.1 Problem Description

Understanding the reaction kinetics mechanism of the pyrolysis of a phenolic/graphite composite is essential to development of an effective expert process control system for carbonization of carbon-carbon materials. During carbonization, material properties are a function the extent of pyrolysis and thus a complete material property/process parameter knowledge base is not possible without an understanding of the pyrolysis mechanism. While several mechanisms have been proposed, they are based on the pyrolysis of neat resin only and are incomplete and/or contradictory. There is currently very little known of the effect of heating rate on the reaction mechanism.

1.2 Objectives

The purpose of the work proposed in this review is to significantly increase the understanding of the reaction kinetics of the carbonization of cured phenolic/graphite composites with an emphasis on the effect of heating rate on the kinetics. Experimental techniques recommended include thermal analysis, both thermogravimetric and differential, gas chromatography/mass spectrometry analysis of gaseous pyrolysis products, and infrared spectroscopy analysis of the intermediate structures present during pyrolysis. Additionally, the development of curves relating carbonization temperature to extent of carbonization is described. All experimental techniques are to be applied at a number of heating rates to determine the effects, if any, of heating rate on pyrolysis.

Resulting experimental data would be used to evaluate the pyrolysis mechanisms reported in the literature, clear up areas of discrepancy, and propose a complete mechanism.

1.3 Anticipated Results

The proposed experimental work and analysis would result in an improved fundamental understanding of the reaction kinetics of the pyrolysis of a phenolic resin in a phenolic/graphite composite. The specific effects of heating rate on the pyrolysis process could be reported. The fundamental reaction mechanism

knowledge would be applicable for use in the development of an expert process control system to control the carbonization of phenolic/graphite composite materials in the production of carbon-carbon components.

2. Review of the Literature

An expert process control system requires a material property/process parameter relationship knowledge base, models describing the processes, and in-situ sensors. A literature review of these components with respect to the carbonization of phenolic/graphite composites reveals that very little work has been accomplished in the determination of process parameter/material property relationships and that the fundamental reaction kinetics of the pyrolysis of the composite phenolic matrix is not well understood. A review of reaction mechanisms of phenolic pyrolysis reported in the literature is included in this section. Also in this section is a review of reaction kinetic analysis techniques.

2.1 Reaction Mechanisms

During curing of phenolic resins, methylene bridges are formed by condensation and water is released. Thus the cured phenolic matrix, prior to first carbonization, consists of phenol rings held together by methylene bridges. During carbonization this structure is pyrolyzed via a complicated set of reactions to form a restructured carbon matrix. The primary reaction products of the pyrolysis are water, methane, hydrogen, carbon monoxide, and carbon dioxide. Due to the complexity of the reactions, little is known about the actual reaction mechanisms. Three mechanisms have been proposed in the literature and are being studied. A brief overview of each mechanism is described in this section. All are incomplete and/or contradictory to one another.

2.1.1 Ouchi Mechanism

Ouchi and Honda [1] studied the carbonization process of neat phenolic resins using x-ray diffraction, gaseous decomposition product analysis, and infra-red spectrometry. Carbonization of powdered resin

was performed in vacuum at a rate of $2^{\circ}\text{C}/\text{min}$. From these studies, Ouchi and Honda propose the following mechanism for phenolic resin pyrolysis.

- a) The phenol-formaldehyde resin structure does not change at temperatures below 300°C .
- b) Above 300°C , water begins to evolve due to a dehydration reaction involving two phenolic hydroxyl groups. Diphenyl-ether linkages between benzene nuclei result.
- c) Extensive reaction occurs between 400°C and 600°C . The dehydration reaction and evolution of water described in b) continues and methylene bridges are eliminated by reacting with pyrolysis products H_2O and H_2 .

As will be discussed below, later work by Jackson and Conley [2] was unable to experimentally verify this mechanism.

2.1.2 Parker/Winkler Mechanism

Parker and Winkler [3] studied the carbonization process of neat phenolic resins in powder form using mass spectrometry, residual elemental analysis and thermogravimetric techniques. Heating rates used were 0.5 and $1.0^{\circ}\text{C}/\text{min}$. From their work they propose the following mechanism.

- a) Initiation of pyrolysis involves breaking the thermodynamically weakest carbon-carbon bond that connects the aromatic pendant group to the main chain. Either a phenol or cresol radical is formed, dependent on which side of the single bonded phenol ring the break occurs. Abstraction of a hydrogen atom from a main chain methylene group occurs quickly to give phenol and cresol vapor products. One terminal radical is formed for each pendant group eliminated and one diarylketone is formed for each hydrogen abstraction. This entire step occurs below 350°C .
- b) Termination of main chain radical pairs produce a stable, thermally cross linked intermediate polymer structure. This intermediate retains all phenol groups which are initially bonded in the main chain to two or more methylene groups.

- c) Above 500°C the crosslinked intermediate polymer structure loses methane and carbon monoxide by chain scission and recombination to form an unstable char consisting of a cross linked diphenyl ether polymer.
- d) A thermally stable char network is formed by continued crosslinking of the aromatic rings present in the unstable char with the elimination of hydrogen and water.

In summary, the mechanism proposed by Parker and Winkler involves formation of a thermally cross-linked intermediate structure with elimination of pendant aromatic rings and retention of all aromatic carbons which are multiple bonded in the initial polymer structure.

2.1.3 Conley Mechanism

Jackson and Conley [2] have investigated the carbonization of neat phenolic resins using infrared spectrometry, vapor-phase chromatography performed at 310°C/sec., thermogravimetry performed at 0.05°C/sec. (3°C/min.), and x-ray analysis. Based on their findings, Jackson and Conley propose the following mechanism.

- a) Water and paraformaldehyde are formed at 400°C and result from the loss of residual methylol groups left after curing.

This proposed mechanism step is in contrast to that proposed by Ouchi and Honda where water evolves due to a dehydration reaction involving two phenolic hydroxyl groups. The Ouchi and Honda mechanism does not predict a loss of methylol groups. The Conley mechanism does not predict a diphenylether linkage.

- b) Carbon dioxide and carbon monoxide are produced by a number of oxidation reactions involving the resin and product oxygen. Oxidation of methylene linkages, residual methylol groups, and dihydroxybenzophenone linkages present in the cured resin with product oxygen all produce carbon dioxide and carbon monoxide.
- c) Oxidative degradation is believed to be the primary mechanism of degradation. Decarboxylation and decarbonylation form CO₂ and CO. Phenol, cresol, and other higher phenolic species are formed

from dehydroxydiphenylmethane and low polymers which have been terminated and trapped in the cured resin system.

- d) Above 450°C, decomposition to produce char and CO, via ring scission, is rapid. The char is formed by the decomposition of the oxidized resin through a quinone-type intermediate which accounts for the appearance of CO.

2.2 Reaction Kinetics Analysis Techniques

A number of techniques have been used to study the reaction kinetics of a polymer system. Among these are thermal methods (including both thermogravimetric and differential analysis), gas chromatography/mass spectrometry, and infrared spectroscopy. The application of these techniques to the study of a polymer system's reaction kinetics will be discussed in this section. Specific examples of their use in the study of pyrolysis of phenolic resin pyrolysis will be cited.

2.2.1 Thermal Analysis Methods

Thermal analysis methods are used in the study of liquid and solid state chemical reactions at elevated temperatures. These methods involve continuous measurement of a change in a physical property such as mass, volume, heat capacity, etc. as the sample temperature is increased at a predetermined rate. *Thermogravimetry measures the change in sample mass as a function of temperature.* Differential thermal analysis measures a thermal property differential between the sample and a thermally inert material.

2.2.1.1 Thermogravimetric Techniques

- Thermogravimetric analysis (TGA) techniques have long been an important tool in the characterization of polymeric materials. Two types of TGA exist: isothermal and dynamic. Isothermal TGA requires measurement of the sample mass as a function of time at a constant temperature. Dynamic TGA requires measurement of the sample mass as the sample is heated according to a predetermined temperature

cycle. Dynamic TGA is the more widely used. The advantages of dynamic over isothermal TGA methods are:

- Fewer data points are required. The temperature dependency of a reaction over wide temperature ranges can be determined from a single dynamic scan while separate isothermal scans must be obtained for each temperature range of interest.
- Continuous weight loss vs. temperature monitoring assures that no features of the kinetics are overlooked.
- Sample-to-sample variations are eliminated with the use of a single scan.
- Premature reactions are possible in isothermal analysis and can make determination of reaction kinetic parameters difficult.

For the remainder of this section, any reference to TGA will imply a dynamic procedure.

From the TGA mass change vs. temperature curve, sample thermal stability and composition information can be obtained. TGA data can also be used to estimate the kinetic parameters of degradation kinetics. The most commonly used analysis techniques for TGA data are the method developed by Freeman and Carroll and the multiple heating rate method. The Freeman and Carroll method calculates the reaction order, n , and the activation energy, E , from one set of TGA data according to:

$$\log R_t = n \log W - (E/2.3R)(1/T) \quad (1)$$

where R_t is the rate of weight loss and W is the weight remaining. The advantages of the Freeman and Carroll method are that only one curve is necessary and both the activation energy and the reaction order can be calculated. The disadvantage of the method is that it is necessary to accurately determine steep slopes.

The multiple heating rate method calculates the activation energy from TGA curves made at a number of different heating rates. If the reaction possesses an activation energy the temperature at which the

maximum rate of weight loss occurs varies with heating rate. This variation can be used to determine the activation energy of the reaction. A plot of the log of the heating rate, RH, versus $1/T_m$, where T_m is the temperature of maximum rate of weight loss, is used to determine the activation energy according to

$$-d \ln(RH) / d(1/T_m) = E/R \quad (2)$$

The disadvantages of the multiple heating rate method are that more than one set of TGA data is needed and no reaction order information is obtained.

TGA techniques have been used in the study of the pyrolysis of phenolic composites. Brown [4] found that TGA of phenolic composites was simple, reproducible, and virtually independent of sample weight and form. He reports TGA data of cured phenolic/carbon laminates at heating rates of 5, 10, and 20 °C/min. (Note that all three rates are considered fast by production standards.) As seen in Table 1, three major weight loss peaks are observed and increased heating rate shifts the weight loss peaks to higher temperatures.

Table 1 Temperature of Weight Loss Peaks - TGA Data			
Heating Rate, °C/min.	Peak 1 °C	Peak 2 °C	Peak 3 °C
5	311	430	527
10	325	446	534
20	342	454	550

A multiple heating rate analysis was used to determine the following activation energies of the three weight loss peaks: 34.4 Kcal/gmole for peak one, 54.4 Kcal/gmole for peak two, and 68.9 Kcal/gmole for peak three.

2.2.1.2 Differential Thermal Analysis

As stated earlier, differential thermal analysis measures a thermal property differential between the sample and a thermally inert material. Thus heat effects associated with physical or chemical changes

are recorded as a function of temperature. In the specific case of differential scanning calorimetry (DSC) both the specimen and an inert reference material are heated at the same predetermined rate. The differential in the amount of heat input required to heat the sample and that required to heat the reference material is recorded. An exothermic reaction in the sample will be evident by a *reduced* amount of input energy required to maintain the specified heating rate. An endothermic reaction in the sample will be evident by an *increased* amount of input energy required to maintain the specified heating rate.

DSC is an effective technique for monitoring chemical and physical processes that are associated with enthalpy changes in any material. When a reaction occurs during DSC analysis, the change in heat content and thermal properties of a sample is indicated by a deflection or peak. As with TGA, if the reaction possesses an activation energy, the temperature at which the peak occurs will vary with heating rate. This variation in peak temperature can be used to determine the activation energy of the reaction. The shape of the DSC curve can be used to determine reaction order. Further analysis of DSC curves by integration of sample peaks can provide heat of reaction information. Integration of the curve can be difficult if there is a large shift in the baseline.

DSC techniques have been used in the study of the pyrolysis of phenolic composites. Brown [4] found it difficult to obtain reproducible results from a composite and ran all data on neat phenolic resin only. It was also found that determination of a shifting baseline, due to the changing heat capacity of the pyrolyzing phenolic, was very difficult. Another concerns involving the use of DSC techniques for the investigation of the phenolic pyrolysis kinetics is the possibility that condensation of volatiles could effect the data. Brown does identify two exothermic peaks which shift with changes in heating rate and reports

- calculated activation energies for both peaks.

2.2.2 Infrared Spectroscopy

Infrared (IR) spectroscopy is used to characterize materials and provide information on the molecular structure of a sample. It is particularly useful for determining the functional groups present in a molecule since many functional groups vibrate at a characteristic frequency independent of their molecular environment. When a sample is hit with a photon beam, infrared radiation is absorbed which can excite molecules into higher-energy vibrational states. This absorption occurs when the energy of the photon exactly matches the energy difference between two vibrational states. Differences in vibrational energy states are unique to the functional group.

In studies of reaction kinetics, IR analysis can be used to identify intermediate structures of the reacting material. This information is used to hypothesize reaction mechanisms.

Ouchi [1] and Jenkins [5] have reported IR spectra of neat phenolic resin pyrolyzed to various temperatures. In both cases pyrolyzed polymer crushed into fine powder and mixed with powdered potassium bromide was compacted into thin disks for IR examination. The main absorption peaks in the spectrums are identified by Ouchi and are summarized in Table 2.

Table 2 Main Absorption Peaks	
Wavenumber cm^{-1}	Functional Group
3400	phenolic OH
3020	aromatic OH
2900,2830	aliphatic CH
1630,1500	aromatic ring
1475	methylene bridge
1100	aliphatic ether link
900, ~700	aromatic CH

2.2.3 Gas Chromatography/Mass Spectrometry

The gas chromatograph/mass spectrometer, GC/mass spec., is the combination of two analytical investigative instruments. This combination instrument is particularly useful for analysis of organic materials. Basically, the gas chromatograph is used to separate a mixture so that each component is individually introduced into the mass spectrometer for analysis.

The gas chromatograph partitions a passing mixture into two phases; stationary and mobile. The proportions in which a component of the mixture distributes itself between the two phases differs for each component. The result is that the different compounds present in a mixture individually exit the gas chromatograph.

The mass spectrometer is useful for identification and determination of the structure of a compound. It provides information on the sample molecular weight and the pieces that constitute the sample molecules. A molecule is passed through an electron beam where it becomes ionized and/or breaks into fragment ions. Once the ions are formed they are detected and analyzed by various methods. Such methods include: time-of-flight, which measures the time it takes for ions to travel a given distance, magnetic sector, which determines the magnetic force needed to bend the ion path, and quadrupole methods, which use oscillating fields to control ion path. In all methods the ion mass is detected.

GC/mass spec. is used in the investigation of reaction kinetics by using the technique to analyze the gaseous by-products of the reaction. Identification of by-products can be useful in understanding the mechanism of a reaction. Ouchi and Honda pyrolyzed a neat phenolic resin at 2°C/min and analyzed the off-gases using mass spectrometry. The species evolved as a function of temperature range are listed in Table 3.

Table 3 Gases Evolved from Phenolic Resin Pyrolysis as a Function of Temperature	
Temperature Increment °C	Species Evolved
100 - 200	H ₂ O
200 - 300	LMS*
300 - 400	H ₂ O, LMS*, CO ₂
400 - 500	H ₂ O, LMS*, CO ₂ , C ₂ H ₆ , H ₂ , CO, CH ₄
500 - 600	H ₂ O, CO ₂ , C ₂ H ₆ , H ₂ , CO, CH ₄
600 - 700	H ₂ O, CO ₂ , C ₂ H ₆ , H ₂ , CO, CH ₄
700 - 800	H ₂ O, H ₂ , CO

* LMS = lower molecular substances. Later work by Shapiro and Asawa [6] identified these LMS as primarily phenol with small quantities of o-cresol and p-cresol.

Brown [4] reports an analysis of pyrolysis products of cured phenolic/graphite composite material using TGA-FTIR, and a chromatography technique. Three zones of volatile formation are identified as shown in Table 4.

Table 4 Gases Evolved from Phenolic/Graphite Composite Pyrolysis as a Function of Temperature	
Temperature Zone	Species Evolved
Centered around 100°C	H ₂ O
250 - 375 °C	H ₂ O, C ₆ H ₅ OH, CH ₂ , CO ₂ , NH ₃
425 - 740 °C	H ₂ O, C ₆ H ₅ OH, CH ₄ , CO, H ₂

3. Experimental Approach

This section will outline the proposed experimental approach to be used to investigate the mechanism and reaction kinetics of the pyrolysis of phenolic/graphite composite materials. The procedures outlined will focus on the determination of the dependency of the pyrolysis on heating rate. Preparation of the cured phenolic/graphite panels and carbonization of a set of samples will be outlined. Finally, the application of thermogravimetric analysis, infrared spectroscopy, and gas chromatography/mass spectrometry to this study will be described.

3.1 Sample Preparation

The objective of the curing procedure is to obtain a set of cured panels which are as consistent as possible using current technology. To meet this objective a modified qualitative process automation (QPA) cure cycle is to be used. The cured panels will be characterized by the following: void and delamination presence, before and after weights, density, fiber volume, thickness, and degree of cure.

Sets of carbonized samples will be made for further evaluation of the intermediate states of the carbonizing panels. A set of samples will be carbonized to temperatures between 200°C and 800°C at 100°C intervals. A number of sets will be made at various heating rates.

3.2 Thermogravimetric Analysis

TGA will be used for two separate studies, the first to investigate weight loss curve dependence on heating rate, and the second to develop extent of carbonization vs. carbonization temperature curves at different heating rates.

To study the dependency of weight loss curves on heating rate, cured composite samples will be run on the TGA from room temperature to 800°C at the following constant linear heating rates: 0.1°C/min., 0.5°C/min., 1.0°C/min., 2.0°C/min. Additional heating rates will be run as required. As reported in section 2.2.1.1, Brown [4] reports that a shift occurs in the weight loss curves with heating rate increase. These shifts have been attributed to the activation energy of the kinetics. However, all rates studied by Brown are very fast compared to industrial carbonization rates. Data reported by Kliner [7] using non-constant heating rates in the industrial rate range showed no shifting of the curves with heating rate.

Brown developed an extent of carbonization as a function of carbonization temperature curve at a constant heating rate of 100°C/min. He assumed that the extent of carbonization is a function of temperature, independent of path, and that the one curve would hold for all heating cycles. The following

study is designed to determine the validity of the assumption and to determine the effects, if any, of heating rate on the extent of carbonization vs. temperature curve. The sets of samples partially carbonized to different extents at different heating rates described above will be used. The extent of carbonization of these panels will be determined as described by Brown using TGA.

Brown assumed that carbonization is complete at 800°C. Percent total weight loss obtained from a TGA of a sample carbonized to 800°C is assigned an extent of carbonization value of 100%. The percent total weight loss obtained from a TGA of a non-carbonized panel (in the as-cured state) is assigned an extent of carbonization value of 0%. The extent of carbonization of the intermediately carbonized panels will be calculated according to:

$$(WL_C - WL_T) / (WL_C - WL_{800}) = \% \text{ of carbonization} \quad (3)$$

Where: WL_C = weight loss of the as-cured panel, WL_{800} = weight loss of panel carbonized to 800°C, and WL_T = weight loss of panel carbonized to temperature T.

3.3 Infrared Spectroscopy

In order to better understand the reaction mechanism occurring during pyrolysis of the phenolic/graphite composite, the molecular structure present at various stages of carbonization will be analyzed using IR techniques. The effect, if any, of carbonization heating rates will also be investigated. This technique has been used to evaluate neat phenolic resin in the past and preliminary experimental work indicates it will be applicable to the pyrolyzed composite as well. The sets of samples described above will be used for this study. Infrared spectroscopy spectrum from wavenumbers of 4000 to 700 cm^{-1} will be obtained from an FTIR spectrometer using potassium bromide discs.

3.4 Gas Chromatography/Mass Spectrometry

To better understand the reaction kinetics of the pyrolysis of the phenolic/graphite composite, the gaseous products produced during carbonization will be studied using GC/mass spec. equipment in conjunction

with a Cahn TGA balance. Off-gas products will be collected and analyzed as a function of carbonization temperature using different heating rates.

4. References

1. K. Ouchi, "Infrared Study of Structural Changes During the Pyrolysis of a Phenol-Formaldehyde Resin." **Carbon**. Vol.4, 1966, pp. 59-66.
2. W.M. Jackson and R.T. Conley, "High Temperature Oxidative Degradation of Phenol-Formaldehyde Polycondensates." **Journal of Applied Polymer Science**. Vol.8, 1964, pp. 2163-2193.
3. J.A. Parker and E.L. Winkler, "The Effects of Molecular Structure on the Thermochemical Properties of Phenolics and Related Polymers", NASA TR R-276, Ames Research Center.
4. S.C. Brown, S.T. Bhe, K.T. Clemons, H.F. Reese and T.N. Sreenivasan, "Process Science for 2D Carbon-Carbon Exit Cones", AFWAL-TR-86-4028, Aerojet Strategic Propulsion Company, July 1986
5. G.M. Jenkins, K. Kawamura and L.L. Ban, "Formation and Structure of Polymeric Carbons." **Proceedings of the Royal Society of London, A**. Vol.327, 1972, pp. 501-517.
6. R.W. Seibold, "Carbonization of Phenolic Resin", MDAC Paper WD 2426, McDonnell Douglas Astronautics Company-West, January 1975
7. K. Kliner, personal communication, 1992.

A STUDY OF THE CHEMICAL VAPOR DEPOSITION
OF A SINGLE FILAMENT

Rose Marie Vecchione
Graduate Student
Department of Chemical Engineering

University of Dayton
300 College Park
Dayton, Ohio 45469

Final Report for:
Summer Research Program
Wright Laboratory

Sponsored by:
Air Force Office of Scientific Research
Wright Patterson Air Force Base, Dayton, Ohio

July 1992

A STUDY OF THE CHEMICAL VAPOR DEPOSITION
OF A SINGLE FILAMENT

Rose Marie Vecchione
Graduate Student
Department of Chemical Engineering
University of Dayton

Abstract

The application of carbon coatings onto ceramic monofilaments using chemical vapor deposition (CVD) was investigated. The conditions for depositing smooth, dense, uniform coatings of variable thicknesses onto ceramic monofilaments were determined. The microstructure of pyrolytic carbon coatings as a function of the process variables including gas flow rates, gas composition, CVD reactor temperature, and position within the CVD coating chamber was examined. Several experiments were conducted to determine the optimum conditions. The results of these experiments are interpreted and recommendations given.

A STUDY OF THE CHEMICAL VAPOR DEPOSITION OF A SINGLE FILAMENT

Rose Marie Vecchione

INTRODUCTION

The purpose of this research project was to investigate the chemical vapor deposition (CVD) of carbon onto silicon carbide monofilaments (Textron SCS0) using the CVD system at WL/MLLM. It was desired to determine the processing conditions necessary to produce smooth, uniform, dense coatings onto the SiC monofilaments. The process variables include gas flow rates, type of carbon source gas, gas composition, operating temperature, and fiber position within the coating furnace. The resulting coatings were examined by a scanning electron microscope (SEM).

DISCUSSION OF PROBLEM

In order to begin this project, a certain amount of literature research on the subject of CVD was necessary. My focal point, Dr. Malas, directed me to Dr. Alam whom I would be assisting throughout my 12 week term. He had a database containing bibliographies of various literature topics, the majority of which being CVD. The format for the bibliographies included an abstract. I was to complete this database and in the process learn about the research which had been done in this area. The software package used was ProCite version 1.34 for the Macintosh.

It was then necessary to become familiar with the CVD fiber coating apparatus located in building 655, WPAFB, OH. Figure 1

illustrates the CVD fiber coating apparatus. In general, there are four main components to the CVD fiber coater system (1): source gas inlet and delivery; heating source; fiber coating chamber/spooling assembly; vacuum generation and control; and exhaust treatment. The source gas control and delivery component delivers the stored gases to the reaction chamber in a controlled fashion.

The purpose of the heating source is to heat the substrate to the desired temperature. The furnace currently used is a wire wound, resistively heated vertical tube furnace capable of operating to 1600 degrees Celsius. The SCR power supply is driven by a PID microprocessor based controller responding to a mV input signal supplied by one thermocouple. The length of the furnace is 17 cm.

The purpose of the fiber coating chamber is to contain the heated section of fiber and coating gases. The coating chamber consists of a 20 mm ID fused silica tube to contain the fiber and gases inside the furnace plus flanges to hold and seal the tube and to provide gas and fiber inlet/outlet ports.

The fiber spooling assembly, which enables the continuous coating of long lengths of fibers, was not used. This is because all of the experiments to be run are static fiber depositions.

The vacuum generation and control system allows for coating processes to be run at any desired pressure and flow rate. The main components are a vacuum pump, a pressure measuring device, and a flow valve. A controller opens or closes the valve to maintain a set point pressure.

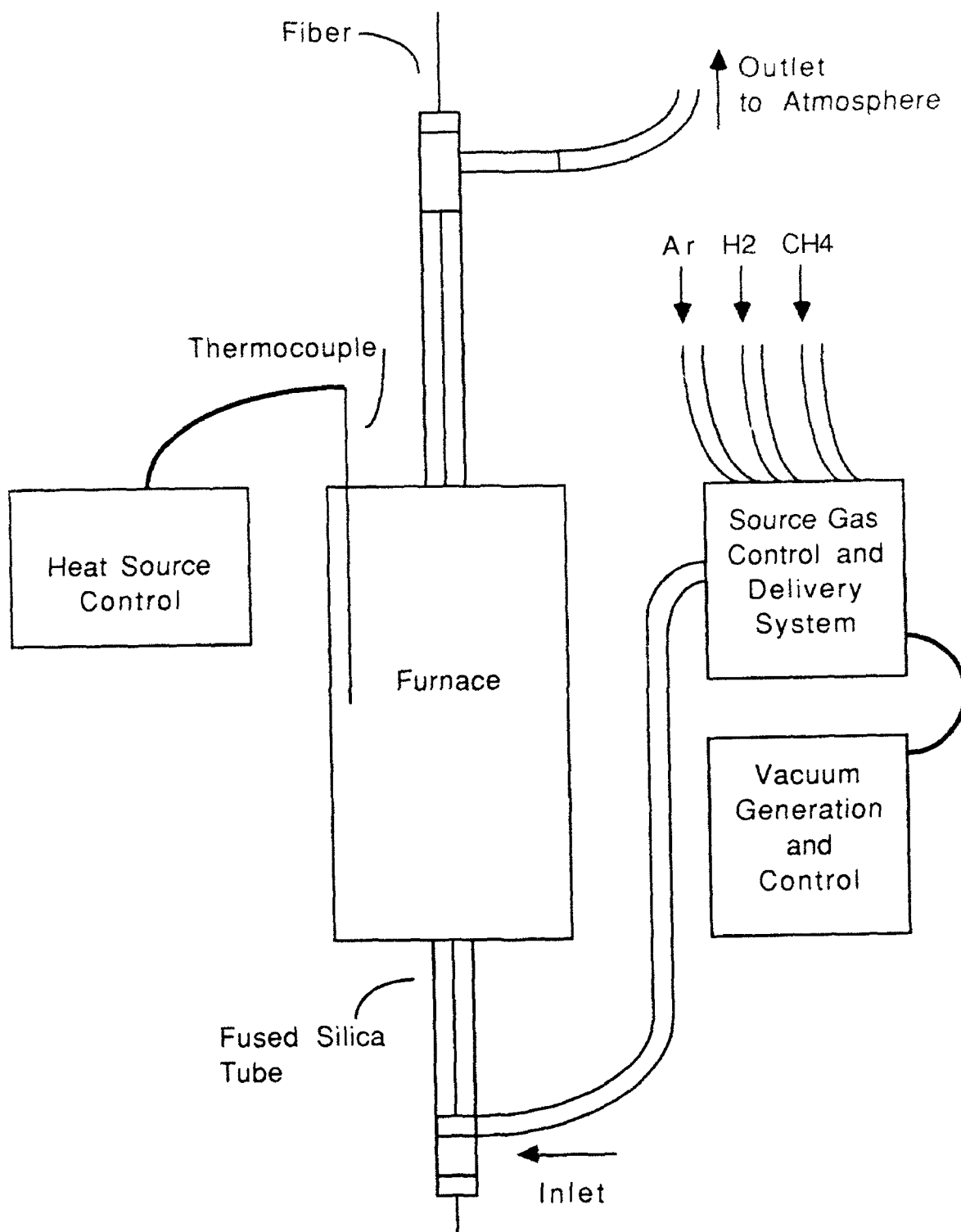


Figure 1. CVD Fiber Coating Apparatus.

The exhaust treatment system provided with the CVD fiber coater includes a solid type adsorbent that captures and neutralizes most acids and starting material gases and a hydrogen burner to controllably burn off the hydrogen carrier gas to prevent buildup and explosion in the room. The gas treatment system provided was not utilized because the gases used in the experiments were not corrosive or toxic in the amounts used. As a safety precaution, the whole CVD fiber coater is located within a hood and the exhaust gases were manually lit at the outlet in order to burn off any hydrogen present.

In order to examine the coated fibers it would be necessary for me to learn how to use the SEM (ETEC) in building 655. An SEM technician, Joe Williams, taught me how to use this machine.

METHODOLOGY

Six experiments were conducted. Each were static fiber depositions using CH_4 as the carbon source gas and H_2 as the carrier gas. The gases entered the bottom of the reactor and exited the top. Argon was used as the purge gas. The conditions for each run are tabulated in Table 1.

Table 1. Conditions used to study carbon deposition.

Date	Substrate	Set Pt. T deg C	CH_4 sccm	H_2 sccm	Time min
5/18/92	TC	1090	12	48	142
5/27/92	TC & 2f's	1090	11	44	133
6/8/92	2f's	1090	12	48	180
6/17/92	TC & 2f's	1098	14	42	197
7/1/92	2f's	1100	16	39	90
7/8/92	1f	1100	24	57	104

* TC=thermocouple, f=fiber

The first step for each run was to obtain a constant furnace temperature. The desired furnace temperature was 1100 degrees Celsius. Because of a known overshoot inherent in the controller the temperature for the furnace was set a little below or at 1100 degrees Celsius. The power input to the furnace was set at 2.6 to 2.7 Amps at the start of each run and was manually manipulated for the duration of the experiment. The manual manipulation was necessary in order to heat the furnace up in a reasonable amount of time. The furnace took on an average of three hours to heat up to the desired temperature with manipulation. Further manipulation of the power input was sometimes necessary even after reaching the desired temperature because of a lack of good control by the controller. During this heat-up period, Ar was used as a purge gas for the fiber coating chamber. Ar was set at 20 sccm.

Fifteen minutes before the start of the reaction of run 5/18/92, CH₄ was used to purge the lines and was then vented to the atmosphere. To begin the deposition, CH₄ was then allowed to enter the chamber along with H₂. The substrate used was a thermocouple wire with the junction positioned approximately at the middle of the length of the furnace. During the deposition, the furnace temperature went as high as 1117 degrees Celsius while the corresponding thermocouple temperature was 1140 degrees Celsius. When a thermocouple was used as a substrate in any of the runs there was always a noticeable difference in temperature between that of the thermocouple and that of the furnace.

Two hours and six minutes before the start of the reaction of run 5/27/92 and twenty minutes before the start of the reaction of run 6/8/92, H_2 was used as a purge gas for the reactor chamber. The thermocouple substrate broke before the start of the reaction of run 5/27/92.

H_2 was used as a purge for the chamber and CH_4 was allowed through the lines (separate from H_2) and was vented to the atmosphere. This was done 2 hours and twenty minutes before runs 6/17/92 and 7/1/92 and was done 3 hours and 3 minutes before run 7/8/92. During run 6/17/92 the reactor began to overheat (up to 1113 degrees Celsius) so H_2 and CH_4 were shut off and Ar was used as a purge until the temperature of the furnace went down and became stable at 1103 degrees Celsius. The reaction was then continued by letting H_2 and CH_4 flow back into the chamber.

One can see from Table 1 that the first three experiments were run for 142 minutes, 133 minutes, and 180 minutes with 20% CH_4 . The first and third of these had total gas flow rates set at 60 sccm while the second had its total gas flow rate set at 55 sccm. The fourth experiment was run for 197 minutes with 25% CH_4 and a total gas flow rate of 56 sccm. The fifth and sixth experiments were run for 90 minutes each with 30% CH_4 and total gas flow rates of 54 sccm and 81 sccm, respectively.

After the deposition and after the furnace cooled down, the substrate was removed from the chamber and prepared for observation. The substrate was cut so that the top of the substrate corresponded to the

top of the furnace and the bottom of the substrate corresponded to the bottom of the furnace. The resulting substrate was then prepared as follows: The thermocouple from run 5/18/92 was cut into 1 inch intervals leaving 1/2 inch near the junction providing eleven samples for observation with the SEM; The two fibers from run 5/27/92 were each cut into 2 1/8 inch intervals providing 8 samples per substrate; The two fibers from run 6/8/92 were each cut into 1 1/2 inch intervals providing twelve samples per fiber; The thermocouple and two fibers from run 6/17/92 were each cut into 4 inch intervals providing 5 samples per substrate; The two fibers from run 7/1/92 were each cut in 2 inch intervals providing 9 samples per fiber; The fiber from run 7/8/92 was cut into 2 inch intervals providing 9 samples for observation. The substrates were cut with a wire cutter and the size of each sample was 1/4 inch.

RESULTS

The results of the coating experiments are detailed below and supported by SEM photomicrographs of the coated fibers as well as similar previously performed experiments (2). In general, the resulting coating on the thermocouple of run 5/18/92 was sooty. The wire appeared bare 0, 2, and 4 inches above the bottom of the furnace. At 6 inches there was mostly soot with a little cauliflower structure buried in. Samples taken 8, 9, 10, and 11 inches above the bottom showed a thick cauliflower coating. At 10 inches the coating appeared to be approximately 10 micrometers. The sample at 10 inches, magnified 300

times, can be seen in Figure 2-A. The coating became a fluff at 13 inches and then subsided completely at 15 and 17 inches.

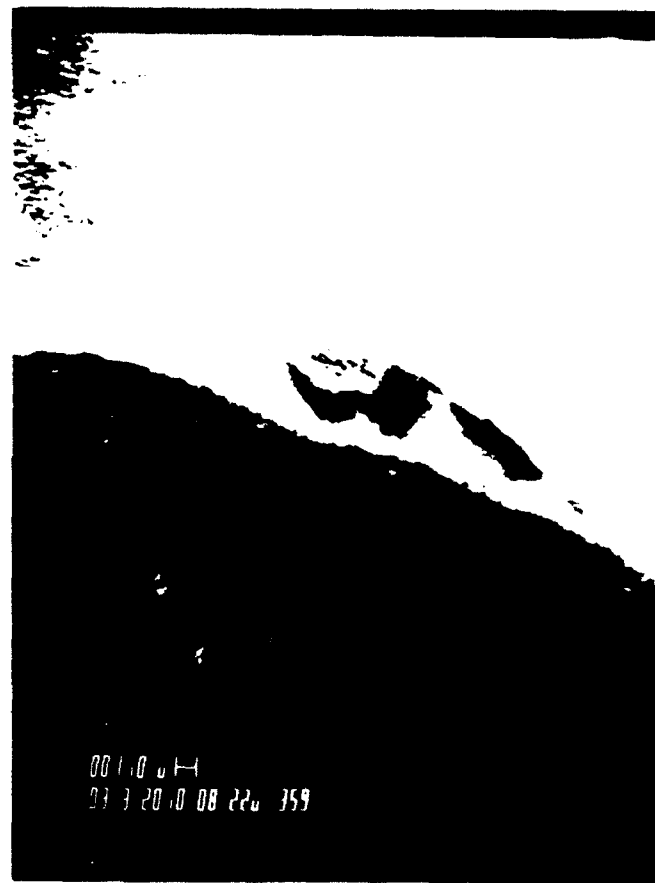
The resulting coatings on the fibers of run 5/27/92 were barely visible with the SEM. Both fibers showed similar coatings. Only a very thin coating at 9 and 11 1/8 inches above the bottom of the furnace could be detected. At 9 inches the coating appeared to be 1 micrometer thick and at 11 1/8 inches the coating appeared to be 1/2 micrometer thick. The coating was dense and uniform. The sample at 11 1/8 inches, magnified 3000 times, can be seen in Figure 2-B.

The coatings on the fibers of run 6/8/92 were thin, dense, and uniform. Similar coatings appeared on each. Nothing was visible at 0, 1 1/2, 3, 4 1/2, and 6 inches from the furnace bottom. There was a very thin coating at 7 1/2 inches and a visible coating of approximately 1 micrometer appeared at 9, 10 1/2, and 12 inches. A thin coating at 13 1/2 inches was also visible. Streaks of soot appeared at 15 inches and the fiber appeared bare at 16 1/2 inches from the bottom of the furnace. Figure 2-C shows the fiber at 9 inches magnified 400 times.

The coating on the thermocouple of run 6/17/92 generally appeared sooty. Nothing appeared at 0 inches, the coating was very thick and sooty at 4, 8, and 12 inches, and a fluffy soot was observable at 16 inches. The coatings on the fibers were thin, dense, and uniform. Both fibers showed similar coatings. No coating was observable at 0 and 4 inches. At 8 inches, magnified 500 times in Figure 2-D, the coating appeared to be approximately 1 micrometer thick. A thin coating was visible at 12 inches and only some soot was visible at 16 inches.



A



B

Figure 2



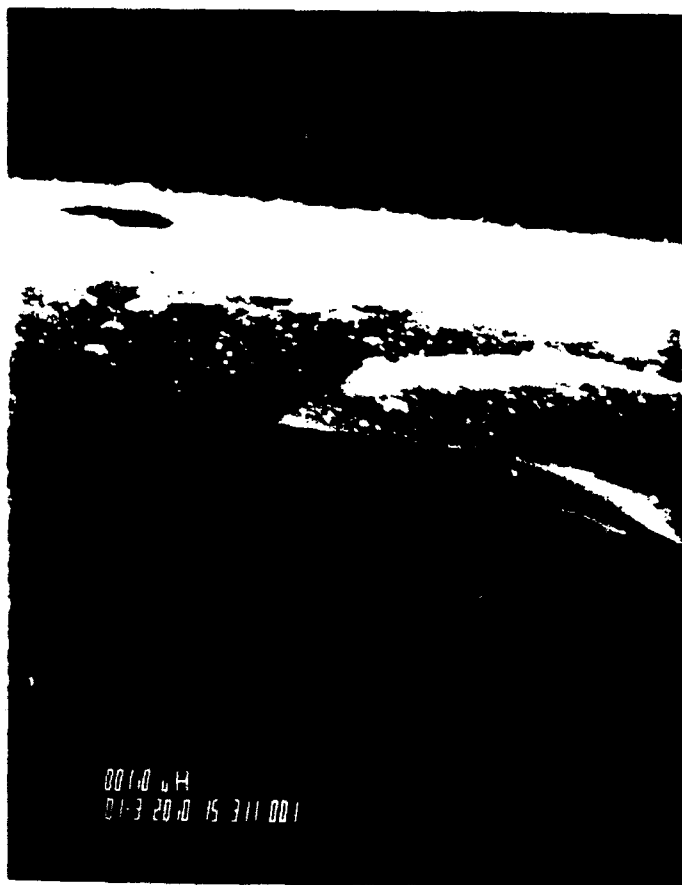
C



D

Figure 2 (cont.)

The coatings resulting from run 7/1/92 appeared uniform and dense, as did the previous runs, but possessed slightly thicker coatings. Both fibers had similar coatings. Nothing appeared 0, 2, and 4 inches from the bottom of the furnace. A very thin coating was visible at 6 inches and a 2-3 micrometer coating was visible at 8 and 10 inches. The sample



A



B

Figure 3

at 10 inches, magnified 1000 times, is shown in Figure 3-A. A very thin, fluffy, 1/2 micrometer coating was visible at 12 inches above the bottom of the furnace. Nothing was visible at 14 and 16 inches.

The coating resulting from run 7/8/92 was similar to that obtained in run 7/1/92 only occurring 2 inches higher. Nothing was visible on the fiber at 0, 2, 4, and 6 inches from the bottom of the furnace. From 8 to 17 inches a coating was visible. At 8 inches the coating appeared very thin. The coating at 10 inches, was much thicker, approximately 3-4 micrometers, and was actually falling off of the fiber. This can be seen in Figure 3-B where the sample is magnified 100 times. At 12 inches a 2 micrometer coating was visible and at 14 inches thin, dense, uniform coating of 1 micrometer was visible. Towards the top of the furnace, at 16 inches, a very thin coating was visible.

CONCLUSION

It is anticipated that the conditions used on 7/1/92 will be sufficient to produce coatings of the desired morphology. The conditions which tend to affect the resulting coating the most are the reactor temperature, the total reactant gas flow rate, and the percent of carbon source gas used.

1100 degrees Celsius appears to be a good set point temperature for the process. For future experiments, one could try running it at even higher temperatures. It has been shown that lower temperatures, even temperatures close to 1100 degrees Celsius, such as 1050 degrees Celsius, do not result in good coatings. In fact, the coatings are not even visible with the SEM (2).

The total reactant gas flow rate seems to have an effect on where the region which produces the best coatings will be. This can best be seen from runs 7/1/92 and 7/8/92. All conditions were identical for these two runs with the exception of the total reactant gas flow rate. It was 54 sccm for run 7/1/92 and 81 sccm for run 7/8/92. The region where the best coating occurred moved up from 8 inches to 10 inches. For the first five runs where the total reactant gas flow rate varied from 54 to 60 sccm the region of best coating occurred, on the average, between 8 and 12 inches above the bottom of the furnace.

The percent of carbon source gas used seems to have an effect on the thickness of coating produced. 30% CH₄ seemed to give the thickest coating. Where the thickness of the coatings resulting from 20% CH₄ and 25% CH₄ varied only from 1/2 to 1 micrometer, the thickness of the coatings resulting from using 30% CH₄ varied from 2-3 micrometers. For further experimentation, a higher percentage of carbon source gas could be used.

Switching to a three zone furnace might reduce the variability of the deposit morphology over the length of the furnace. This will provide a much more uniform temperature profile as well as a longer region of constant temperature.

REFERENCES

1. CVD Fiber Coater Description and Operation Instructions (provided with the CVD Fiber Coater).
2. "Carbon Coating of Ceramic Monofilaments", Final Technical Report for Contract F33615-89-C-5609, Task 91, Jason R. Guth.

A SWITCHED RELUCTANCE MOTOR DRIVE USING MOSFETS,
HCTL-1100, and MC6802 MICROPROCESSOR

Shy-Shenq P. Liou
Assistant Professor

Lawrence Vo
Graduate Student

Division of Engineering

San Francisco State University
1600 Holloway Avenue
San Francisco, CA 94132

Final Report for:
AFOSR Summer Research Program
Wright Patterson Laboratory

Sponsored by:
Air Force Office of Scientific Research
Bolling Air Force Base, Washington, D.C.

September 1992

A Switched Reluctance Motor Drive Using MOSFETs, HCTL-1100, and MC6802 Microprocessor

Shy-Shenq P. Liou
Assistant Professor

Lawrence Vo
Graduate Student

Division of Engineering
San Francisco State University

Abstract

A switched reluctance motor drive is designed, built, and tested using the MOSFETs as the power switches. The control function of this drive is done by a general purpose motion control chip HCTL-1100 from Hewlett Packard. The interface between the HCTL-1100 and the user is through a Motorola microprocessor MC-6802. A bang-bang current control circuit is also built into the drive to limit the motor current to be less than or equal to the rated motor current. Simple assembly program can enable the user to input the command velocity, command position, and command profile to the HCTL-1100 motion control chip. The test run of this drive is very successful.