

AFRL-PR-WP-TR-2004-2135

**CRYO POWER AND HEAT
TRANSFER**

**L.C. Chow, M. Bass, J. Du, Y. Lin, T. Chung, C. Walsh,
B. Glassman, and K. Sundaram**

**University of Central Florida
Department of Mechanical, Materials and Aerospace Engineering
Orlando, FL 32816-2450**



SEPTEMBER 2004

Final Report for 01 August 1999 – 31 July 2004

Approved for public release; distribution is unlimited.

STINFO FINAL REPORT

**PROPULSION DIRECTORATE
AIR FORCE MATERIEL COMMAND
AIR FORCE RESEARCH LABORATORY
WRIGHT-PATTERSON AIR FORCE BASE, OH 45433-7251**

NOTICE

USING GOVERNMENT DRAWINGS, SPECIFICATIONS, OR OTHER DATA INCLUDED IN THIS DOCUMENT FOR ANY PURPOSE OTHER THAN GOVERNMENT PROCUREMENT DOES NOT IN ANY WAY OBLIGATE THE U.S. GOVERNMENT. THE FACT THAT THE GOVERNMENT FORMULATED OR SUPPLIED THE DRAWINGS, SPECIFICATIONS, OR OTHER DATA DOES NOT LICENSE THE HOLDER OR ANY OTHER PERSON OR CORPORATION; OR CONVEY ANY RIGHTS OR PERMISSION TO MANUFACTURE, USE, OR SELL ANY PATENTED INVENTION THAT MAY RELATE TO THEM.

THIS REPORT HAS BEEN REVIEWED BY THE AIR FORCE RESEARCH LABORATORY WRIGHT SITE OFFICE OF PUBLIC AFFAIRS (AFRL/WS/PA) AND IS RELEASABLE TO THE NATIONAL TECHNICAL INFORMATION SERVICE (NTIS). AT NTIS, IT WILL BE AVAILABLE TO THE GENERAL PUBLIC, INCLUDING FOREIGN NATIONALS.

THIS TECHNICAL REPORT HAS BEEN REVIEWED AND IS APPROVED FOR PUBLICATION.

/s/

BRIAN D. DONOVAN
Mechanical Engineer
Plans & Analysis Branch

/s/

BRIAN G. HAGER
Chief
Plans & Analysis Branch

/s/

CYNTHIA A. OBRINGER
Deputy Chief
Power Division

Do not return copies of this report unless contractual obligations or notice on a specific document require its return.

REPORT DOCUMENTATION PAGE					Form Approved OMB No. 0704-0188	
The public reporting burden for this collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Department of Defense, Washington Headquarters Services, Directorate for Information Operations and Reports (0704-0188), 1215 Jefferson Davis Highway, Suite 1204, Arlington, VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to any penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number. PLEASE DO NOT RETURN YOUR FORM TO THE ABOVE ADDRESS.						
1. REPORT DATE (DD-MM-YY) September 2004		2. REPORT TYPE Final		3. DATES COVERED (From - To) 08/01/1999 – 07/31/2004		
4. TITLE AND SUBTITLE CRYO POWER AND HEAT TRANSFER				5a. CONTRACT NUMBER F33615-96-C-2681		
				5b. GRANT NUMBER		
				5c. PROGRAM ELEMENT NUMBER 62173C		
6. AUTHOR(S) L.C. Chow, M. Bass, J. Du, Y. Lin, T. Chung, C. Walsh, B. Glassman, and K. Sundaram				5d. PROJECT NUMBER 1651		
				5e. TASK NUMBER 01		
				5f. WORK UNIT NUMBER 06		
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) University of Central Florida Department of Mechanical, Materials and Aerospace Engineering Orlando, FL 32816-2450				8. PERFORMING ORGANIZATION REPORT NUMBER		
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES) Propulsion Directorate Air Force Research Laboratory Air Force Materiel Command Wright-Patterson AFB, OH 45433-7251				10. SPONSORING/MONITORING AGENCY ACRONYM(S) AFRL/PRPA		
				11. SPONSORING/MONITORING AGENCY REPORT NUMBER(S) AFRL-PR-WP-TR-2004-2135		
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution is unlimited.						
13. SUPPLEMENTARY NOTES Report contains color.						
14. ABSTRACT Over the past eight years, researchers at the University of Central Florida have studied high flux heat transfer by flow boiling, pool boiling and spray cooling. These techniques are targeted for cooling high power density devices such as power MOSFETs and diode lasers arrays at both cryogenic and room temperatures. Seven tasks were included. The results of flow boiling and pool boiling with liquid nitrogen were included in earlier annual reports for this contract and therefore are not reported here. This final report includes a study on the simulation and testing of power MOSFETs, thermal management of diode laser arrays, development of compact spray nozzles, and fluid management of spray cooling with multiple nozzles for large surface areas. Several examples of the significant achievements include the development of a microstereolithography apparatus so that micro spray nozzles can be made with great precision; a new way of packaging diode laser arrays which is vastly superior to the state of the art; and a unique and effective way to solve the flooding problems associated with multiple nozzle spray cooling.						
15. SUBJECT TERMS						
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT: SAR	18. NUMBER OF PAGES 158	19a. NAME OF RESPONSIBLE PERSON (Monitor) Brian Donovan 19b. TELEPHONE NUMBER (Include Area Code) (937) 255-5735	
a. REPORT Unclassified	b. ABSTRACT Unclassified	c. THIS PAGE Unclassified				

Table of Contents

TABLE OF CONTENTS	iii
LIST OF FIGURES	vi
LIST OF TABLES	x
NOMENCLATURE	xi
1. INTRODUCTION	1
2. FABRICATION AND SIMULATION OF THE POWER MOSFET	
2.1 INTRODUCTION.....	2
2.2 POWER MOSFETS	4
DIFFERENT TYPES OF POWER MOSFETS	4
MERITS OF DMOSFET STRUCTURE.....	7
OPERATION OF POWER MOSFET	8
POWER MOSFET DOPING PROFILE	9
THRESHOLD VOLTAGE	10
ON-STATE CONDUCTION	10
2.3 PROCESS SIMULATOR	15
PROCESS SIMULATOR (DIOS) INTRODUCTION	15
PROCESS SIMULATION	16
DESCRIPTION OF COMMANDS USED	16
2.4 DMOSFET FABRICATION AND DESIGN CONSIDERATION	26
DESIGN OF MASKS FOR DMOS	26
POWER MOSFET FABRICATION	29
OPEN WINDOWS USING FIRST MASK	29
ION IMPLANTATION	29
BORON DIFFUSION	29
PHOSPHORUS PREDEP	30
GATE OXIDE	30
CONTACT WINDOWS AND METALLIZATION	30
BORON IMPLANTATION AND DIFFUSION	30
2.5 RESULTS	31
TEMPERATURE EFFECTS	31
EXPERIMENTAL RESULTS	32
POSSIBLE REASON FOR IN-OPERABILITY OF DMOSFET	34
2.6 CONCLUSION	35
3. RESEARCH AND DEVELOPMENT OF COMPACT SPRAY NOZZLE	
3.1 INTRODUCTION.....	36
3.2 LITERATURE REVIEW	39
SCANNING METHOD	42
INTEGRAL PROCESS METHOD	41
3.3 METHODOLOGY	45
SYSTEM REQUIREMENTS	45
LOCATION	46
SPATIAL LIGHT MODULATOR	46
LIGHT SOURCE	48
OPTICS AND OPTO-MECHANICS	49
LIGHT ENGINE	50
TRANSLATION STAGES.....	52
VAT SETUPS	53
VAT BOTTOM COATING	54

PHOTOPOLYMERS	55
COMPUTER CONTROL SYSTEM	56
FINAL DESIGN SETUP	57
3.4 RESULTS	60
RANDOM SHRINKAGE	60
LASER POWER CURVE	61
CERASET CURE DEPTHS	63
SINGLE LAYER BUILDS	65
MULTIPLE LAYER PARTS (STANDARD BUILD)	67
ADHESION FORCES	69
MULTIPLE LAYER PARTS (INVERTED BUILD)	69
MANUFACTURING ERROR	70
3.5 CONCLUSION	71
4. A FLUID MANAGEMENT SYSTEM FOR A MULTIPLE NOZZLE ARRAY SPRAY COOLER..	
4.1 INTRODUCTION	73
DESIGN PROBLEM	73
DESIGN & SOLUTION	74
4.2 EXPERIMENT SETUP	77
FLUID DYNAMIC ANALYSIS & SETUP	77
SUCTION EFFECTIVENESS TESTING	81
THERMAL DESIGN & FEM ANALYSIS	81
4.3 THERMAL TESTING PROCEDURES	84
THERMAL CALCULATIONS & RESULTS	86
THERMAL RESULTS	87
UNCERTAINTIES	88
4.4 CHAMBER DESIGN	88
4.5 CONCLUSION	90
5. THERMAL MANAGEMENT OF DIODE LASER ARRAYS	
5.1 THERMAL MANAGEMENT AND SPRAY COOLING	92
THERMAL CONDUCTION AND CONVECTION	92
SPRAY COOLING	93
SPRAY COOLED DIODE LASER ARRAY	95
SPRAY COOLED STACK PACKAGING RESULTS AND ANALYSIS	100
5.2 BEAM CONTROL PRISM PACKAGING DESIGN (BCPP)	101
BASIC DESIGN CONCEPT	101
OPTICAL ISSUES	104
1. CHANNEL DEPTH EQUATION	104
2. PITCH EQUATION	104
3. N-ELECTRODE NON-BLOCKING CONDITION	104
4. OUTPUT BEAM RADIUS EQUATION	104
5. BCP FOCAL LENGTH	104
ELECTRIC CONNECTION	105
THERMAL ISSUES	107
MECHANICAL ISSUES	108
EXPERIMENTAL DESIGN	108
BCPP TEMPERATURE DISTRIBUTION FEA RESULTS	112
EXPERIMENTAL RESULTS	116
5.3 CONCLUSION	121
6. CONCLUDING REMARKS	122
APPENDICES	
A. CALCULATIONS FOR ION IMPLANTATION	124
B. FABRICATION PROCESS PROCEDURE	127
C. COMMAND FILE OF DIOS	133

D. BEAM SHAPING DESIGN	135
REFERENCES	138

List of Figures

Figure 2.1: An n-channel MOSFET	4
Figure 2.2: An Vertical V-groove MOSFET	5
Figure 2.3: An truncated V-groove vertical MOSFET	6
Figure 2.4: An vertical DMOS structure	6
Figure 2.5: The UMOSFET Structure	7
Figure 2.6: Voltage Current Limitations of MOSFET, BJT and Thyristor	7
Figure 2.7: Typical doping profile of a power MOSFET	9
Figure 2.8: The resistance components within DMOSFET structure	11
Figure 2.9: Cross section of DMOSFET for the specific on resistance analysis	12
Figure 2.10: Boron (p-region) diffusion in silicon epilayer	22
Figure 2.11: Boron (p-region) and Phosphorus (n ⁺ region) diffusion	23
Figure 2.12: Full device structure	25
Figure 2.13: Full device structure with triangle option <i>off</i>	26
Figure 2.14: (a) Mask –1 for channel length of 15 μ	27
Figure 2.14: (b) Mask-1 for channel length of 10 μ	27
Figure 2.14: (c) Mask-1 for channel length of 5 μ	28
Figure 2.15: Overhead view of Mask-2	28
Figure 2.16: Overhead view of Mask-3	28
Figure 2.17: Overhead View of Mask-4	29
Figure 2.18: Effect of Temperature on Device Characteristics	32
Figure 2.19: Characteristic of MOSFET at room temperature	35
Figure 2.20: Characteristic of MOSFET at liquid Nitrogen temperature	35
Figure 3.1: STL file estimation of a curve	37
Figure 3.2: Typical STL file	37
Figure 3.3: 3D Systems SLA-250	38
Figure 3.4: IH Process	39
Figure 3.5: IH Process sample part	40
Figure 3.6: Two-photon process	40
Figure 3.7: Two-photon sample part	41
Figure 3.8: University of Sussex MSL process	42
Figure 3.9: University of Sussex sample part	43
Figure 3.10: EPFL MSL setup	44
Figure 3.11: EPFL sample part	44
Figure 3.12: Exploded view of spray nozzle	45

Figure 3.13: Digital Micromirror Device	47
Figure 3.14: SEM photograph of DMD	47
Figure 3.15: Lens used for intensity apodization	49
Figure 3.16: Two lens beam expander	50
Figure 3.17: Imaging optics for LCD	51
Figure 3.18: Imaging system for DMD.....	51
Figure 3.19: TIR prism setup	52
Figure 3.20: Example of screen stamping	52
Figure 3.21: Sweeping method for making layer	53
Figure 3.22: Waiting method for creating layer thickness	54
Figure 3.23: Inverted build layer creation	54
Figure 3.24: The structure of Ceraset™ SN.....	55
Figure 3.25: Computer communication for MSL system	56
Figure 3.26: MSL software main menu	56
Figure 3.27: Manual test operations GUI	57
Figure 3.28: ProE Drawing of Final Design	59
Figure 3.29: Photo of Final Design	59
Figure 3.30: SEM Photo of Indentation Mark	60
Figure 3.31: Low vs. High repetition rate	61
Figure 3.32: Laser power as a function of Repetition Rate	62
Figure 3.33: Relationship between internal and external power meters	63
Figure 3.34: Ceraset Cure Depths	64
Figure 3.35: Commercial 7120 Cure Depths	65
Figure 3.36: Static mask design	66
Figure 3.37: Ceraset single layer sample part	66
Figure 3.38: Ceraset single layer from a spray nozzle	67
Figure 3.39: Multiple layer part standard build	68
Figure 3.40: Cross section view of layers using standard build	68
Figure 3.41: Adhesion force limitation on geometry	69
Figure 3.42: Nozzle section compared to standard ruler	70
Figure 3.43: Beam expander design	71
Figure 4.1: Single spray: horizontal	73
Figure 4.2: Multiple spray interaction: horizontal	73
Figure 4.3: Fluid build up points	74
Figure 4.4: Fluid flow observed from the side of the spray cooler with the un-modified siphons, flooding between spray cones reduced	75
Figure 4.5: Spray Nozzle Configuration	75

Figure 4.6: Overall siphon placement	76
Figure 4.7: Overall spray cooler design	76
Figure 4.8: Experimental setup	77
Figure 4.9: Observer flow from bottom of grooved plate: un-modified siphons used	78
Figure 4.10: Modified siphon designs	79
Figure 4.11: Fluid dynamics visualization around the base of the siphons	79
Figure 4.12: Siphons placement and siphon type Refer to Figure 10 for type of siphon	80
Figure 4.13: Suction Effectiveness for a 16 Nozzle Array: No Heat used during tests	81
Figure 4.14: Heater design with thermocouple locations	82
Figure 4.15: Sectional temperature profile of heater	83
Figure 4.16: Experimental Setup for Multiple Nozzle Spray cooler	85
Figure 4.17: Vacuums	85
Figure 4.18: Custom Lab view Layout	86
Figure 4.19: Q vs. $T_w - T_{sat}$ 20 PSI Head Pressure and Flow Rate of 9.8 Liters/min	87
Figure 4.20: Q vs. $T_w - T_{sat}$ for 30 PSI head at Flow rate of 11.2 Liter/min	88
Figure 4.21: Dimension of manufactured Chamber	89
Figure 4.22: Physical Location of Center of Gravity	90
Figure 4.23: Front view of manufactured chamber	90
Figure 5.1: Coolant flow rate comparison of different cooling methods	93
Figure 5.2: The system pressure versus water boiling temperature	94
Figure 5.3: Low pressure water and ammonia spray cooling experimental results	94
Figure 5.4: The arrangement of the diode array and the spray nozzle	96
Figure 5.5: The spray cooling experimental set up with the cooling coil design	97
Figure 5.6: The spray cooling experimental set up with the mist cooling design	97
Figure 5.7: Diode laser array input and output power versus current	98
Figure 5.8: Average spectrum at different input current levels	99
Figure 5.9: Wavelength and temperature of different emitters at different driving currents	99
Figure 5.10: Heat flux versus average diode temperature	100
Figure 5.11: FEA result of temperature distribution of diode array stack with 2.7 mm and 1.5 mm thick copper spacers	101
Figure 5.12: The Basic structure of the BCPP	103
Figure 5.13: Parameters of the BCPP	103
Figure 5.14: Output beam radius, a , versus BCP index of refraction for different ϕ angles ...	104
Figure 5.15: Current directions in the p- and n-electrodes	105
Figure 5.16: Packing example for larger array with “vertical” electrical connections	106
Figure 5.17: Packing example for larger array with “horizontal” electrical connections	106
Figure 5.18: BCPP with shaped n-electrodes	107

Figure 5.19: Sandwiched copper-BeO-copper BCPP substrate design	109
Figure 5.20: Detail dimensions of BCPP substrate employing BCP stand concept	109
Figure 5.21: Detail dimensions of BCPP substrate employing a dual BCP	110
Figure 5.22: BCP channel side views of different cutting methods	110
Figure 5.23: The n-electrode surface plot after cutting the channel	111
Figure 5.24: The n-electrode design	112
Figure 5.25: FEA computed temperature distribution of a single emitter in a 60W bar	114
Figure 5.26: FEA computed temperature distribution of a single 60 W bar	114
Figure 5.27: FEA computed temperature distribution of an entire BCPP array built with of 10 60W bars	115
Figure 5.28: FEA computed temperature distribution of a single diode bar for different waste heat levels	115
Figure 5.29: FEA computed temperature distribution in the emitter plane of the BCPP array for different waste heat levels	116
Figure 5.30: Top view of the BCPP array	116
Figure 5.31: Front view of the BCPP array	117
Figure 5.32: IR image of the BCPP array	118
Figure 5.33: Output spectrum of the BCPP diode laser array for different current levels	118
Figure 5.34: P-I curve and efficiency of the BCPP array	119
Figure 5.35: Experimental setup for measuring the output spectrum of each emitter	120
Figure 5.36: Measured temperature distribution of a BCPP diode laser array	121

List of Tables

Table 2.1: Values of different parameters used in fabricating the wafers	32
Table 2.2: Values of resistances for different devices on a fabricated wafer 1	33
Table 2.3: Values of resistances of devices on fabricated wafer 2	33
Table 2.4: Values of resistances of devices on fabricated wafer 3	34
Table 2.5: Values of resistances of devices on fabricated wafer 4	34
Table 3.1: DSM Somos Properties	56
Table 3.2: Random shrinkage numbers	62
Table 5.1: Coherent 808 nm B1-40C-19-30-A diode specification	95
Table 5.2: BCP positioning tolerance	108
Table 5.3: LaserTel LT1200-40W and 1300-60W diode laser specifications	113

Nomenclature

Micro-Electrical and Mechanical Systems (MEMS)

- The integration of mechanical systems with microelectronics

Stereolithography Apparatus (SLA)

- Commercial rapid prototyping machine from 3D Systems

Aspect Ratio

- The length and/or width to height ratio of a part

Microstereolithography (MSL)

- Process based on SLA used to make micro-parts

Standard Tessellation Language (STL)

- File format used to simplify solid models to be built in a rapid prototyping machine

Vector-by-vector

- MSL process technique where each layer is made using a small laser beam and hatching in the layer features

Ultraviolet (UV)

- Light with a wavelength shorter than 400 nm

Raster Scanning

- Same as vector-by-vector technique for making layer patterns

XYZ Stage

- Motorized translation stages used to position the MSL system

Critical Exposure

- The amount of energy needed to polymerize the liquid photopolymer

Spatial Light Modulator (SLM)

- A device used to carry an optical image

Liquid Crystal Display (LCD)

- An array of liquid crystals that are used for light modulation

Fill Ratio

- The ratio of the area of a pixel to the area of the gap between pixels

TEM₀₀

- Mode of a laser with a Gaussian shaped intensity profile

Polymer-Derived-Ceramic (PDC)

- A polymer which when exposed to high temperatures, converts to a ceramic material

Digital Micromirror Device (DMD)

- A MEMS device that consists of tiny mirror used for light modulation

Teflon® AF

- A fluoropolymer version of DuPont's Teflon® material

Photopolymer

- A liquid polymer that solidifies when exposed to certain wavelength light

Pyrolyzed

- The process of raising the temperature of the PDC and converting it to a ceramic structure

Static Mask

- A glass plate with a pattern used to make sample parts

Microdisplay

– LCD's and DMD's used in consumer electronics

A_{out}

– Outside area of the heater (cm^2)

A_h

– Top pedestal heater area (cm^2)

E_{in}

– Electrical Input (W)

Eff_{suc}

– %Effectiveness of suction system

E_{cooled}

– Energy absorbed by spray cooling (W)

E_{loss}

– Energy loss to the shell of the heater (W)

k_{inter}

– Interpolated value of thermal conductivity (W/cm-K)

q''

– Heat flux (W/cm^2)

X_{1-2}

– Distance between T.C.1 and T.C. 2 = 10.16mm

X_{1-3}

– 17.78mm

X_{2-3}

– 7.62 mm

X_{1-w}

– Distance T.C. is from the cooled surface 6.08 mm

T_w

– Cooled surface temperature ($^{\circ}C$)

T_{b1-b3}

– T.C. temperature on right side 1 -3 ($^{\circ}C$)

T_{1-3}

– Temperature of thermocouples 1 through 3 ($^{\circ}C$)

T.C.

–Thermocouple

$\dot{Volume}_{edges}$

– Volume flow rate over the edges of the spray cooler (L/min)

1. INTRODUCTION

This final report gives a detailed description of the four remaining tasks (Tasks 7 – 10) in the Air Force Contract to University of Central Florida under Contract Number F33615-96-C-2681. The other tasks were already reported in earlier annual reports and are therefore not repeated here. The four tasks are organized and reported as Chapter 2 through Chapter 5 as shown below:

Chapter 2: Task 7 – Fabrication and Simulation of the Power MOSFET

Chapter 3: Task 9 – Research and Development of Compact Spray Nozzle

Chapter 4: Task 10 – Fluid Management System for a Multiple Nozzle Array Spray Cooler

Chapter 5: Task 8 – Thermal Management of Diode Laser Arrays

The rest of the report includes some concluding remarks and Appendices and References.

2. FABRICATION AND SIMULATION OF THE POWER MOSFETS

2.1 INTRODUCTION

The applications for power semiconductor devices are quite diverse. The power ratings extend over a tremendous range from the level of 100 W at microwave frequencies to 100 MW at low frequencies. An ideal power switch used for power conditioning should therefore be capable of handling high currents and voltages, and be able to switch at a high speed.

Another approach to classification of the applications for power semiconductor devices is in terms of the voltage and current ratings that the device must satisfy for each application. The first category requires relatively low breakdown voltage (less than 100 volts) but requires high current handling capability. Two examples are automotive electronics and switch mode power supplies. The second group of applications lies along a trajectory of increasing breakdown voltage and current handling capability, i.e. a substantial increase in power handling capability. At lower power levels there are applications such as display drives. These applications are being served by monolithic power integrated circuits containing multiple high voltage drive transistors. At higher power levels, smart power technology is being developed to provide monolithic chips for lamp ballasts and fractional horsepower motor control. An ideal power device must be able to control the flow of power to loads with zero power dissipation (Baliga, 1995a).

Many of the merits of MOSFETS are as important in devices designed for power application as they are in those designed for large-scale integration. Of course, the decision as to what constitutes a power device is quite arbitrary. For the purpose of definition we shall apply this term to any device capable of switching at least 1 A (Grant, 1989a).

Power metal-oxide-semiconductor field-effect transistors (MOSFETs) are the most commercially advanced devices. These devices have evolved from MOS integrated –circuit technology. Prior to the development of power MOSFETs, the only device available for high-speed, medium-power applications was the power bipolar transistor. Despite the attractive power ratings achieved for bipolar transistors, there exist several fundamental drawbacks in their operating characteristics. First, the bipolar transistor is a current-controlled device. A large base drive current, typically one-fifth to one tenth of the collector current, is required to maintain them in the on-state. Even larger reverse base drive currents are necessary for obtaining high-speed turn off. These characteristics make the base drive circuitry complex and expensive. The bipolar transistor is also vulnerable to a second breakdown failure mode under the simultaneous application of a high current and voltage to the devices commonly required in inductive power circuits. Furthermore, it is difficult to parallel these devices. The forward voltage drop in bipolar transistors decreases with increasing temperature. This promotes diversion of the current to a single device unless emitter ballasting schemes are utilized.

The power MOSFET was developed to solve the performance limitations experienced with power bipolar transistors. In this device, the control signal is applied to a metal gate electrode that is separated from the semiconductor surface by an intervening insulator (typically silicon

dioxide). The control signal required is essentially a bias voltage with no significant steady-state gate current flow in either the on-state or the off-state. Even during the switching of the devices between these states, the gate current is small at typical operating frequencies (<100 kHz) because it only serves to charge and discharge the input gate capacitance. The high input impedance is a primary feature of the power MOSFET that greatly simplifies the gate drive circuitry and reduces the cost of the power electronics (Baliga, 1995b).

Power MOSFETs can also be paralleled easily because the forward voltage drop increases with increasing temperature, ensuring an even distribution of current among all components. However, at high breakdown voltages (>200 V) the on-state voltage drop of the power MOSFET become higher than that of a similar size bipolar device with similar voltage rating. This makes it more attractive to use the bipolar power transistor at the expense of the worst high frequency performance.

The power MOSFET is a unipolar device. Current conduction occurs through transport of majority carriers in the drift region without the presence of minority carrier injection required for bipolar transistor operation. No delays are observed as a result of storage or recombination of minority carriers in power MOSFETs during turn-off. Their inherent switching speed is particularly attractive in circuits operating at high frequencies where switching power losses are dominant.

Power MOSFETs have also been found to display an excellent safe operating area; that is, they can withstand the simultaneous application of high current and voltage (for a short duration) without undergoing destructive failure due to second breakdown. These characteristics of power MOSFETs make them important candidates for many applications. They are being used in audio/radiofrequency circuits and in high-frequency inverters such as those used in switch mode power supplies. Other applications for these devices are for lamp ballasts and motor control circuits.

Power MOSFETs contains a high impedance input which greatly simplifies drive circuitry and allows a bias voltage with very low current (on the order of 100nA) to act as the gate control signal. The gate control signal is applied to a metal gate electrode, which is separated from the semiconductor surface by an intervening insulator such as silicon oxide. Due to the low gate current and the device switching operation, the inherent switching speed orders of magnitude greater than the older conventional bipolar transistors, making it an attractive device for circuits operating at high frequencies (where switching power losses dominate)

These advantages of the power MOSFET devices, however, are somewhat diminished by their conduction characteristics which are highly dependent upon temperature and voltage rating. A solution to this problem was the advent of the insulated gate bipolar transistor or IGBT.

The double-diffusion MOSFET, or DMOSFET has been the most commercially successful structure of the MOSFET family; it is more stable than the older V-groove MOSFET, and the new U-groove MOSFET has just begun to become commercialized. The DMOS structure is constructed through the use of planar diffusion technology utilizing refractory gates (such as using polysilicon as a mask). The P- region and the N⁺ source regions are diffused through a

common window defined by the polysilicon mask, and the difference in the lateral diffusion between the P- and N+ source defines the surface channel.

2.2 POWER MOSFETS

Power MOSFET are the most commercially advanced devices. The operation of the power MOSFET relies upon the formation of a conductive layer at the surface of the semiconductor. Until recently, the development of discrete devices has followed the basic concept of the lateral channel structure used in the earlier applications. Such devices have the drain, gate, and source terminals on the same surface of the silicon wafer. Although this feature makes them well suited for integration, it is not optimum for achieving a high power rating. The vertical channel structure, with source and drain on opposite surface of the wafer, is more suitable for a power device because more area is available for the source region and because the electric field crowding at the gate is reduced (Baliga, 1995c).

Different Types of Power MOSFETs

For power applications many individual devices may be connected together in parallel during the final metallization process. However, there are two major reasons why the planar structure of Figure 2.1 is unsatisfactory if it is simply scaled up for higher powers.

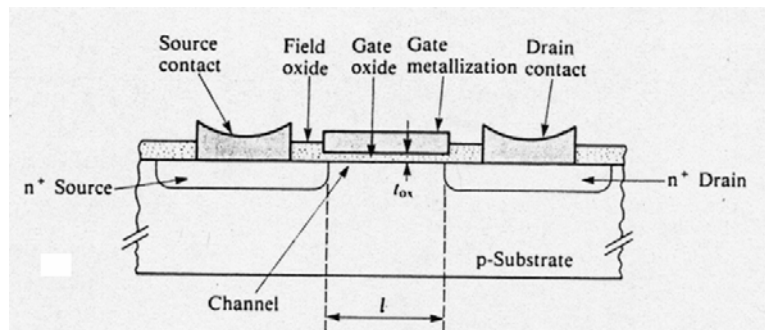


Figure 2.1: An n-channel MOSFET

First, the drain-source spacing has to be increased in order to obtain a high voltage blocking capability. If l were kept small, the drain depletion layer would eventually punch through to the source at high values of V_{DS} . By having the channel more heavily doped than the adjacent drain region, the depletion layer occupies the drain region preferentially. The second disadvantage of the lateral power MOS transistor arises from the need to make all three connections (to source, gate, and drain) on the same, upper surface. While this facilitates the monolithic integration of component, it complicates the metallization required for a single power device. Both effects reduce the area of silicon usefully used to form the active transistor region. There is thus a low silicon utilization factor. As a result, lateral power MOSFETS are rarely used as discrete devices, except in some linear applications. They are being used increasingly in power integrated circuits.

During 1970s some radically different MOSFET configurations were evolved. These eventually enabled the two major disadvantages of the lateral power MOS transistors to be avoided. The essential step was to use the substrate material to form the drain contact. As a result the current flows “vertically” through the silicon from drain to source. A technique that was more successful and that led to the production of the first commercial power devices was to use an anisotropic

etch to produce a V-shaped groove in the silicon surface. In order to make a power FET, the V-shaped groove is etched into the surface after successive p and n⁺ diffusions have been completed. This produces the structure shown in Figure 2.2. Such devices have become known as VMOS transistors.

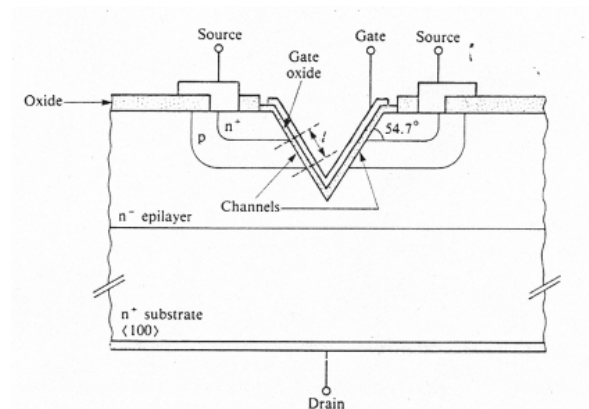


Figure 2.2: An Vertical V-groove MOSFET

The angle of the groove is determined by the crystal structure of silicon. On the heavily doped n⁺ substrate a lightly doped n⁻ epitaxial layer is grown, and into this, successive p and n⁺ layers are diffused, just as though an npn bipolar junction transistor were being made. The channel length l is determined by the relative depths of the successive diffusions. A penalty that has to be paid for the VVMOS structure is the reduced electron mobility in the channels under the {111} faces of the V-grooves, in comparison with that of normal {100} MOS devices. Each gate of the VVMOS transistor controls the current from the two sources, one on either side of the groove, as it flows to the common drain. Current crowding at the apex of the groove can limit the useful current rating of the device. In the OFF state the sharp apex causes a local region of high electric field to develop, and this may limit the voltage rating. By arranging for the groove etch to stop before an apex is formed, the structure shown in Figure 2.3 can be obtained and both of these problems are reduced.

Because of problems in controlling the critical etching processes, VVMOS FETS are difficult to produce. To a large extent they have been superseded by a different type of vertical MOS transistor. The different lateral *spreads* of the two diffusions can be used in exactly the same way. This technique was applied first to the lateral MOS device. The channel length is no longer dependent on the resolution of the photomicrography, but on the control of the lateral spreads of successive phosphorus and boron diffusions through the same oxide window.

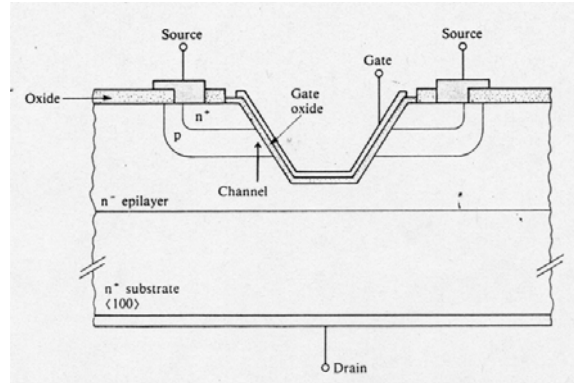


Figure 2.3: An truncated V-groove vertical MOSFET

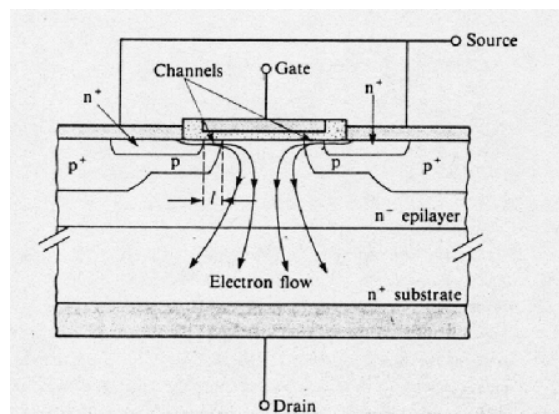


Figure 2.4: An vertical DMOS structure

The vertical, double-diffused MOS structure which has become so important and which we shall refer to as VDMOS, fuses together these two concepts. As Figure 2.4 shows, it uses the double-diffusion technique to determine the channel length l , and it supports the drain voltage vertically in the n^- -epilayer. The current flows laterally from the source through the channel, parallel to the silicon surface, and then turns through a right angle to flow vertically down through the drain epilayer to the substrate and the drain contact. The p-type “body” region, in which the channel is formed when a sufficiently positive gate voltage is applied, and the n^+ source contact regions are diffused successively through the same window etched in the oxide layer. The channel length can be controlled to submicrometer dimensions if required. Because of the relative doping concentrations in the diffused p-channel region and the n^- epilayer, the depletion layer which supports V_{DS} extends down into the epilayer rather than laterally into the channel.

The third power MOSFET structure named the U-MOSFET structure is shown in Figure 2.5. A U-shaped groove is formed in the gate region by using reactive ion etching. The fabrication of this structure can be performed by following the same sequence as described earlier for the VMOSFET with the V-groove replaced by the U-groove. The U-groove structure has a higher channel density than either the VMOS or DMOS structures which allows significant reduction in the on-resistance of the device.

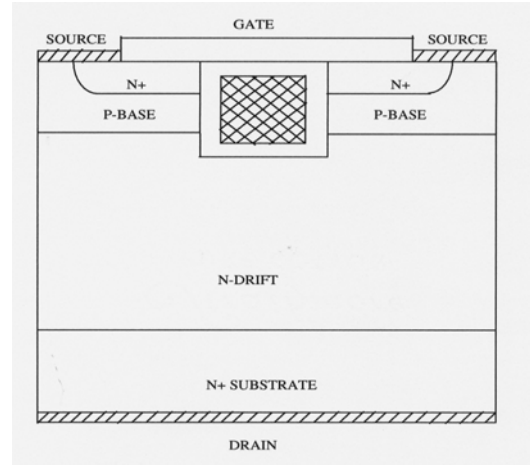


Figure 2.5: The UDMOSFET Structure

Merits of DMOSFET Structure

DMOS has number of significant features. These include the vertical geometry, the double diffusion process, the polycrystalline silicon gate, and the cellular structure. Following these, several important but more subtle technological refinements have been introduced in different ways by different manufacturers and so have steadily improved device performance.

The range of voltage, current, and switching frequency over which the different types of semiconductor device can operate are illustrated in Figure 2.6. The parameter in which the DMOSFET has a marked superiority over other device types is the important one of frequency or switching speed. Not only does it maintain gain to much higher frequency, it also has a linear transfer characteristic. In common with other types of MOSFET, its input resistance is very high. Negligible power is needed to maintain it in the ON state, and it can often be driven directly from a CMOS or TTL logic output. Like other MOSFETs, it is liable to catastrophic failure through electrostatically induced gate overvoltage, but the thicker gate oxide and greater gate capacitance of power devices means that they are less vulnerable than integrated circuits.

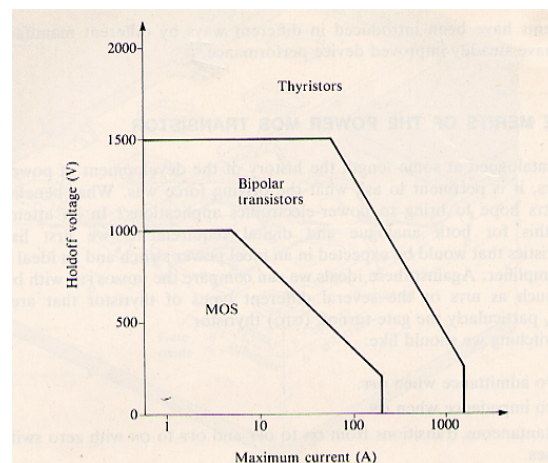


Figure 2.6: Voltage Current Limitations of MOSFET, BJT and Thyristor

Vertical MOSFETs, while they are much improved in their voltage and current capabilities over lateral MOS devices, do not yet match the combined high voltage and high current ratings

possible with bipolar devices. One reason for this is the steep rise in the ON-state drain-source resistance, $R_{DS(on)}$, with the FET voltage rating. The minimization of the ON-state voltage drops is an important consideration in power devices. For voltage ratings less than about 100 V, VDMOS FETs and bipolar devices occupying the same silicon area have similar ON-state voltage drops. This is not true for high voltage ratings, particularly above about 200 V.

The FET ON resistance has the very desirable feature that it increases with increasing temperature. This aids the establishment of a uniform current density throughout the device and means that FETs operated in parallel share the current. Current hogging is avoided. Equally important is the fact that the second breakdown, which so limits the power handling capacity of BJTs, is normally avoided (Grant, 1989b).

Operation Of Power MOSFET

P-N junction between the P-base region and the N-drift region of all three power MOSFET device structures provides the forward blocking capability. P-base region is connected to the source metal by a break in the N^+ source diffusion that establishes a fixed potential to the P-base region during device operation. When the gate electrode is externally shorted to the source, the surface of the P-base region under the gate, which is the channel region remains unmodulated at a carrier concentration determined by the doping level. Now, application of a positive gate voltage reverse biases the P-base/N-drift region junction. This junction supports the drain voltage by the extension of a depletion layer on both sides. The depletion layer extends primarily into the N-drift region due to the higher doping level of the P-base region. A lower drift region doping and a larger width are required for getting higher drain blocking voltage capability. Gate electrode is always connected to the source to establish its potential at the lowest point during the forward blocking state. Otherwise its potential can rise via capacitive coupling to the drain potential. This induces modulation of the channel region, which can produce an undesirable current flow at drain voltage well below the avalanche breakdown limit. Power MOSFETs cannot support large drain voltages unless the gate is grounded. Application of a positive bias to the gate electrode creates a conductive path extending between the N^+ source region and N-drift region. The gate bias creates strong electric field normal to the semiconductor surface through the oxide layer and modulates the conductivity of the channel region. The gate induced electric field attracts electrons to the surface of the P-base region under the gate. This electric field strength is sufficient to create a surface electron concentration that overcomes the P-base doping. The resulting surface electron layer in the channel provides a conductive path between the N^+ source regions and the drift region. Now if a positive drain voltage is applied, current will flow between drain and source via N^- drift region and the channel. This current flow is controlled by the resistance of these regions.

Power MOSFETs can be switched off by reducing the gate bias voltage to zero, that is by externally shorting the gate electrode to the source electrode. When the gate voltage is removed, the electrons are no longer attracted to the channel and the breaks the conductive path from drain to source. This switching from on-state to off-state takes place rapidly without any delay caused from minority carrier storage and recombination which are experienced in bipolar devices. This turn-off time is controlled by the rate of removal of the charge on the gate electrode because this charge determines the conductivity of the channel.

There are parasitic $N^+ - P - N^+$ bipolar transistor in all three power MOSFET structures. This parasitic bipolar transistor is kept inactive by shorting the P-base region to the N^+ source regions by the source metallization. The resistance between the P-base region between the shorts can become large and any lateral current flow in the P-base, due to the capacitive currents arising at high applied $[dV/dt]$ to the drain can lead to forward biasing the N^+ / P junction at locations remote from the shorts. Forward biasing of the N^+ / P junction activates the parasitic bipolar transistor and leads to the initiation of minority carrier transport. This not only can slow down the switching of the power MOSFET but also can lead to second breakdown. Due to the high $[dV/dt]$'s observed in the high frequency applications, it is common practice to form the short in every cell and minimize the length of the N^+ source region from the edge of the channel to the short.

Power MOSFET Doping Profile

Diffusion profile for a typical power MOSFET is shown in Figure 2.7. The solid lines indicate the dopant distribution and the dashed lines are the net carrier concentration profiles. This net carrier concentration profile differs from the dopant profiles due to compensation effects. N_D is the constant doping concentration in the N^- drift region. N_{SP} and N_{SN^+} are the surface doping concentration of p-base region and the N^+ -source region. The peak doping in the P-base region is determined by the surface concentration of the P-base diffusion and the N^+ source depth. This peak doping in the P-base region is an important parameter for the doping profile because it controls the threshold voltage of the power MOSFET.

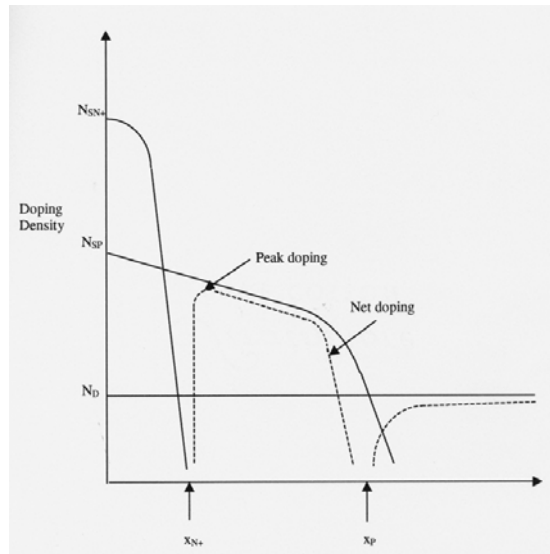


Figure 2.7: Typical doping profile of a power MOSFET

Another important design parameter of power MOSFET is the channel length. MOSFET on-resistance and the transconductance is highly influenced by channel length. The channel length is determined by the difference in the depths of the P-base and N^+ source diffusions, i.e. $(X_p - X_{N^+})$.

A parasitic $N^+ - P - N$ bipolar transistor exists in the power MOSFET despite the shorting of the N^+ source to the P-base. When the P-base/ N -drift junction is reversed biased, the depletion layer in the P-base can extend to the N^+ source/ P -base junction and cause premature reach-through

breakdown. It is important to design the P-base diffusion profile so that sufficient charge is resident in the P-base to prevent reach-through of the depletion layer to the N⁺ source.

Threshold Voltage

The voltage on the gate electrode at which strong inversion begins to occur in the MOS structure is an important design parameter for power MOSFETs because it determines the minimum gate bias required to induce a conductive channel. This voltage is called the threshold voltage. For proper device operation, the threshold voltage should not be very large or very small. If it is large, a high gate bias voltage will be needed to turn on the power MOSFET. This might cause problems with the design of the gate drive circuitry. It is also very important that the threshold voltage not be too low. Due to the existence of charge in the gate oxide, it is possible for the threshold voltage to be negative for n-channel power MOSFETs.

This is unacceptable condition because a conductive channel will not exist at zero gate bias voltage, i.e., the device will exhibit normally-on characteristics. Even if the threshold voltage is above zero for an n-channel power MOSFET, its value should not be too low because the device can then be inadvertently triggered into conduction either by noise signals at the gate terminal or by the gate voltage being pulled up during high speed switching. Typical power MOSFET threshold voltage are designed to range between 2 and 3 volts. Threshold voltage can be expressed by (Streetman, 1995):

$$V_T = \phi_{ms} + 2\phi_f - \frac{(Q_i + Q_d)}{C_{ox}} \quad (2.1)$$

Here ϕ_{ms} is the work function difference between metal and the semiconductor, ϕ_f is the surface potential, Q_i and Q_d are the oxide interface charge and charge in the depletion region per unit area respectively. C_{ox} is the oxide capacitance per unit area. Surface Potential ϕ_f is temperature dependent and can be expressed as:

$$\phi_f = \frac{kT}{q} \ln\left(\frac{N_A}{n_i}\right) \quad (2.2)$$

On-State Conduction

A conductive path is created across the P-base region underneath the gate by applying a positive voltage at the gate electrode for a n-channel device. The current flow is limited by the total resistance between the source and drain. This resistance consists of many components, which determines the on-state voltage drop when the device is carrying current. Figure 2.8 presents the same structure of DMOSFET with different components of the resistance. The resistance of the N⁺ source (R_{N+}) and substrate (R_{SUB}) regions are negligible for high voltage power MOSFETs that have high drift region resistance. They became quite significant if the drift and the channel resistance became small as in the case of low (<100 volts) breakdown voltage devices. The channel resistance (R_{CH}) and accumulation layer resistance (R_A) are determined by the conductivity of the thin surface layer induced by the gate bias. These resistances are functions of the charge in the surface layer and the electron mobility near the surface.

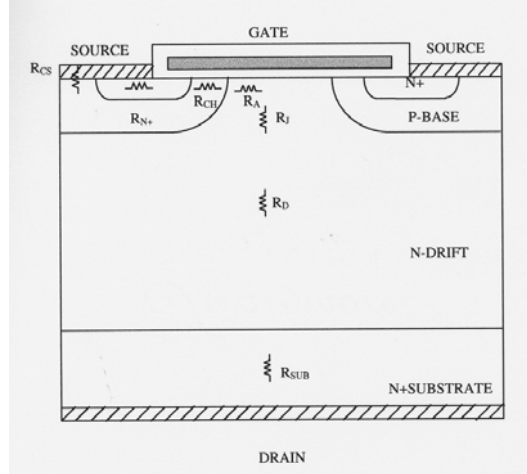


Figure 2.8: The resistance components within DMOSFET structure

The drift layer contributes two components to the total on-resistance. The portion of the drift region that comes to the upper surface between the cells contributes a resistance R_J that is enhanced at higher drain voltages due to the pinch-off action of depletion layers extending from adjacent P-base regions. This phenomenon has been termed the JFET action. Finally the main body of the drift region contributes a large series resistance (R_D) especially for high voltage devices. The on-resistance of power MOSFET is the total resistance between the source and drain terminals in the on-state. The on-resistance determines the maximum current rating of the device. The power dissipation in the power MOSFET during current conduction is given by:

$$P_D = I_D V_D = I_D^2 R_{on} \quad (2.3)$$

Expressed in terms of chip area (A):

$$\frac{P_D}{A} = J_D^2 R_{on,sp} \quad (2.4)$$

Where (P_D/A) is the power dissipation per unit area; J_D is the on-state current density and $R_{on,sp}$ is the specific on-resistance, defined as the on-resistance per unit area. These expressions are based upon the assumption that the power MOSFET is operated in the linear region at a relatively small drain bias during current conduction. The specific on-resistance of the power MOSFET is determined by the all resistance components.

$$R_{on} = R_N + R_{CH} + R_A + R_J + R_D + R_{SUB} \quad (2.5)$$

Additional resistances can arise from a non-ideal contact between the source/drain metal and the N^+ semiconductor regions as well as from the leads used to connect the device to the packages. Now each part of the specific on-resistance is discussed very briefly.

Substrate Resistance

For high voltage power MOSFET, this resistance is negligible but it contributes significantly for the power MOSFETs with low breakdown voltages. For rapid current spreading at the drift interface, the current density within the substrate could be assumed uniform. The specific resistance contributed by the substrate is given by:

$$R_{SUB} = \rho_{SUB} t_{SUB} \quad (2.6)$$

Where ρ_{SUB} is the resistivity of the substrate and t_{SUB} is the thickness.

Source Resistance

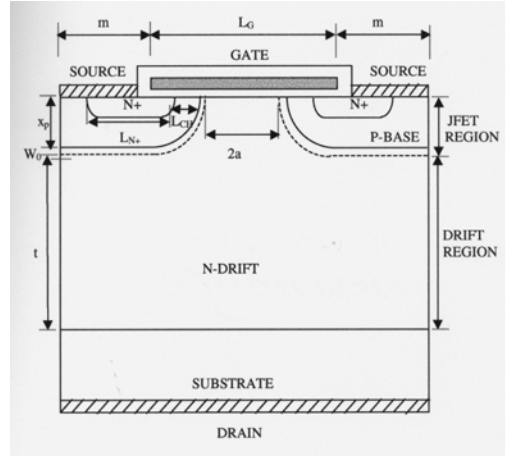


Figure 2.9: Cross section of DMOSFET for the specific on resistance analysis.

Figure 2.9 shows the cross section of the DMOSFET cell. In this cell $2m$ is the cell diffusion window and L_G is the length of the gate electrode between the adjacent cells. L_{N+} is the length of the N^+ source region. The specific resistance contributed by the source region is given by:

$$R_{N+,sp} = \frac{1}{2} \rho_{SN} + L_{N+} (L_G + 2m) \quad (2.7)$$

Where ρ_{SN+} is the sheet resistance of the N^+ diffusion and $(L_G + 2m)$ is the cell repeat spacing. The specific resistance of the N^+ source region is negligible compared with all other resistance in the structure.

Channel Resistance

The specific resistance contributed by the channel is given by:

$$R_{CH,sp} = \frac{L_{CH} (L_G + 2m)}{2\mu_{ns} C_{ox} (V_G - V_T)} \quad (2.8)$$

The channel resistance decreases when the cell repeat spacing is reduced. Decreasing the gate oxide thickness while maintaining the same gate drive voltage can also reduce it.

Accumulation Layer Resistance

Accumulation layer resistance depends on the charge in the accumulation layer and the mobility for free carriers at the accumulated surface. The specific resistance contributed by the channel is given by:

$$R_{A,sp} = \frac{K(L_G - 2x_p)(L_G + 2m)}{2\mu_{nA}C_{ox}(V_G - V_T)} \quad (2.9)$$

Here x_p is the diffusion depth of the p-base region. The factor K is introduced to account for the two dimensional nature of the current flow from the channel into the JFET region via the accumulation layer. Good agreement with experimental results has been observed for $K = 0.6$, which implies that the effective resistance to the drain current flow is 60 percent of the total accumulation region resistance.

JFET Region Resistance

The resistance of the drift region between the P-base diffusions is referred to as the JFET resistance because the current flow resembles that in a junction field effect transistor with the P-base regions acting as the gate regions. Under the assumption that the current is flowing uniformly down from the accumulation layer into the JFET region, its contribution to the specific resistance is then given by:

$$R_{J,sp} = \frac{\rho_D(L_G + 2m)(x_p + W_o)}{(L_G - 2x_p - 2W_o)} \quad (2.10)$$

Drift Region Resistance

Many models for the current spreading into the drift region has been proposed. One such model that allows a reasonably accurate estimation of the drift region spreading resistance, is based on the current spreading from a cross section of a ($a = L_G - 2x_p$) at a 45 degree angle. The cross section for the current flow then increases with depth through the drift region. According to this assumption there is no overlap in the current flow path. The specific resistance contribution for the drift region is obtained as:

$$R_{D,sp} = \frac{\rho_D(L_G + 2m)}{2} \ln\left(\frac{a+t}{a}\right) \quad (2.11)$$

But, in the case of higher breakdown voltage devices with larger drift region thickness, the current flow paths will overlap. The drift region resistance must then be modeled as the sum of a region where the cross section increases with depth and a second region with uniform cross section equal to the cell width. This lead to a drift region specific resistance is give by:

$$R_{D,sp} = \frac{\rho_D(L_G + 2m)}{2} \ln\left(\frac{L_G + 2m}{L_G - 2x_p - 2W_o}\right) + \rho_D(t - m - x_p - W_o) \quad (2.12)$$

Contact Resistance

The drain contact covers the entire back surface of the device. The drain contact resistance contribution to the specific resistance is given by:

$$R_{CD,sp} = \rho_C \quad (2.13)$$

Where ρ_C is the specific contact resistivity.

For source region, the contact region depends upon the area where the source metal and the source N^+ region overlap. If the resulting contact area is A_{CS} , the contribution to the specific resistance by the source contact is given by:

$$R_{CS,sp} = \frac{A_{cell}}{A_{CS}} \rho_C \quad (2.14)$$

Ideal Specific On-Resistance

In the ideal case where the resistances of the N^+ source, N^+ substrate, n-channel region, accumulation region and the JFET region are negligible, the specific on-resistance of the power MOSFET will then be determined by the drift region alone. In addition if it is assumed that the current flows uniformly through the drift region without current spreading effects, the resistance of the drift region is referred to as the ideal specific on-resistance for the power MOSFET. For n-channel devices ideal specific on-resistance is:

$$R_{on,sp} = \frac{W_D}{q\mu_n N_D} \quad (2.15)$$

If the doping level N_D (cm^{-3}) is required to support a given breakdown voltage V_B and depletion width W (cm) at the breakdown can be calculated as follows (Shams, 1998):

$$N_D = \frac{\epsilon E_C^2}{2qV_B} \quad (2.16)$$

$$W = \frac{2V_B}{E_C} \quad (2.17)$$

The specific on-resistance (ohm-cm^2) associated with the drift layer to support V_B is

$$R_{on,sp} = \frac{W}{qN_D \mu_n} \quad (2.18)$$

$$R_{on,sp} = \frac{4V_B^2}{\epsilon E_C^3 \mu_n} \quad (2.19)$$

Where ϵ is the permittivity, E_C is the breakdown field, q is the electronic charge and μ_n is the electron mobility. Both the electron mobility μ_n and the breakdown electric field E_C are dependent on N_B .

2.3 PROCESS SIMULATOR

The major focus of this thesis is the simulation and fabrication of the Power MOSFET (DMOS). The DMOS was simulated using the Technical Computer Aided Design software developed by Integrated Systems Engineering AG (ISE-TCAD). This is a popular semiconductor process, device and circuit simulator used by many major companies in the semiconductor industry. The latest version of ISE-TCAD, version 6.0, incorporates a graphical interface as well as having ability to edit using a text editor. Both process and device simulations are sequenced in the operational window by utilizing drag and drop icons. GENESIS is ISE's graphical front-end to design, organize and automatically run complete TCAD simulations projects. It is aimed at providing the user with a convenient framework to drive the large variety of ISE simulation and visualization tools and other third party tools, and to automate the execution of fully parameterized projects.

The ISE-TCAD software is a combination of many individual programs, which the user defines in a particular order. Because ISE-TCAD has so many applications and uses, this discussion will only consider those programs used to design and simulate the operation of the DMOS.

The tools in ISE –TCAD are given below along with their functions:

TESIM: - 1D process simulator

DIOS: - 2D and 3D process simulator

MDRAW: -Grid editor

DESSIS: -Device simulator

Process Simulator (DIOS) Introduction

Here DIOS is used for the process simulations. The program DIOS takes as its input a sequence of commands which may be entered from standard input (i.e. at the prompt in a command window) or composed in a command file. An optional additional input is a PROLYT mask file containing details of geometry's for the various mask levels. The simulation of a process flow is thus achieved by issuing a sequence of commands corresponding to the individual process steps. In addition a number of control commands are provided to allow the user to select physical models and parameters, gridding strategies and graphical output properties if desired.

In DIOS, input information can be given in two ways:

Interactive Command Input

DIOS is interactive with user- it is possible to simulate a whole process flow by entering commands line-by-line as standard input and observe the results step-by-step in the graphical output.

Command File input

A command file may be written entirely by the user or generated using the ISE user interface tool LIGAMENT. To save time and reduce syntax errors, it is recommended to either use LIGAMENT to create a template or to copy and edit an example command file from GENESIS database.

Interactive Graphics

DIOS version 6.0 provides an enhanced graphical user interface. If the graphical output is turned on, it presents the user with a console, allowing direct control of many DIOS plot properties such

as selection of species, display of mesh, layers and contour mapping. Other features include plotting of cutlines, point sampling of all available variables and control of program execution. The interactive command input has been used for the process simulation of DMOS.

Process simulation

In the interactive mode of DIOS commands are entered one by one at the dios prompt and the results are seen in the graphic window. First a rough grid was chosen for simulation in proportion to the actual device size. The actual device size was reduced by 30 times to fit into the chosen grid. Only half part of the symmetrical MOSFET is simulated to reduce the simulation time. All the process steps were carried out in the same sequence as those while fabricating the device in the lab. A description of the substrate was given i.e the concentration, orientation and type. Then a wet oxidation was carried out using the diffusion command. Here the required thickness and temperature were specified. For all the photolithography steps, mask command was used in which the material and the left and right locations were specified. The oxide and resist were removed by using etch command. Then a dry oxidation is carried out in which the required thickness of 100nm is specified. Boron implantation is done by implant command in which all the necessary data for implantation is given such as dose, energy and tilt angle. Then boron diffusion is carried out for 58 minutes. This diffusion time is found out by calculations so that the junction depth is 1.5μ deep. Then windows for n^+ region are opened using mask command and then predeposition of phosphorus is done. This is followed by gate oxidation for 7 minutes is done to get an oxide thickness of $350\text{\AA}$. Then contact windows are opened in the same manner as described before.

Now to get the full structure of the simulated device, the device is reflected around y axis. Here the fourth mask level which is used for metallization is substituted by mask command in which the coordinates of the metal contacts are specified. The simulated device is saved for further device simulation. The results from process simulation can be used in MDRAW and DESSIS to get the characteristics of the fabricated device. Various concentrations of all the species can be observed in the graphics window.

Description of Commands Used

Coordinate System

In 2D DIOS uses a right-handed Cartesian coordinate system. In the 2D DIOS simulation plane and in the X11 graphics the DIOS-X axis points laterally to the right and the DIOS-Y axis points vertically to the top. The lanes are in microns. In DIOS the 1D or 2D geometry or layer structure is stored as a system of polygons.

Title

Title ("dmos")

Title must be the first command in each DIOS input file. Note that command words may be abbreviated and DIOS is not case sensitive. During the input parsing, DIOS will check for complete parameter names first and then check for abbreviations. In case of ambiguous input DIOS will print out warnings and inform about the selected parameter name.

Grid

Grid (x (0.0,12.0) , y(-4.0,0.0) , nx=27)

The simulation grid in DIOS is “constructed” in several steps:

A first coarse triangulation is constructed from scratch, covering the entire simulation domain (but not resolving any gradients or material interfaces). This is called initial or User – grid.

A refinement triangle tree is built by subsequent splitting of the triangles. This tree is called ITRI-grid.

Extraction of the leaf elements of the refinement tree.

The grid is adapted to the layer structure.

Extraction of mesh points on the left side of the grid and construction of a special ID mesh (only for 1D layer structure, 1D data profiles and only if Control (1D=on) (default).

Post-processing, mesh optimization (surface-parallel refinement, mesh smoothing, Delaunization)

Re-interpolation of all data sets.

A rectangular domain defined by the x-y coordinates is tessellated using nearly equilateral triangles. Initial triangle spacing is defined by either dx or nx. In the above case of nx=27, the user-defined starting mesh has the domain divided into 27 triangles with bases of 0.44 μ aligned along the x-axis. Alternatively, the smallest possible triangle size may be defined explicitly using dx, which then determines nx indirectly. If the user wants to precisely control the interpolation error, it is recommended to define the initial grid large enough to contain the entire simulation domain during the entire process simulation.

Substrate

Substrate (orientation =100, element =P, conc=2E14, ysubs=0.0)

The substrate command is used to initialize the layer system and to define the properties of the wafer material. The crystal orientation of the wafer surface, the background-doping element and its (constant) concentration are defined. The location and extension of the 1D or 2D layer system in the X-and Y-direction can be defined also directly by XLeft, XRight, YBottom, YTop. The initial (vertical) position of the substrate surface YSubs can be prescribed as well. By default <100> material is assumed. The crystal orientation can be specified only for the substrate. Other (e.g. deposited) silicon layers are treated with the same crystal orientation. The background doping element (ELEMent=B) and the resistivity (RHO) (in Ω cm) can be specified. If RHO=undefined, the background doping concentration can be specified explicitly. If concentration is also undefined (or if a non-positive value is specified) no background is assumed for the simulation.

Etching

Etch (material = OX, stop = sigas, rate (Isotropic = 60), over=20%)

In DIOS no rigorous physical/chemical simulation of etching is done. Instead a set of geometry operations is provided which allows to define local “etching rates” that can be used to approximate the modifications of the structure during the etching process. In a rigorous sense DIOS can not predict the shapes after etching, they are entirely determined by the user supplied etching parameters and by the discretization of the layer structure. The Etching command allows to remove material which is in contact to gas. If only a material (but no etching rates) are specified, all regions (of this material) in contact with the gas region are removed. All inclusions in these regions are deleted as well.

Etching (Material=OX)

If etching rates are specified, by the evolution of the etching front is simulated in a sequence of time steps. An etch stop is defined by a list of materials or boundary sorts, Stop. The isotropic and nonisotropic components of the etching rates can be specified as:

Rate (Isotropic =... A0 = ...A1 =.....A2 =)

Several materials can be etched at the time. In this case one can define the etching time or etch stops.

Mask

Mask (material = resist, thickness=0.8 μ , Xleft=0.0, Xright=4.25)

Lithography processes are not simulated in DIOS. Instead mask regions composed of photoresist (or other material) can be defined with the Mask command. The lateral begin and end position of a mask region can be defined by the user. Several masks can be specified in one command. The mask positions can also be read from an external input file. If the user wants to extract the mask positions directly from the layout, ISE's Ligament tool should be used to specify the process flow, the mask file and the position of the cutline of the DIOS simulation. In the Mask, command the Material, Thickness and the lateral position of the mask corners Xleft, Xright and/or X = (...) can be specified. If nonsymmetric masks are required, separate values DXRight and DThRight can be specified. The mask command is thought for relatively planar surfaces or infinitely high masks. The position of mask edges can alternatively be read from a file. The name of the mask file can be specified as

dios – mask = “file” command

on the command line, when starting DIOS or in the mask command.

Implantation

Implant (element = B, dose=5.66E12, energy=30kev, tilt = 7)

The distribution of the implanted ions and implantation damage can be computed using analytical distribution functions or Monte Carlo simulation. The implantation dose is defined in DIOS as particles per area of the wafer surface. The specified values or Rotation and Tilt define the (3D) incident ion beam with respect to the wafer. Tilt is assumed around the X-axis and Rotation is assumed around the Z-axis of the (tilted) wafer coordinate system. The orientation of the wafer and the two angles Tilt and Rotations are sufficient to describe completely the physical system during the implantation. Implant damage in the form of point defects and amorphization is simulated by default.

Point Defects

During the implantation step, point defects (interstitials, vacancies) and extended defects (dislocations) are created. In DIOS, the “effective” number of interstitials and vacancies created per ion are controlled by Ifactor and Vfactor respectively. In the default model (damage=+1), the point defect distributions thus follow the as-implanted profile. The alternative model is the Hobler function (damage=Hobler). The parameter damage can also be set to Monte Carlo damage (damage=Mcdamage) when a MC implantation is performed.

Amorphization

By default, the amorphization of the silicon is computed based on the Holber function (amorphization-Hobler). Wherever the damage induced by implantation exceeds a threshold value. (1.15e22 cm⁻³), the silicon is assumed to be amorphous. The alternative amorphization

model is based on the as-implanted profile (amorphization=+1). Dose conservation in layered structures is achieved by converting the layer thickness according to the ratio of the projected ranges and then rescaling the profiles accordingly to ensure dose conservation.

Diffusion

Diffusion (element = B, temperature = 1100degc, time=58min, atmosphere=O2)

During all high temperature processes the dopant redistribution needs to be computed. The redistribution is caused by dopant, point defect diffusion, chemical reactions at the interfaces and inside the layers, convective dopant transport due to internal electrical fields as well as due to material flow and moving material interfaces. In DIOS the Diffusion command is used to simulate all high temperature steps. The various process atmospheres are described by the parameter Atmosphere-O2 | HC1 | H2O | H2O2 | N2 | Epitaxy | Prebake | Mixture. The diffusion model can be selected with a global model switch.

DIFFusion (ModDiff = Conventional | Equilibrium | LooselyCoupled | SemiCoupled | PairDiffusion).

Diffusion Models in DIOS

Diffusion models may be selected either globally,

DIFFusion: (ModDiff=PairDiffusion)

Or for an individual diffusion step:

DIFFusion (time=10s, temperature=1050, ModDiff=PairDiffusion)

The key point of the point defect assisted diffusion models in crystalline silicon is the coupling between dopants and point defects (interstitial and/or vacancies). This coupling is described by three types of equations:

transient clustering model for each dopant,

pair formation, ionization and diffusion of pairs

ionization and diffusion of unpaired point defects

transient (or equilibrium) (de-)clustering of dopants

Equilibrium (or transient) (de-) clustering of point defects (storage of majority of point defects generated as implantation damage in immobile <311> clusters with subsequent transient dissolution).

In the following discussion:

D is the diffusion coefficient, it usually depends on the carrier concentrations.

CS is the concentration of substitutional doping, e.g. BActive for boron.

CC is the carrier concentration (electron (n) for donors and holes (p) for acceptors)

I0 is the concentration of neutral unpaired point defects. Identical driving forces and equations are assumed for interstitials and vacancies.

PairDiffusion

The PairDiffusion model is the most complete model in terms of the coupling effects between point defects and dopants. The main assumptions for the diffusion are the following:

point defect-dopant pairs: mobile species, numerous charge states

unpaired point defects: mobile species, numerous charge states

unpaired dopant on lattice site (substitutional): immobile species, fully activated

The diffusion flux is given by: $j = -D (CC) * \text{grad} (CS * CC * I_0)$

For example, for boron the diffusion flux is given by: $j = -D (p) * \text{grad} (B_{\text{active}} * p * I_0)$

The PairDiffusion is recommended if transient or nonlocal coupling between dopant species or point defects need to be simulated, e.g. for reverse short channel effect simulation.

SemiCoupled

The main assumption of the SemiCoupled model are the following:

point defect-dopant pairs: mobile species, numerous charge states

unpaired point defects: mobile species, numerous charge states

unpaired dopant on lattice site (substitutional): immobile charge states

These assumptions are the same as the PairDiffusion one. However, the major difference is the cancellation of one driving force. For the PairDiffusion we had:

$$J = -D (CC) * \text{grad} (CS * CC * I_0) = -D (CC) * I_0 * \text{grad} (CS * CC) - D (CC) * CS * CC * \text{grad} (I_0)$$

This second gradient is completely neglected in the SemiCoupled model and we obtain:

$$J = -D (CC) * I_0 * \text{grad} (CS * CC)$$

For example, for boron the flux is given by: $j = -D (p) * I_0 * \text{grad} (B_{\text{active}} * p)$

The main consequence of this definition of the driving force is the absence of diffusion if there is no doping gradient, even if a gradient of point concentration exists. In this case, the PairDiffusion model and the SemiCoupled model give different results. The parameter values are the same as for the PairDiffusion model. It is not recommended to use the SemiCoupled model, due to the rather arbitrary cancellation of a driving force. The model has been implemented mainly to understand the nonlinear coupling.

LooselyCoupled

Assumptions:

point defect-dopant pairs: not existing

unpaired point defects: mobile species, numerous charge states

unpaired dopant on lattice site (substitutional): mobile species, fully activated

These assumptions are fundamentally different from the two previous models. In this model the balance equations for both point defects and dopant species do not contain any more dopant point defect pairs. However, the neutral unpaired point defect concentration remains as factor outside of the gradient in the driving force, similar to the SemiCoupled model. The model assumptions are consistent and motivated empirically: the dopant diffusivity depends on the local point defect concentrations.

The diffusion flux is given by: $j = -D (CC) * I_0 * \text{grad} (CS * CC)$

For example, the boron flux is given by: $j = -D (p) * I_0 * \text{grad} (B_{\text{active}} * p)$

The parameters for the LooselyCoupled model are the same as for the Equilibrium and Conventional models. These parameters are distinct from the ones used in the PairDiffusion and SemiCoupled models. The LooselyCoupled model can be used to account for transient diffusion behavior.

Equilibrium

Assumptions:

neutral point defect concentration: prescribed externally, e.g. constant,

point defect dopant pairs: not existent

unpaired point defects: not balanced. Can be derived from assumed neutral point defects, carried concentrations and substitutional dopant concentrations.

Unpaired dopant on lattice site (substitutional): mobile species, fully activated.

This model can be derived from the LooselyCoupled model if the point defects are assumed to diffuse very fast. For an inert diffusion the constant boundary value of the neutral interstitial concentration extends through the entire simulation domain. During oxidation steps an inhomogeneous interstitial profile is computed, which depends on the local oxidation rate at a “nearby” silicon surface.

The diffusion flux is given by: $j = -D(CC) \cdot I_0 \cdot \text{grad}(CS \cdot CC)$

For example, for boron the diffusion flux is given by: $j = -D(p) \cdot I_0 \cdot \text{grad}(BActive \cdot p)$

The assumption for the diffusivity corresponds to the assumptions made in Suprem-3. The same parameters as in the LooselyCoupled and Conventional model are used. The equilibrium model can NOT(!) be used to simulate transient or non-equilibrium coupled diffusion effects.

Conventional

Assumptions:

neutral point defect concentration: prescribed externally, e.g. constant,

point defect dopant pairs: not existent

unpaired point defects: not balanced. Can be derived from assumed neutral point defects, carried concentrations and substitutional dopant concentrations.

Unpaired dopant on lattice site (substitutional): mobile species, fully activated.

The basic model assumption are equivalent to the Equilibrium model. During oxidation steps an inhomogeneous interstitial profile is computed, which depends on the local oxidation rate of a “nearby” silicon surface.

The diffusion flux is given by: $j = -D(CC) \cdot I_0 \cdot \text{grad}(CS \cdot CC)$

For example, for boron the driving force is given by: $j = -D(p) \cdot I_0 \cdot \text{grad}(BActive \cdot p)$

There are two major differences between the Conventional and the Equilibrium models.

In the Conventional model the diffusion equation for each single dopant is solved separately and an Gummel-like outer iteration process is used. (In the Equilibrium model the full-coupled Newton-problem is solved.) Some convergence problems may occur for diffusion in polysilicon with the Conventional model.

The second difference is the selection of empirical models for the dopant diffusivities. In the Conventional model for each dopant species the model can be selected individually: Diffusion (B (ModDif=Suprem-2 | Suprem-3 | FairTsai | DC | DEFF)). In the Equilibrium model for all dopants always a Suprem-3 diffusivity model is used; only the coefficients can be modified.

Although the diffusion models are available for any temperature range, the default parameters are valid for temperatures greater than 600°C and are most reliable in the range 750°C - 1250°C. The initial value of the total concentrations of point defects and dopant species are defined or modified externally, e.g. during an ion implantation step. The total concentrations are solution variables in the global system of balance equations (initial –boundary value problem for the system of nonlinear parabolic partial differential equations for all dopant species to be solved on the mesh). The assumed diffusion fluxes depend on the chosen diffusion model.

Given the total concentrations of point defects and dopant species, the concentrations of substitutional dopants, neutral point defects, clusters, electrons etc. (local solution variables) and their derivatives with respect to the total concentrations are solved in a local Newton iteration for each of the mesh points. The local system of equations is composed of nonlinear algebraic and/or ordinary differential equations. The number and type of equations in the local system depends on the material and the chosen models (e.g. clustering). For all models, except PairDiffusion and SemiCoupled by default an equilibrium arsenic-clustering model is assumed. For the PairDiffusion and SemiCoupled model a transient boron (de) clustering model is selected by default.

For ModDiff = Conventional the diffusion model can be specified individually for each material and dopant species. In crystalline silicon the default model has been described above. In polycrystalline materials (Po, MS) a so called two-stream model is used by default, which takes into consideration the effects of grain growth and of grain and grain boundary diffusion. For, other material either a constant diffusivity (DEFF) or a concentration dependent (DC) diffusivity is assumed. For all diffusion models other than Conventional the same type of equations in crystalline silicon is selected and the type of equations for materials other than crystalline silicon is fixed; only the coefficients can be modified. In polycrystalline materials a two stream model and in other materials a concentration dependent diffusivity are used (Mauriello, 1998). Figure 2.10 shows boron (p type) diffusion and Figure 2.10 shows phosphorus diffusion (n type) in the p-well.

Oxidation

Diffusion (temperature = 1100degC, atmosphere = O₂, thickness = 100nm)

A rigorous modeling of oxidation and other thermal processes that change the layer structure (silicidation) includes the chemical reactions and segregation at interfaces (dissolution of particles, reaction of dissolved particles with a layer material, production of a new layer material), the diffusion, convection and (if appropriate) volumetric reactions of dissolved particles, the screening property of some interfaces or layers for

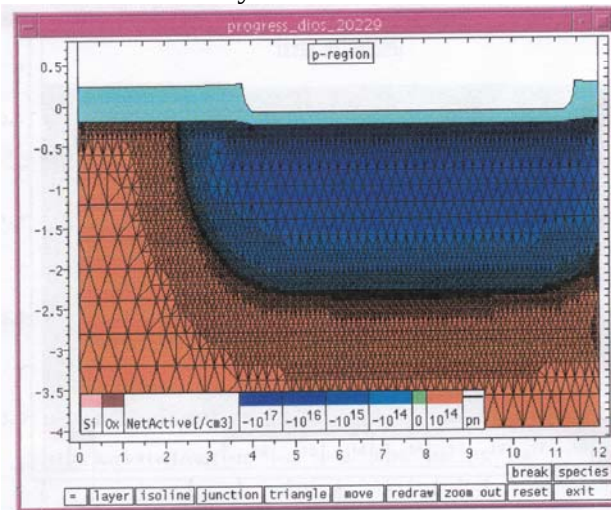


Figure 2.10: Boron (p-region) diffusion in silicon epilayer

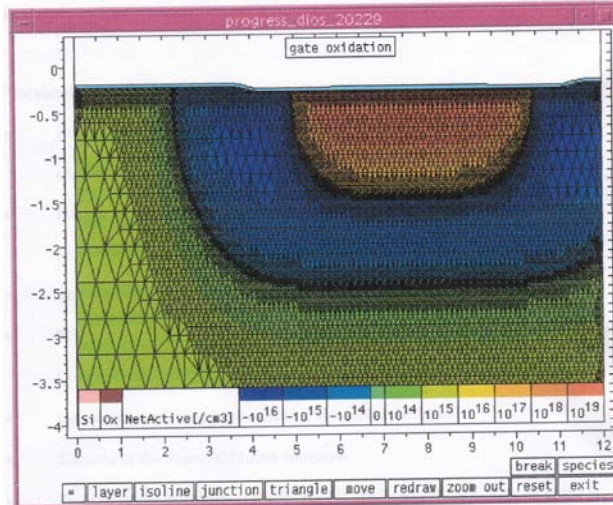


Figure 2.11: Boron (p-region) and Phosphorus (n+ region) diffusion

particle fluxes, and a result of interface or bulk reactions, a mechanical deformation of the entire layer structure.

In process simulation one may assume (quasi-) stationary reaction-diffusion equations for the oxidizing species, and decouple their solution from the solution of the (quasi-) stationary mechanical problem (again, the variations of the velocity with time are negligibly small and mainly defined by the slowly changing structure).

The simulation of an oxidation process is subdivided into several steps:
 solution of the reaction-diffusion-convection equation for the dissolved species (oxidant diffusion and interface reactions as boundary conditions)
 Evaluation of interface reaction terms and computation of boundary conditions for mechanical problem.

Solution of mechanical problem

Transformation of the grid (for NewDiff=0), execution of the convection step for Control (NewDiff=1, Convection≠Regrid) with interpolation of the concentrations

Computation of boundary conditions for dopant diffusion.

Solution of the dopant diffusion equations

Transfer of the discontinuous velocities at the interfaces from the grid to the layer structure (passive deformation)

Application of consumption rates to define final layer structure, topology test and modification (active deformation)

Local update of the grid in the vicinity of the moving interfaces (for NewDiff=1) interpolation of concentrations, performing convection step for Control (Convection=Regrid)

Test and, if necessary, full readaptation of the grid

By default (MODOX=Massoud), the initial oxide thickness is chosen 1.5nm. For MODOX=DealGrove a temperature dependent initial oxide thickness is computed as suggested

in (ISE TCAD Manual). For thin (gate) oxides and high oxidation temperatures this initial oxide thickness might already exceed the desired final oxide thickness.

Several simplified oxidation models are implemented in DIOS. They differ mainly with respect to the complexity and coupling of the physical models involved. The thin oxide model of Massroude (an elaboration of the Deal-Grove equation) is used by default; oxidation rate is Si crystal orientation-dependent.

diff: establishes a new set of global defaults:

dthickness = 2nm – maximal increase in oxide thickness per time step.

Save

Save (file = dmos)

The SAVE command is used to write output files for subsequent evaluation of the results, continuation of the simulation (dmp *) or for off-line coupling to other simulation tools.

SAVE (File=xxx,
TYPE = (dmp, exp, prf, plf, dmp.gz, bound, dp, cmd, geb, MdrawAndLines, dmp.Z, kg,dom,
User, Itri, Picasso, MESHDP, lay, lai, KPIF, meshbuild, Solidis, 3D, Gip, Mdraw, VISE, DFISE))
By default, a binary DIOS save file is written.

Files can be saved regularly at the end of a process step and after a certain number of time steps.

Replace (Control (NSave=100))

File can also be saved regularly after a certain interval of wall clock time (not simulated time, not CPU time):

Replace (Control (Saveeach = 2h))

With the

Save (File = xxx, Type = Mdraw)

Command, DIOS results can be used in the mesh generator MDRAW or in the device simulator DESSIS. The four necessary files are saved:

boundary description: xxx_mdr.bnd

command file for MDRAW: xxx_mdr.cmd

process simulation DF-ISE grid file: xxx_dio.grd [.gz]

process simulation DF-ISE doping file: xxx_dio.dat [.gz]

Contact

Up to 20 contacts can be defined in the data record

Contacts (Contact1 (name =, x=, y=, xe=, ye=, Location=) ...)

The contact names are saved to the geometry description file xxx_mdr.bnd. They can be used subsequently in the device simulation with DESSIS.

For Location = Bottom | WellLeft | WellRight | TopLeft | TopRight, the begin and end points (x, y), (xe, ye) are defined automatically. These contacts are placed on the outer contour of the grid.

Graphic

Graph (triangle = on, plot)

In the interactive mode the command

Graphic (

Calls a local command loop, where graphical output can be done. If the closing parenthesis is entered, the simulator leaves this local command loop. Graphic commands in the command files are executed. By default a 2D plot of a layer system, net doping and p-n junctions (if present) are shown:

Graphic (Plot)

The pictures are drawn into a separate XII window.

Replace, Control

Replace (control (ngraphic = 10))

This command can be used, to force DIOS to redraw a picture every 10 time steps and at the end of each process step.

The parameter record Control is used for general control purposes, in particular to specify data for the grid adaptation. The parameters can be specified in the Replace command at any time after the Title command:

Replace (Control (name = value)).

The parameters of the control record can often also be specified locally just for one process step in most of the command:

Diffusion (Control (LPRot = 2)).

Reflect

Reflect (reflect = 0.0)

With DIOS only half of a symmetric structure may be simulated. The Reflect command can be used to expand, shrink, shift or reflect the layer structure, the grid, the functions defined on the grid and, if possible the refinement rectangles. The symmetry line must be outside of the structure, resp. and the chosen window. The reflect command can be used repeatedly, but only one reflection at either the left or right side is allowed per Reflect command. To get the full device structure as shown in Figure 2.12, the simulated structure is reflected around y axis. Figure 2.13 shows the full device structure with the triangle option put to *off*.

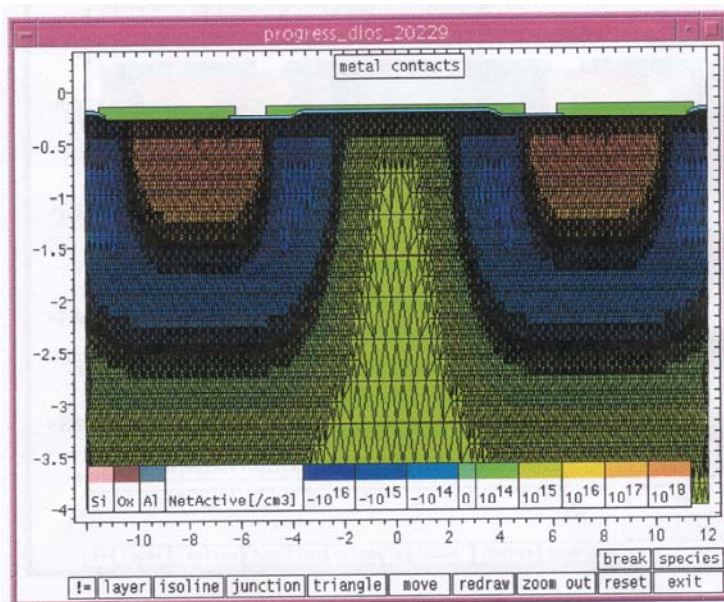


Figure 2.12: Full device structure

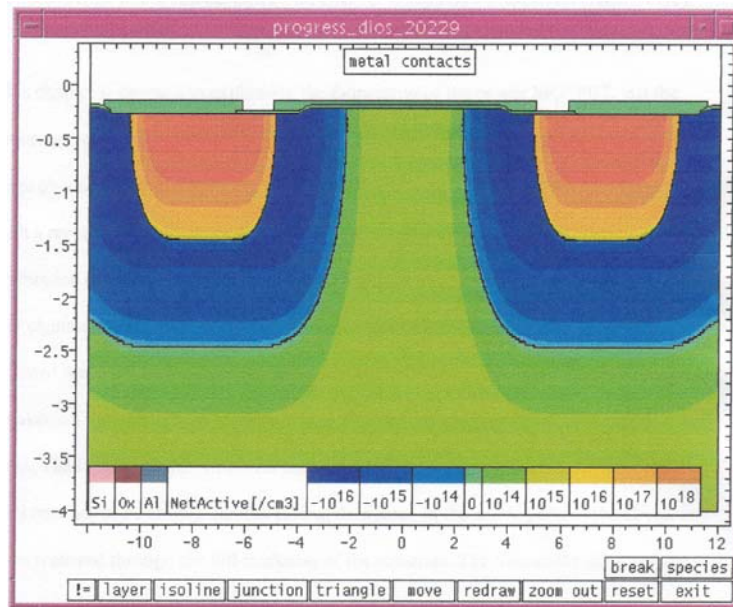


Figure 2.13: Full device structure with triangle option off

2.4 DMOSFET FABRICATION AND DESIGN CONSIDERATION

This chapter is devoted to explaining the fabrication of the power MOSFET. All the processes involved in fabricating the device are described in brief. In order to achieve a large channel width for good on-state characteristics, power MOSFETs are fabricated with a repetitive pattern of small cells. For the fabrication of power MOSFETs, the wafer orientation can be chosen to provide the highest surface mobility in order to achieve a low channel contribution to the on-resistance. DMOSFETs are fabricated on (100) oriented wafers. Due to the extremely rapid increase in on-resistance with increasing breakdown voltage, power MOSFETs are confined to blocking voltages of less than 1000 volts. The ideal depletion width for these devices is less than 100 microns. The wafer thickness has to be chosen carefully. Heat dissipated in the active parts of the device has to be removed through the full thickness of the substrate. The thicker the slice, the greater is the thermal resistance of the final device. However, too thin a wafer will lead to breakage during manufacture and a lower yield. Generally, wafer thickness between 250 and 500 μ are used. The wafer thickness also adds to R_{ds} (on). However heavy doping with phosphorus ensures that the resistivity is low. DMOSFETs were fabricated with class 100 clean room conditions. Following steps were performed for fabrication.

Design of Masks for DMOS

From the study of cross sectional view of the DMOS structure, it was determined that four mask levels are required to produce the device. Mask level one is used to open window in the thermally deposited silicon oxide for the boron diffusion used for the p-base regions. Mask level two is used to open the window in the silicon oxide for the phosphorus diffusion for n+ source regions. Mask level three is used to open the contact windows for the aluminium metallization. Mask level four is used for metal patterning to form the gate and source contacts.

All the four levels of masks were first designed using Corel Draw. The final design was then made using Auto-Cad. The masks were drawn originally on a 100x scale. This design was then sent to Photomasks Company in Canada (Adtek Photomasks). The drawings were reduced 100 times and then produced on a Chrome glass plate of 4 X 4 inch. All the four levels of masks were arranged on the 3 x 3 inch area of the glass plate to fit accordingly in the Karl-Suss mask aligner in fabrication lab. Figure 2.14 gives overhead view of each mask. The dimensions shown are in centimeter. There are totally 120 devices arranged in 10 columns and 12 rows. The first block of 40 devices has channel length of 5μ , the second one has 10μ and the last block has 15μ as channel length. Devices with different channel lengths are fabricated to study the characteristics of devices with varying channel length.

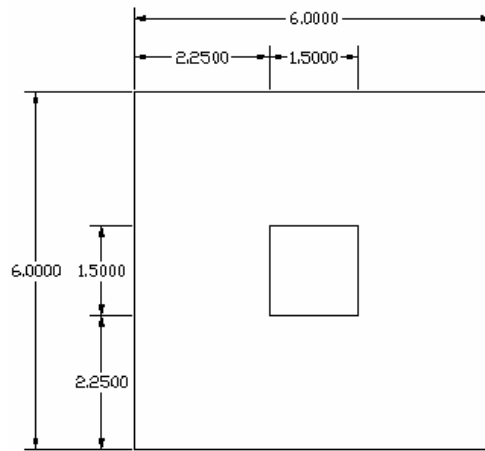


Figure 2.14: (a) Mask –1 for channel length of 15μ

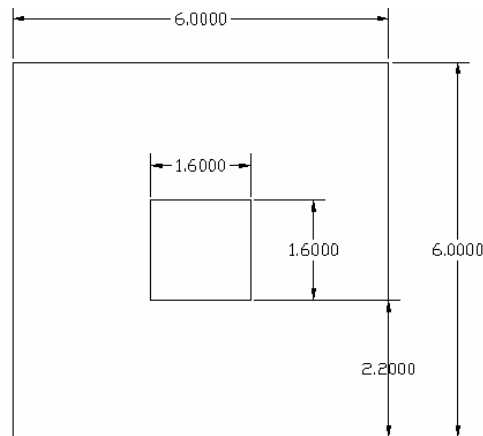


Figure 2.14: (b) Mask-1 for channel length of 10μ

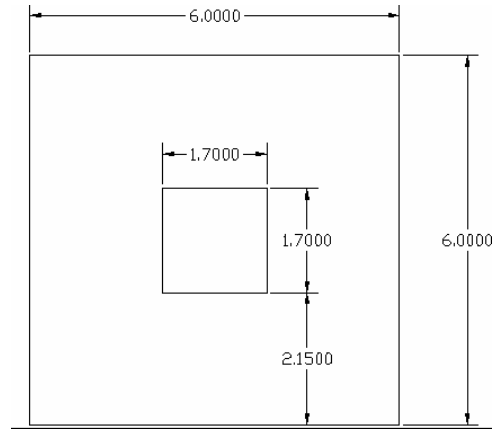


Figure 2.14: (c) Mask-1 for channel length of 5μ

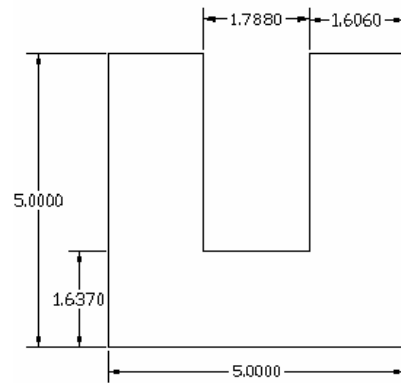


Figure 2.15: Overhead view of Mask-2

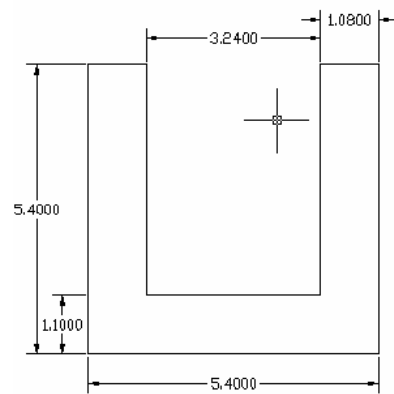


Figure 2.16: Overhead view of Mask-3

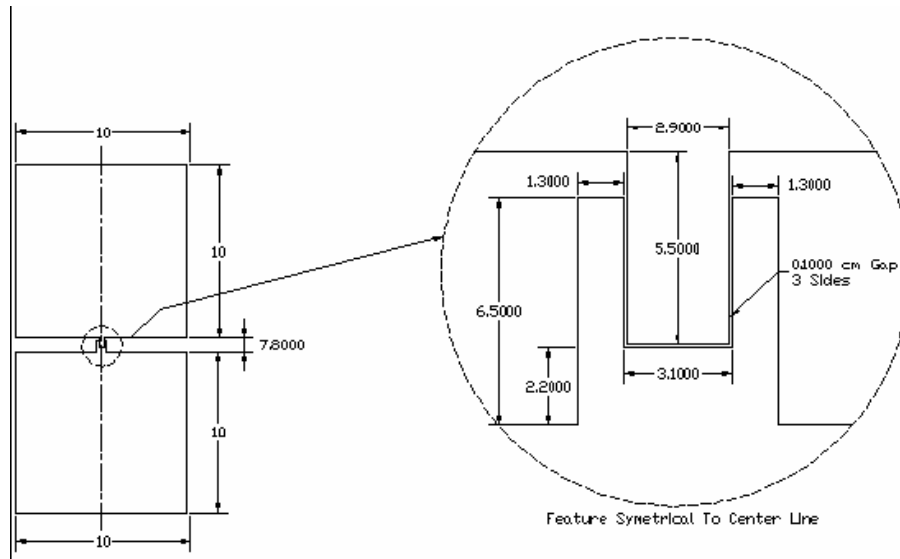


Figure 2.17: Overhead View of Mask-4

Power MOSFET Fabrication

The complete step by step procedure of the fabrication of DMOS is given in the appendix. Here the processes will be described in brief. The silicon wafer selected was a $2 \times 10^{14} \text{ cm}^{-3}$ n-type epilayer ($\rho = 17 - 23 \text{ ohm-cm}$) with a thickness of $3.5 \mu\text{m}$ on a n+ silicon. The orientation of the wafer was $\langle 100 \rangle$. The wafer was cleaned as per the steps given in Appendix B. A thermal oxide of $4000 \text{ \AA}$ was grown on the cleaned wafer at 1100°C in wet oxidation furnace for 27 minutes.

Open windows using First Mask

Negative photoresist was then spun for 30 seconds at 3000rpm and soft bake was performed for 3 minutes. First mask was used to open windows for p+ region. A Karl Suiss mask aligner was used and the wafer was exposed for 10 seconds using UV light.

The photoresist was then developed for 5 minutes 30 seconds and then was observed under microscope to see if the patterns were clearly visible. A buffered oxide etch was then performed for 7 minutes to etch the oxide. The photoresist was washed away with acetone and deionized water. The wafer was inspected for proper opening for p-region.

Ion Implantation

Dry oxidation is then performed for 51 minutes to grow an oxide of $1000 \text{ \AA}$. This oxide thickness makes the implanted peak at silicon- silicon dioxide interface. Ion implantation of boron is done to get a lower concentration at the surface. A surface concentration of 10^{16} is obtained by implantation whereas if boron diffusion is done by thermal diffusion, a high surface concentration of the order 10^{19} is obtained. The calculations for dose as shown in Appendix A are done. Here the implanter energy is 30keV. The values of Projected range (R_p) and Normal Straggle (ΔR_p) are found from the table for the corresponding energy. The value of R_p is found to be $0.0987 \mu\text{m}$. The wafer was sent to Core Systems, California for ion implantation. All the required data for implantation was supplied to the company.

Boron Diffusion

Boron drive-in was performed for 58 minutes after implantation. To find the diffusion time calculations were made as shown in Appendix A. All the calculations were done for a junction depth of 1.5μ. This diffusion was carried out in dry oxidation furnace. So, little oxide grows on already deposited oxide.

Phosphorus predep

A similar procedure was performed for the second mask level as for the first one for the n-region opening. A phosphorus predep is performed for 10 minutes. All the grown oxide is then etched by performing a borosilicate etch for 10 minutes.

Gate Oxide

Dry oxidation is done for 7 minutes to grow an oxide layer of 350°A as gate oxide.

The wafer is subjected to oxidation for the above-mentioned time, which is found by doing process simulation. This oxidation process acts as drive-in for phosphorus.

The surface concentration of phosphorus is found to be 10^{19} .

Contact Windows and Metallization

The third mask level is used to open contact windows using the same procedure as for the first and second masks level. Aluminum is deposited on the entire wafer. For the fourth mask level, positive photoresist is spun for thirty second. Soft bake is performed for 3 minutes. The fourth mask is then aligned with all the other mask levels and is exposed for 10 seconds. Photoresist is then developed for 4 minutes and 30 seconds. After inspection of the developed pattern, the wafer is hard baked for 3 minutes. To get the required pattern aluminium is etched using aluminum etch. Photoresist is then removed by acetone and deionized water. Aluminum is also deposited on the backside of the wafer to form the drain contact. In order to ensure good contact formation, aluminum was annealed in an inert (nitrogen) atmosphere at a temperature of 400° C for 30 minutes following deposition and patterning.

Boron Implantation and diffusion

Some considerations were taken into account not to punch through the Silicon epi-layer while diffusing Boron and Phosphorus into the epi-layer. The thickness of epi-layer was 3.5μ, having concentration of $2 \times 10^{14} \text{ cm}^{-3}$. Boron was thus diffused 1.5μ deep, keeping in mind that the junction will move deeper as various high temperature processes are carried out. Taking the junction depth as 1.5μ calculations were made for diffusion time for boron drive-in. Boron diffusion follows Gaussian distribution.

$$x_j = 2\sqrt{Dt \ln \frac{N_0}{N_B}} \quad (2.20)$$

Where

D = Diffusion Coefficient

N_B = Background Concentration

N_0 = Peak Concentration

t = Diffusion time

Diffusion coefficient was found out from the following equation:

$$D = D_0 \exp\left(\frac{-E_A}{kT}\right) \quad (2.21)$$

Where

$$T = \text{Temperature}$$

The dose is given by following equation:

$$Q_{diff} = N_0 \sqrt{\pi Dt} \quad (2.22)$$

The total dose for boron implantation is double of the above dose.

The boron concentration after implantation is given by (Jaeger, 1993):

$$N_p = \frac{Q}{\sqrt{2\pi} \Delta R_p} \quad (2.23)$$

Where

$$\Delta R_p = \text{Straggle}$$

Following boron diffusion, phosphorus predep is performed for 10 minutes.

2.5 RESULTS

The contributions to the on-resistance from various terms are dependent upon the device geometrical design parameters. The dominant components of the on-resistance are the channel resistance, the accumulation layer resistance, the JFET region resistance, and the drift resistance. When the gate length (L_g) is small, the JFET and drift region resistances become large due to the small width through which the current must flow into the channel. At the same time, the accumulation layer resistance becomes small because of the shorter path along the surface and the channel resistances become small because of reduction in the cell pitch that is equivalent to an increase in channel density. The opposite trends occur when the gate length is increased. Thus there is an optimum gate length at which the specific on-resistance has a minimum value. It is possible to approach ideal specific on-resistance with higher breakdown voltages.

Temperature Effects

The threshold voltage decreases with an increase in temperature. Carrier mobility and saturation velocity also decrease as the temperature rises, and in the bulk semiconductor regions the resistivity increases. If V_{gs} is kept constant, the results of rise in temperature are as follows: In the subthreshold region, the drain current after pinchoff increases. Before pinchoff the increase in the various parasitic resistances tends to reduce the slope of the I_d - V_{ds} characteristic. In the normal operating region, the decrease of the carrier mobility offsets the consequences of lowering of V_t . As a result, I_d decreases above and below pinchoff. In the velocity saturation region $I_d(\text{sat})$ again decreases because of fall in V_s .

The combination of all these effects is shown in Figure 2.18. $I_d(\text{sat})$ increases with temperature below a certain current level and decreases with increasing temperature at higher currents.

Operating at elevated temperature increases the junction breakdown voltage, by about 1% for each 10° C rise. It is also likely to reduce the expected device lifetime. At higher temperatures, a lower forward bias voltage is needed to support a given current flowing across any of the p-n junctions. The parasitic bipolar junction transistor is more easily turned on at high temperatures, because of the pinch-base resistance is increased, and the voltage needed to forward –bias the emitter-base junction is reduced. At the same time the BJT is rapidly sent into second breakdown failure.

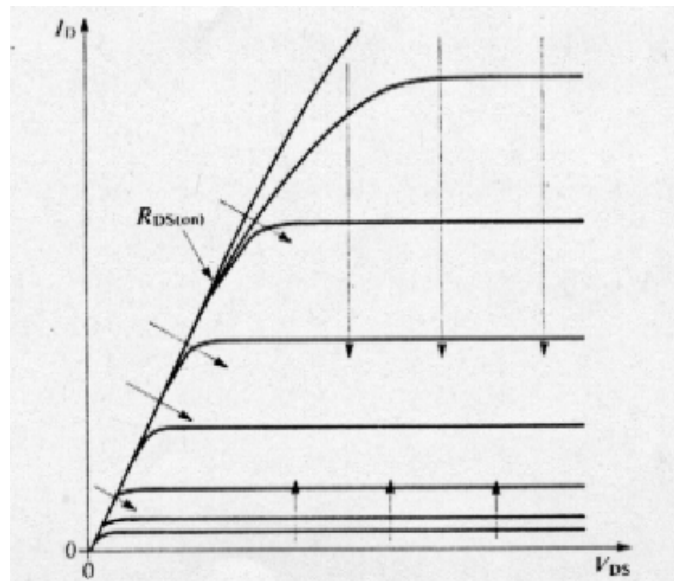


Figure 2.18: Effect of Temperature on Device Characteristics (Arrows indicate the shift of the different parts of the characteristics with an increase in temperature.)

Experimental Results

Eight samples were fabricated using different conditions and were tested using the curve tracer. The following parameters were changed in the fabrication process of the eight wafers.

Table 2.1: Values of different parameters used in fabricating the wafers

Wafer	1 (core)	2 (core)	3 (nc)	4 (core)	5 (nc)	6 (nc)	7 (nc)	8 (core)
Boron Drive-in Time (minutes)	58	30	30	40	35	35	35	40
Gate oxide thickness (Å)	350	300	300	900	700	400	380	350
Gate oxidation temperature (° C)	1100	1000	1000	_____	_____	1000	1000	1000
Gate oxidation time (min)	7	32	32	_____	_____	50	45	40

The fabricated wafers were tested using Curve Tracer. The devices did not behave as power MOSFETs. The testing results are summarized for all the wafers.

Wafer 1

Some devices behaved as diodes and some as resistors. When gate terminal was removed the device still showed characteristics of a diode. The resistance values were tested between source and drain, gate and drain, source and gate, to check if they have been short. Some devices showed resistance in the range of few K ohms and hence they were shorted. For a device to show MOSFET characteristics, the resistance between all the three terminals should be infinite (or should be in range of Mega ohms).

The resistance values noted in few devices of this wafer are shown in Table 2.2.

Table 2.2: Values of resistances for different devices on a fabricated wafer 1

	Gate – Drain resistance	Gate – Source resistance	Source – Drain resistance
Device 1	9.7 K	20 M	20 M
Device 2	1.3 M	1.3 M	12.7 K
Device 3	16 K	20 M	20 M
Device 4	50 K	25 M	25 M

The maximum gate voltage that can be given on the curve tracer is a 2 V. Since the threshold voltage of the power MOSFET might be above 2 V, an external power supply was connected to the gate, wherein the gate voltage was applied in the range of 3 – 6 V. Only one device showed the characteristic of MOSFET. The gate voltage applied was 3.5 V. As the yield was less the device was destroyed while carrying out different tests on it.

Wafer 2

Most of the devices on this wafer behaved as resistors. The resistance values are noted in Table 2.3.

Table 2.3: Values of resistances of devices on fabricated wafer 2

	Gate – Drain resistance	Gate – Source resistance	Source – Drain resistance
Device 1	17 M	20 M	20 M
Device 2	25 M	16.5 M	20 M
Device 3	27 M	20 M	24 M
Device 4	20 M	16 K	25 M

Even though some devices showed no shorts between any pair of terminals, the S-D terminals behaved as diodes without any gate voltage.

Wafer 3

The devices behaved as a diode without the gate voltage. When the gate voltage was increased slowly the curve moved from a diode to a straight line showing that the device now behaves as a resistor. The resistance values of some of the devices are shown in the Table 2.4

Table 2.4: Values of resistance of devices on fabricated wafer 3

	Gate – Drain resistance	Gate – Source resistance	Source – Drain resistance
Device 1	190 Ω	Infinite	Infinite
Device 2	0.8 K	1 K	0.2 K
Device 3	0.7 K	1.3 K	1 K
Device 4	5 K	20 K	15 K

From this wafer only one device worked as a MOSFET. The gate voltage applied was 3 V. Testing of the device at liquid Nitrogen temperature could not be done as the device got shorted.

Wafer 4

Almost all the devices on this wafer showed same characteristics. In the absence of gate voltage they behaved as diodes and when gate voltage was applied they behaved as resistors. The resistance measurement between all the terminals showed that the terminals were shorted. The resistance values of few devices are given in Table 2.5.

Table 2.5: Values of resistance of devices on fabricated wafer 4

	Gate – Drain resistance	Gate – Source resistance	Source – Drain resistance
Device 1	0.5 K	0.4 K	39 Ω
Device 2	0.8 K	0.8 K	0.08 K
Device 3	4.6 K	18 K	12.5 K
Device 4	5 K	20 K	16 K

Wafer 8

One device from this wafer showed the characteristic of a MOSFET. The device had 15 μ channel length. It was also tested at liquid Nitrogen temperature. The on-resistance was calculated from the curve obtained and was found to be in the range of 5 K Ω . This high resistance may be due to poor ohmic contact. The gate voltage applied was 2.5 V. The characteristics of the device at room temperature and liquid Nitrogen temperature are shown in Figure 2.19 and Figure 2.20 respectively.

Possible Reasons for In-Operability of DMOSFET

There can be various possible reasons for the devices to not work as a MOSFET. Some of them are summarized as follows:

The thickness of the epilayer is 3.5 μ . The junction depth of the p-well (Boron) is 2 μ . But during the drive-in process at 1100 $^{\circ}\text{C}$ in dry conditions and other high temperature processes, there is a possibility of the junction moving up to 3.5 μ and touching the epilayer. Due to this the device will behave as a diode instead of MOSFET.

During the gate oxidation process, the phosphorous region can move lower and touch the p region (Boron). Due to this the device will not behave as a MOSFET.

If the gate oxide is too thin it will rupture and there will be a leakage current. So the device will be independent on gate voltage.

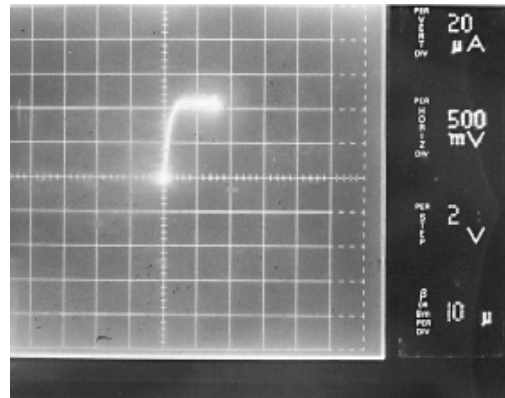


Figure 2.19: Characteristic of MOSFET at room temperature.

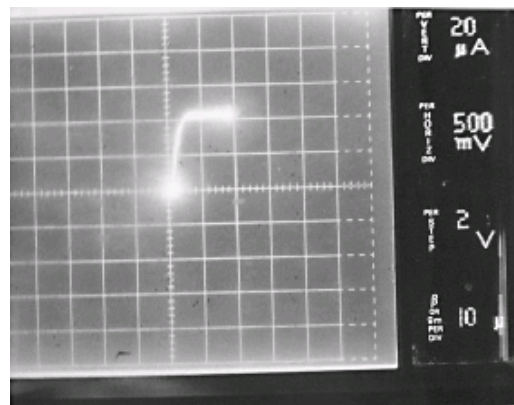


Figure 2.20: Characteristic of MOSFET at liquid Nitrogen temperature

When there are no proper ohmic contacts for the S, D, and G the characteristics shows as a diode instead of MOSFET.

2.6 CONCLUSION

Process simulation of power MOSFETs was carried out using ISE-TCAD simulation tools. The results obtained from the simulation were used as a guideline to fabricate a DMOSFET.

Number of wafers were fabricated with different parameters. The diffusion times and the gate oxide thickness were changed for each wafer. Ideally, it is aimed that all the devices fabricated show the characteristics of a MOSFET with some variations. However, in reality, due to some factors such as precision in fabrication stages like mask alignment, oxide etching, etc., it is very difficult to obtain a high yield. As a result, the real devices would show huge variations. Many of the devices showed diode characteristics and some behaved as resistors. Surprisingly enough, only two such devices behaving as a power MOSFET were obtained.

A number of tests were performed on the devices fabricated. The I_d - V_d curves were observed on the curve tracer. In terms of further work, these results can be used to get a better yield of working MOSFETs on a wafer. An epilayer of greater thickness than 3.5μ , say around 5μ or 10μ can be used.

3. RESEARCH AND DEVLEOPMENT OF COMPACT SPRAY NOZZLES

3.1 INTRODUCTION

An important growing field in today's engineering is the development of Micro-Electrical and Mechanical Systems (MEMS). Over the years there have been many different fabrication processes that have been developed to meet these new demands. Most of these fabrication techniques work well for simple two dimensional components. However, creating parts with a high aspect ratio, the ratio of the height to the width and/or length, is difficult when using these traditional manufacturing methods. The demand for more complicated three dimensional parts requires that a new approach must be taken to make these ideas a reality.

A novel fabrication of parts was developed in the late 1980's called stereolithography. This process involved using a low power laser to introduce free radicals in a polymer in which the molecules would cross-link and become solid. Using this idea the stereolithography process was developed by 3D Systems to fabricate parts in a layer-by-layer manner.

This process was developed to be used to make simple prototypes used to visualize parts and test for fit and function. When the 3D Systems Stereolithography Appartus (SLA) system was released in 1987 it lead to a new classification of manufacturing processes called Rapid Prototyping (RP). Today the RP industry has many different methods for creating 3D parts in a very short period of time.

Building parts in a layer-by-layer manner has distinct advantages over traditional methods. Some things can be made on a RP machine that cannot be made any other way. For example, creating the parts layer-by-layer makes it possible to build a chain with no cuts or breaks, hence adding strength. Also very complex internal geometries can be made without having to have external access for tools. Also, since the parts are built layer-by-layer there is not a size limit on the number of layers that can be built. This means that parts with high aspect ratio can be built easily.

Due to these advantages manufacturing MEMS using a similar process would allow for more complicated and useful parts to be made. Research for a micro-scale stereolithography apparatus has begun. This process is known as microstereolithography.

This paper discusses the previous research in the field of microstereolithography. From this literature an MSL system was designed and built. The methodology of the design process is laid out and a final design is built. Various results and parts were successfully made on the MSL system and were measured to determine the manufacturing error.

Currently there are a few different organizations that are working on developing MSL systems to meet the current demand. In order to understand these processes a good base knowledge is needed of the stereolithography method. In the late 1970's and the early 1980's three different individuals began research in selectively curing a photopolymer to create layers that could be stacked to make three dimensional parts. A. Herbert of 3M and H. Kodama of the Nagoya

Prefecture Research Institute in Japan had trouble maintaining funding for their research and were forced to stop their work in this area. However, C. Hull of Ultra Violet Products, Inc. (UVP) continued to work until he finally developed a complete system that could automatically build detailed three dimensional parts directly from a computer model. This system was patented and Hull, R. Freed, and the stockholders of UVP founded the company 3D Systems, Inc. The SLA-1 was first introduced in November of 1987. This was the beginning of a growing technology that would begin to take off over the next few years. The most common rapid prototyping system found today is the SLA-250 which was released in 1989 (Jacobs, 1992).

This traditional stereolithography process occurs in many steps. First a model of the part must be drawn in the computer. Almost any three dimensional CAD (Computer Aided Design) program can be used, such as AutoCAD, Pro/Engineer, IDEAS, Solidworks, or Catia. After the model is drawn it must be saved into a different file format called STL(Standard Tessellation Language). This is a simple format that converts the complex surfaces into simple triangles. This method provides a simplified approximation to the actual geometry. The accuracy of the mesh is altered by the distance at the midpoint, also known as the chord height (Figure 3.1). Usually, this value is around .0005 units.

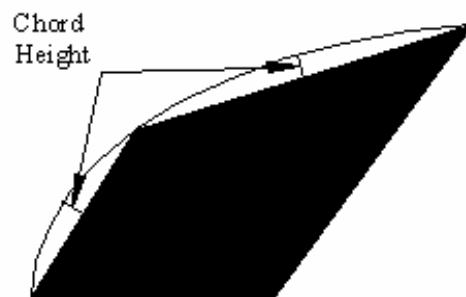


Figure 3.1: STL file estimation of a curve

The smaller the mesh means the more accurate the final result will be, however the file size will also increase and the calculation time for the slicing and hatching algorithms will also increase. With current computer power these are not important issues as they were in the early 1990's. Figure 3.2 shows how a given solid model is converted into STL file.

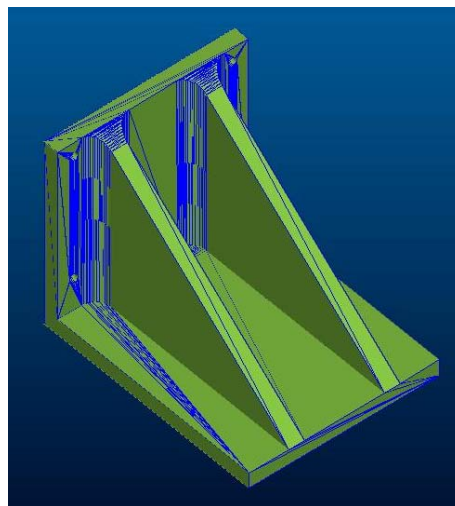


Figure 3.2: Typical STL file

After the appropriate STL file is made then the support structure must be calculated and also made into a STL file. Support structure is built in order to maintain part accuracy and for easy removal from the platform. Usually this structure is simply a cross pattern of lines to support the model when a layer is built as a cantilever or island. The support structure and model are then sliced into the different layers. The slicing algorithm calculates the laser vectors and speed that is necessary to get the desired layer thickness. After slicing the STL files the SLA-250 requires the user to input other criteria and then the build process begins.

The first step in the build process is that the level of the polymer must be exactly known. The machine has a laser sensor that measures the current polymer height and adjusts it with simple displacement plunger attached to the side of the build vat. For each layer the polymer level is measured and the platform is placed one layer thickness below the current level. A recoating blade then sweeps across the vat to level the viscous polymer to a uniform surface.

The stereolithography machines that are commercially available operate in a vector-by-vector manner. This means that a laser light source is reflected by a mirror. This mirror is controlled by the computer and reflects the laser beam to trace out a vector on the polymer surface. The layer is built by tracing the border of the layer and then hatching in the surface. This process can be slow for large cross sections because the laser can only cure its diameter therefore hatching a large area takes many passes of the laser beam.

After the layer has been completed the entire model is dipped into the polymer deep enough to cover the entire part. It is then raised back to one layer thickness below the polymer surface and the process is started again. After all the layers have been completed the part is removed from the polymer and the support structure removed. Also the laser doesn't fully polymerize the liquid so a post cure is done in a UV oven. Once the part is fully cured then usually there is some sanding and polishing of the part.



Figure 3.3: 3D Systems SLA-250

3.2 LITERATURE REVIEW

This chapter will discuss the current progress made by other institutions and companies on the MSL technology. These machines fall into different categories depending on how the micro-parts are made. Some of the systems are very similar to a traditional stereolithography machine and uses a raster scanning method. The laser beam is focused to a very small point exactly at the resin surface and the vat is moved by translation stages to trace the border and hatch in the layer. These were the first type of MSL machine.

Scanning Method

One of the first MSL machines was developed at Nagoya University in Japan. This was very similar to a macro scale stereolithography system except that the laser beam is focused and stationary. Instead the polymer is moved to trace the layer in a raster method this is known as the vector-by-vector method.

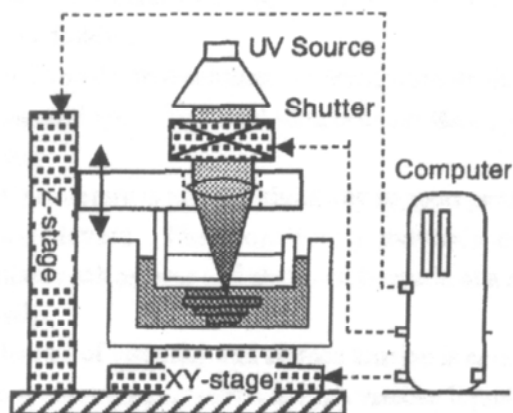


Figure 3.4: IH Process (Maruo, 2001c)

There are a few major things with this method that need to be determined in order to get a high aspect ratio. The main issue is that the width of the solidified polymer along the thickness is affected by the shape and intensity distribution of the irradiated beam and diffraction and absorption of light within the polymer. In order to produce micro-polymer structures with high aspect ratio, a constant width of the polymer along the depth is essential (Nakamoto, 1996). An analytical study of the factors that affect the solidified portion made by the laser was done at Nagoya University. The affects of beam wavelength, the aperture and focal length sizes, absorption coefficient, and defocusing were analyzed. After the analytical values were determined for these parameters experiments were performed to see how accurate the solutions were. Examples of parts made using this machine can be seen if Figure 3.5. These parts are highly accurate however they are very small in size, less than one millimeter long.

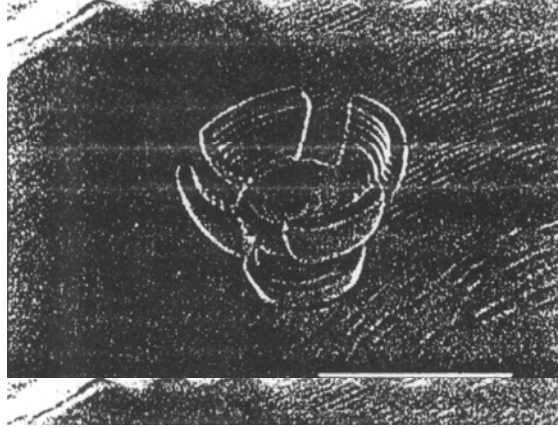


Figure 3.5: IH Process sample part (Nakamoto, 1996)

This small size is the major downfall of this technique. In order to create parts that have larger cross sections in the order of cm^2 it would take a very long time to hatch in a solid layer. A quick estimate to make a 5 cm^2 100 micron layer was done. It is found that with a scan speed of 40 microns/second and a line thickness of 10 microns a 5 cm^2 layer would take over 2 weeks to build. From this point of view it is clear that this application is great for very small and simple parts, but is not feasible to make anything larger than a millimeter.

At Nagoya University another process was developed also. This MSL method is called the Super IH process. This is very similar to the two-photon process also developed at this university. The fabrication system consists of a laser, shutter, galvano-scanner set, xyz-stage, and objective lens. The beam from the laser is introduced into the galvano-scanner set to deflect its direction in two dimensions, and then is focused into the UV polymer with the objective lens (Ikuta, 1998). The focused beam causes polymerization inside the liquid resin at any point. Therefore a true 3D structure can be made directly and not in a layer-by-layer manner.

In this method, the liquid UV polymer is solidified only at the focus, although the laser beam is focused inside the polymer. This is because the solidification reaction isn't proportional to the exposure of the light, even if the initiation reaction of the photopolymerization is proportional to that. The solidification doesn't start until the exposure is over the critical exposure. This nonlinearity to the exposure can make the solidification to be limited near the focused spot where the intensity of the light is higher than out-of-focus regions. (Ikuta, 1998)

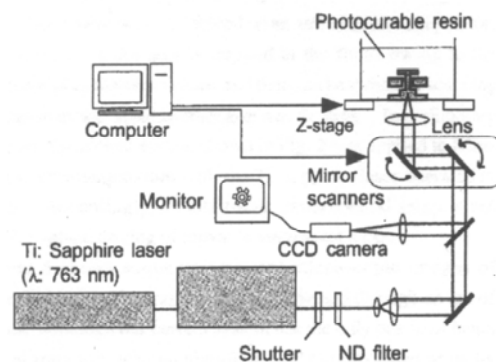


Figure 3.6: Two-photon process (Maruo, 2001b)

The super IH process is not like traditional stereolithography techniques because it doesn't build parts layer-by-layer. Instead the laser beam is focused to the exact point within the polymer that needs to be solidified. The solidified polymer and liquid polymer have roughly the same densities, therefore there is no need for support structure. In other words parts can be made directly and require little post processing. Due to the nature of this process an optically clear polymer must be used in order to control the exposure times, which severely limits material selection. As with the previous method the super IH process cures an incredibly small volume at a time. It would take an extremely long time to fabricate anything larger than a few millimeters. The large benefit is that resolution of this technique is less than 1 micron.

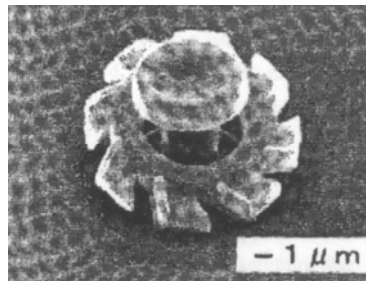


Figure 3.7: Two-photon sample part (Maruo, 2001a)

Integral Process Method

After finding that the vector-by-vector method was not acceptable for fabrication of larger parts another technique was developed for MSL. This technique employed the use of a spatial light modulator. These devices can be found all over the world today in display applications, such as projection television sets, computer projectors, as well as near eye applications, such as heads up displays in military aircraft. A few educational institutions have used these common displays as an “active mask” for a layer-by-layer MSL technique, called the integral process.

The idea of using a mask to make each layer is not unique. This is how current integrated circuit chips are fabricated on silicon wafers. A series of masks are usually printed on soda lime glass and projected onto the silicon wafer or a photoresist to transfer the pattern of the mask. These masks range in size but are usually 4 – 5 inches square and placed into the wafer fabrication machines. A light source is used to alter the chemical structure of the photoresist and therefore transfer the mask pattern to the wafer. After the pattern is transferred other processes can be done such as chemical vapor deposition or etching. This is a standard fabrication process for many MEMS devices. The down side of fabricating a device in this fashion is that for some devices many masks need to be made to obtain the desired geometry. This is where the idea for an active mask began.

One of the current institutions involved in MSL is the University of Sussex in the United Kingdom. This system contains five major components: an ultraviolet laser light source; an optical shutter; a spatial light modulator; a multi-element lithographic lens system; a high resolution translation stage (Chatwin, 1998).

The spatial light modulator is the critical interface between electronic and optical systems (Huang, 1998). For this particular machine a Super VGA (800 x 600) resolution device with a

pixel size of 26 μm x 24 μm . This spatial light modulator is called a transmissive liquid crystal display (LCD). The total active area of this device is 26.4 mm x 19.8 mm (Chatwin, 1998). There are a few flaws in this kind of SLM, first is that the fill factor is only 50% (Chatwin, 1998). This means that only half of the entire SLM are pixels and can be controlled, the other half is the spacing between pixels, which is opaque. This SLM can be run at a maximum of 40 Hz due to the rise and fall time of the liquid crystals (Farsari, 2000). It is stated in the literature that the LCD device is not damaged by wavelengths longer than 350 nm, however shorter wavelengths will cause damage to the SLM's indium tin oxide electrodes and the liquid crystal material (Farsari, 2000). This statement will be discussed in full detail in later chapters.

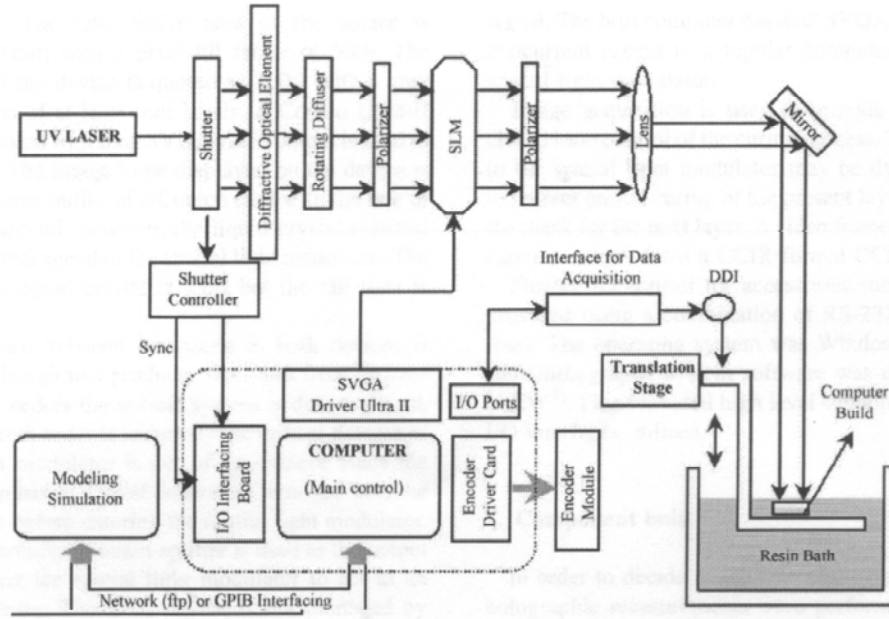


Figure 3.8: University of Sussex MSL process (Huang, 1998)

The second important component in this MSL setup is the light source. The University of Sussex machine uses an argon ion laser with a wavelength of 351.1 nm. This laser beam has a nominally Gaussian intensity distribution and is linearly polarized. In order to obtain a uniform distribution across the SLM an eight-level diffractive optic element was designed by Digital Optics Corporation (Chatwin, 1998). This optical element is considerably more energy efficient than sampling the central portion of the beam using an aperture. Also the low beam quality produced by the diffractive optic element can be improved by the incorporation of a rotating holographic diffuser, which destroys spatial coherence and reduces speckle (Chatwin, 1998). However, this diffuser causes the beam to become depolarized and not collimated.

After the light is repolarized and passed through the spatial light modulator it is then sent through a reduction lens system. This reduces the image from the SLM by two hundred times. This image is then focused on the polymer surface and the entire layer is built in one exposure. In this machine commercially available polymers were used, Ciba-Geigy Cibatool SL 5180 and DuPont Somos 6100 (Farsari, 2000). The laser irradiance was determined using the Beer-Lambert's law.

$$E(z) = E_0 e^{-z/D_p} \quad (3.1)$$

Very precise and intricate parts were created with this system. The layer thickness was 50 μm . The gear in Figure 3.9 was made using 50 layers. Due to the fact that the layers are so thin it is impossible to use a sweeping blade to create the layer of polymer, therefore a wait time was used between each layer. This time is not stated specifically in the literature.

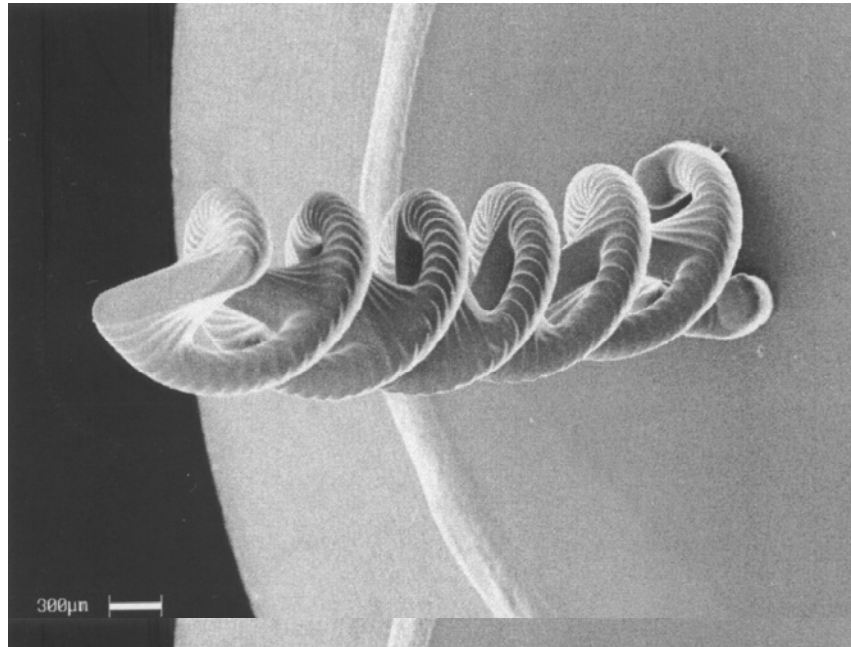


Figure 3.9: University of Sussex sample part (Farsari, 2000)

Another major institution that is developing a MSL machine using a dynamic mask is the Swiss Federal Institute of Technology (EPFL). This system has some of the most extensive literature available to date. This system is only slightly different from the University of Sussex.

The spatial light modulator used at EPFL is a Infocus System 1600GS LCD projection panel (Bertsch, 1997b). This is a VGA resolution dynamic pattern generator (640 x 480) pixels. After the reduction lens it allows parts to be made that have a 2.5 mm x 2.5 mm cross section (Bertsch, 1998). At EPFL they have stated that these LCD displays are damaged by UV light sources and have placed their spatial light modulator between four glass windows which are opaque to ultraviolet light (Bertsch, 1997b). This contradicts the literature published by the University of Sussex.

In order to preserve the LCD a visible light source was used at EPFL. An Argon laser (Coherent INOVA 90) emitting at a visible 515 nm wavelength and in TEM_{00} mode was used (Bertsch, 2001). The laser output power was 1.5 watts. The major disadvantage of this light source is the Gaussian distribution intensity profile. At EPFL this problem was solved by adding a diaphragm along the optical path so that only the central part of the Gaussian beam was used (Bertsch, 1997a). It was stated that once the research using a monochromatic source was complete then it would be worthwhile to use a lamp as a light source (Bertsch, 1997b).

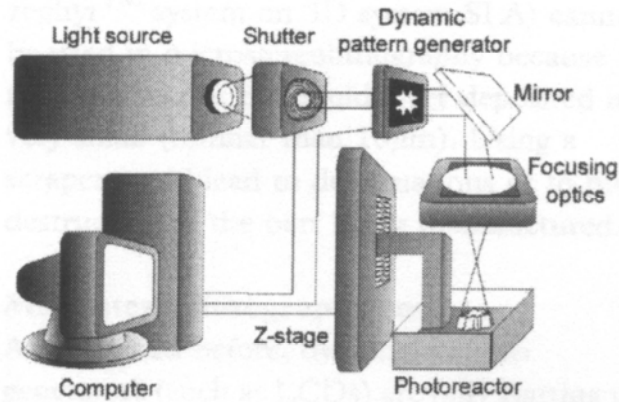


Figure 3.10: EPFL MSL setup (Bertsch, 2000)

This machine was able to build extremely accurate parts. Figure 3.11 shows an example which consists of 673 layers of a 5 μm thickness. The resolution in the x-y plane is determined by the size of the pixels of the pattern generator and the reduction factor of the optical components. With the machine at EPFL a resolution of 5 microns in the three dimensions of space is possible (Bertsch, 1997a).

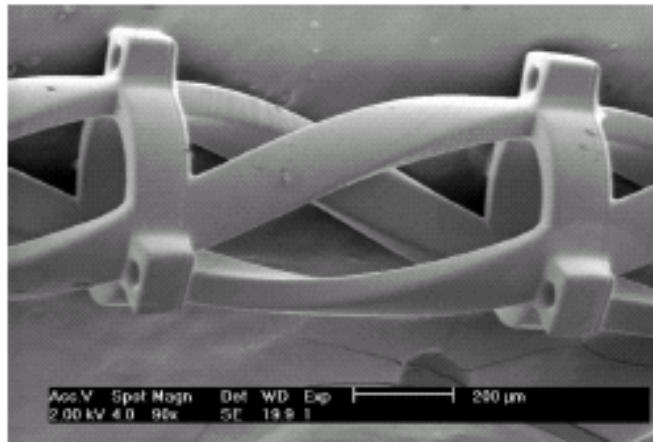


Figure 3.11: EPFL sample part (Beluze, 1999)

As with the MSL machine at Sussex the layer of polymer is leveled using gravity. As the recoating process depends on gravity forces, longer leveling times have to be used for larger cross sections, which also limits the building speed (Bertsch, 2000). The part in Figure 3.11 took roughly three hours to build its 673 layers. Working from these numbers it is found that the time per layer is around 16 seconds. Exposure takes roughly 1 second (Bertsch, 1997b) therefore about 15 seconds are needed to wait for the polymer to level off. Due to this it is necessary to have polymers with low viscosities (Bertsch, 1997a; Farsari, 2000).

3.3 METHODOLOGY

System Requirements

The first step to developing our MSL system was to determine the requirements of the machine. Why were we building this machine? What type of parts do we want to make? How big do we need our build area to be? What resolution do we require for our application? These are the types of questions that needed to be answered before a design was started.

There are many parts that we would like our MSL machine to build, however there is one part that it must build, spray nozzles. The funding for this project is part of a contract to deliver a spray cooling device. These nozzles are roughly 2.5 cm in diameter however only the atomizer needs to be built with a higher accuracy. The atomizer is roughly 12 mm in diameter and needs to be fabricated with a resolution of around 50 μm in order to function properly. Traditional stereolithography cannot meet these build requirements. Also, all of the other MSL machines that were developed at other institutions cannot fabricate anything that is this large. Therefore a new system needs to be developed to build this part.

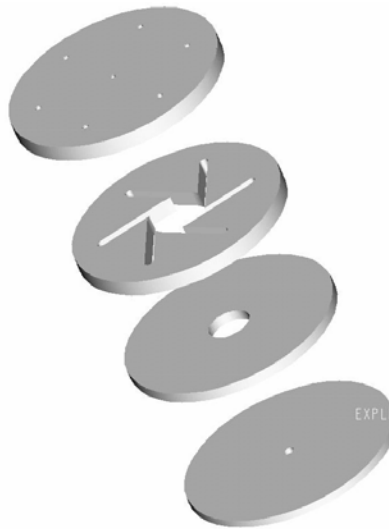


Figure 3.12: Exploded view of spray nozzle

There are other parts that we would like to make on our MSL machine. A micro-turbine is another component that could be realized on our MSL machine. These parts have very complex geometry that is extremely difficult and costly to make by traditional machining. However, this turbine cannot be made out of a stereolithography resin due to the high operating temperatures. As part of the MSL project a high temperature polymer-derived-ceramic is being developed to be used in the MSL process.

The list continues with many other parts that are likely candidates for micro-systems. A ceramic heat exchanger can be built with very thin walls for optimal heat transfer. Also the ceramic is good for low temperatures to be used in the cryogenic cooling units. The main requirements of our system are to be able to make larger parts than other MSL systems and have a resolution of less than 20 μm in the x-y plane and a layer thickness that is 10 μm – 100 μm .

Location

The first challenge that was tackled was finding a place to build our machine. A recently acquired lab was built and it was decided that this lab is where the machine would take shape. With these high accuracy requirements and a complex optical system it was necessary to use an optical table for mounting the MSL system on. A car or truck passing by on the street, high winds, or rain will cause slight vibrations in the building structure. In order to counteract this, the optical table would need to be isolated from the floor so that vibrations from the building would not affect the performance of the machine.

A large 12' x 4' optical table was available however this table weighed 2 tons and was located about one mile away from the lab. The table was taken out of research park and put it in to the engineering building through a window by crane. This was a process that went smoothly because the moving company had experience moving optical tables. After the optical table was assembled in the room the next important step could begin.

Spatial Light Modulator

Choosing a spatial light modulator is the most important step in the design of the MSL machine. This item is the only link between the computer and the optical system. Also, the SLM specifications will be the limiting factor in the accuracy of the machine. There are two types of spatial light modulators commercially available, the liquid crystal display (LCD) and the digital micromirror device (DMD).

The liquid crystal display has been around for many years. Most liquid crystal displays in use today can trace their origins to the invention of the twisted-nematic display. These liquid crystals have only one property untypical for a liquid: their elongated molecules prefer to be aligned with one another (Nelson, 1997). These molecules can be forced to twist relative to one another creating a helical structure known as a twisted nematic. When light passes through this twisted nematic it is rotated through the helical shape. With a polarizer before and after the liquid crystal it is possible to make it appear dark by applying an electric field which forces the twisted nematic to disappear. This concept lead to an array of liquid crystals where a voltage applied to electrodes at the appropriate row and column can turn on any given pixel.

There are two main types of LCD's that are commercially available, transmissive and reflective. A transmissive LCD has two separate polarizers. One is placed before the LCD and one is placed at a 90° angle at the rear of the LCD. The light is passed straight through the LCD and the image is created on the opposite side. These types of LCD's have the one main disadvantage, the electrodes for turning on and off the pixels are located between each pixel. This makes the ratio of active space to dead space lower. This is called the fill ratio. Having a high fill ratio means that more of the screen can be controlled actively and that the space between pixels is reduced. This is where a reflective display becomes desirable.

The reflective LCD has a mirror behind the liquid crystal display. This allows one polarizer to also be used as the analyzer. The electrodes for this type of display can be built into the mirror surface so that a higher fill ratio can be achieved. This high fill ratio is needed for a MSL system, because the gap between pixels may not fully cure the photopolymer.



Figure 3.13: Digital Micromirror Device

The digital micromirror device is much more complex than a LCD. It was developed by Dr. Larry Hornbeck of Texas Instruments in 1987. Even though this technology is relatively young it is beginning to dominate the consumer market of microdisplays for projection televisions and computer projectors. The DMD is a complex MEMS device that is made up of an array of tiny mirrors that can be tilted using electro static forces. These micromirrors are roughly $16\ \mu\text{m}$ square and tilt $\pm 10^\circ$. The digital micromirror device has the ability to operate much faster than LCD's and are extremely reliable. Texas Instruments recently released a new DMD that is not damaged by UV light. The gap between pixels on the DMD is about $1\ \mu\text{m}$, therefore it has a high fill ratio.

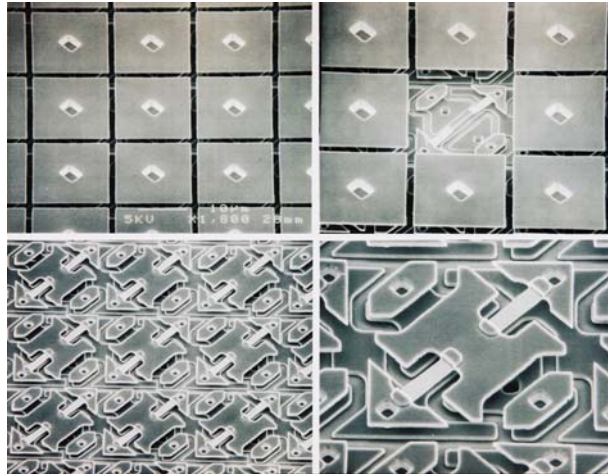


Figure 3.14: SEM photograph of DMD

Choosing between the DMD and the LCD is not an easy task. DMD's are still new, expensive and only available from Texas Instruments; while LCD's are inexpensive and readily available from multiple vendors. In order to decide which type of microdisplay should be used it is desirable to know what wavelength the system will operate at. Due to the complexity of this project it was determined that using as much commercially available material as possible would benefit the project. Currently, there are many macro-scale stereolithography machines that operate in the UV range, because the photoinitiators work better at these shorter wavelengths. Since it was possible to buy off-the-shelf photopolymers in the UV range it was desirable to use this wavelength for our MSL system. Keeping this in mind we begin to look at different light sources for the MSL machine.

Light Source

This process is one of the most difficult because the different light sources, such as lasers, are extremely expensive. From our literature search it is seen that some institutions are using visible light sources while others are using UV sources. However, all previous research was using laser light instead of a lamp or other illumination source. Due to this only laser light sources were considered for use in our MSL system. There are two types of lasers that may be able to be used for our system, solid-state lasers and gas continuous wave lasers

First it was necessary to calculate the amount of power our laser would need in order to illuminate the entire spatial light modulator surface. To get an estimate we used the properties of a commercial polymer to determine the amount of energy needed to make a 250 μm layer. These numbers are in the range of 50 mJ/cm^2 for a 250 μm layer. In order to illuminate the spatial light modulator, which will be no larger than 25 mm by 25 mm, a spot of a diameter 3.5 cm will be needed. Almost 40% of the beam is wasted because the beam must illuminate the entire SLM. This new beam is roughly 10 cm^2 in size, therefore for a one second exposure time the average power of the laser should be at least .5 watts. There will be high losses through the beam shaping process and the various optical components in the system. Also, this is a commercial polymer made especially for laser exposed polymerization. The custom PDC material that we will use with the system may need much more energy to polymerize in a short amount of time. Due to these considerations it was determined that any laser that was over 2 watts should be able to satisfy our needs. For our machine there were two specific lasers that were being considered. Both lasers were roughly the same cost. The first option is a continuous wave gas laser. It is known that the shorter the wavelength the more reactive the photoinitiators are. The shortest wavelength available without entering the UV range is a Krypton gas laser that has a strong line of power at 413.1 nm. At this wavelength an average power of 2.5 watts can be delivered. This laser also produces many other wavelengths so a narrow band filter would need to be used to block the unnecessary wavelengths. Typically lower power gas lasers are simple to use and operate, however these larger lasers need tremendous amounts of power. This laser requires a 480V 3-phase power source and access to a chilled water line or a chilling system and pump. These lasers also have a certain life until the gas chamber needs to be replaced. The benefit of gas lasers is that the output is continuous wave. This means that at anytime a measurement is taken from the laser the average power is very close to the instantaneous power. Also, due to the visible wavelength the beam can be seen clearly and optical alignment can be done easily. Lastly, traditional inexpensive optics can be used that are usually made from BK7 optical glass.

The other laser choice is a diode pumped solid state laser. This is an IR laser that passes through a doubling and tripling crystal which converts the 1064 nm to a single wavelength of 355 nm. Due to the single wavelength no additional filters would be needed. At this wavelength the average power can be up to 4 watts. This average power can vary depending on the repetition rate used by the q-switch in the laser cavity. Unlike the gas laser the solid state laser can use a standard 110V or 220V outlet. It does require a very small water cooling system that comes with the laser head. These lasers also require the diodes to be replaced after approximately 2000 hours of use. Due to the UV wavelength more expensive components will need to be used throughout the optical system. For UV light usually fused silica is the material that the lenses are made from.

Optics and Opto-Mechanics

The beam from these lasers is small (around 1mm) and has a Gaussian intensity profile known as TEM_{00} . These two factors need to be changed in order to create a beam that uniformly covers the area of the SLM. We can begin with the complex task of reshaping the intensity profile of the beam into a uniform or square-top beam.

The first method of creating a uniform beam, known as intensity apodization, is to just absorb more of the higher intensity beam and let the lower intensity edges pass through. This uses a plano-convex gray lens cemented to a plano-concave clear lens. Of course this design has a major disadvantage, only 30% of the beam is transmitted through the system. This fact makes this method to shaping a beam from gaussian to uniform a very poor approach.

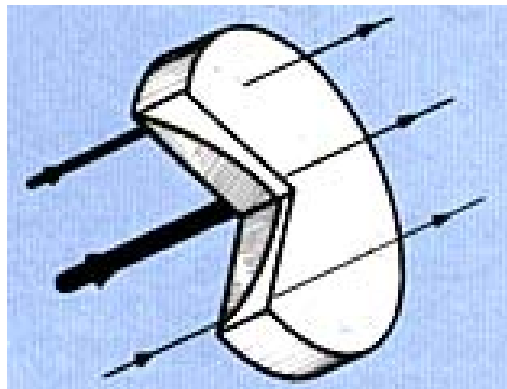


Figure 3.15: Lens used for intensity apodization

Using geometrical methods to shape a laser beam profile involves application of geometrical optics to solve the optical design problem. Specifically, the laws of reflection and refraction are used along with ray tracing, conservation of energy within a bundle of rays, and constant optical path length condition to design a laser beam profile shaping optical system. Interference or diffraction effects are not considered as part of the design process. Only lenses and mirrors are used for the optical components of the laser beam profile shaping system. The design results in two custom aspherical lenses that must be made specific to the wavelength and size of the beam.

The last method of beam shaping is to use a diffractive optic element. This design is based on a Fourier Transform relation between the input and output beam functions. This solution can be obtained using geometrical optics however, the diffraction approach introduces a parameter that contains the product of the widths of the input and output beams. The efficiency of the solution is shown to depend on this parameter. The quality of the solution improves asymptotically with increasing value of this parameter. These elements are much more difficult to design than geometrical methods and still require custom optics to be made for the specific laser.

A two lens system was designed for our setup but due to the cost involved in manufacturing custom aspheric lenses it was decided that temporarily it was best to just sample the center portion of the beam where the intensity can be assumed to be approximately equal. The beam will be expanded to larger than necessary and only the center portion of the beam will be used.

Next a two lens beam expansion is needed to make the beam large enough to illuminate the entire SLM. This is a simple procedure that involves using two plano-convex lenses. The two lenses are designed to create a given expansion ratio. The expansion is determined by the focal length of the lenses. For the laser to illuminate the entire SLM it must be expanded 20 times. The smallest focal length lens available was 10 mm (nominal). Therefore the large lens must have a focal length of 200 mm. These lenses would be separated by the sum of these distances to create the beam expander.

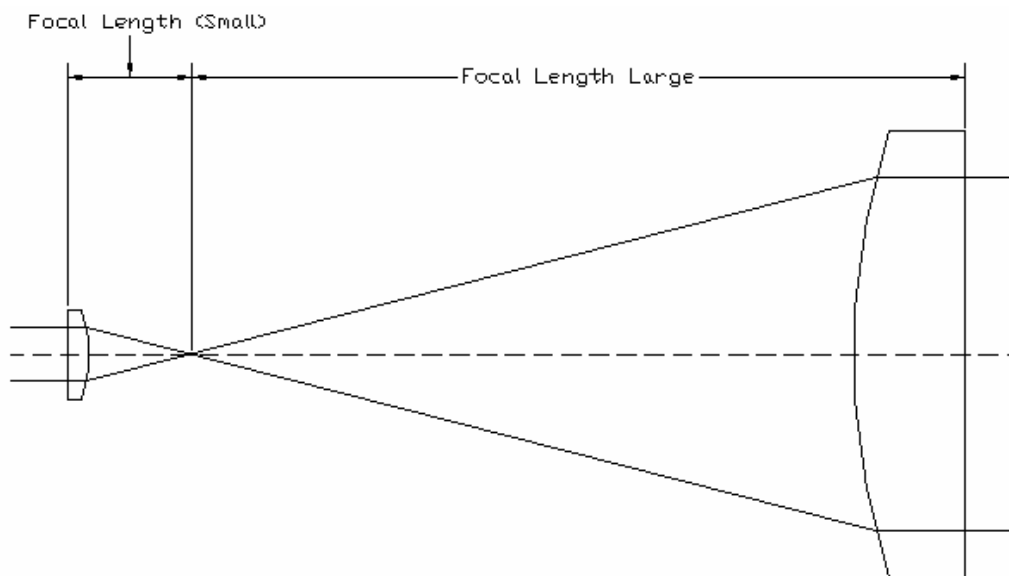


Figure 3.16: Two lens beam expander

The beam expanding lenses need to be mounted in special lens holders that will be used to align the lenses. The smallest lens is the most critical component to have aligned with the laser beam. The mount for this lens will require an 5 axis mounting system that allows for adjustment in the X, Y, Z, θ , and ϕ directions. These adjustments will be used to align the laser expansion along the optical path. The large lens is less critical and will only need a simple X-Y adjustment. These lens holders will need to be mounted onto an optical rail with graduations of millimeters marked. The lenses need to be a precise distance apart to create the beam expansion. The rail is used to keep the optics aligned and used to define the coarse distance between the lenses. The fine distance adjustment is done with the small lens holder. Now that the beam is large enough to illuminate the SLM a light engine needs to be created to do the optical imaging.

Light Engine

The series of optics that are used to illuminate the SLM we call the light engine. Depending on the type of SLM there are different ways of creating the light engine. The simplest possible light engine is using a polarizing cube beam splitter along with a LCD. This system will look like Figure 2.17. The polarized laser light initially passes through the beam splitter. When the light is reflected back through the LCD the polarization of the “on” pixels is rotated 90 degrees. The pixels that are “off” absorb the light and do not reflect. When the reflected image re-enters the beam splitter it is now reflected orthogonal to the original path. An image is created using this light engine along with a LCD.

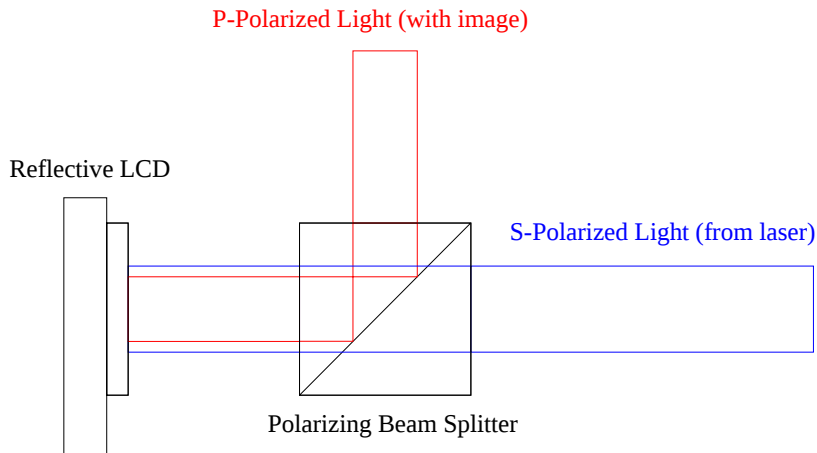


Figure 3.17: Imaging optics for LCD

A DMD requires a much more complex optical system. With a DMD no light is absorbed therefore it is necessary to have the “on” image projected in one direction while the inverted image is sent into a dump. The reason this is complicated to do is because the pixels only rotate ± 12 degrees. This means that a series of prisms with custom angles need to be made to project the image to a plane. This system of optics (Figure 3.18) is known as an offset method for imaging a DMD. There are other methods of creating a light engine for a DMD such as a total internal reflection (TIR) prism. These TIR prisms (Figure 3.19) are more complicated than the offset system because custom optics need to be fabricated at specific angles for the specific DMD.

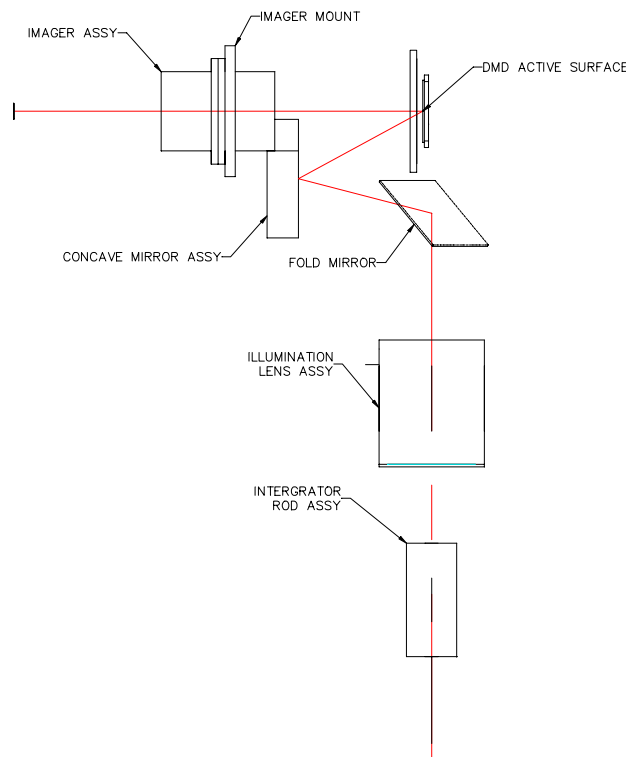


Figure 3.18: Imaging system for DMD

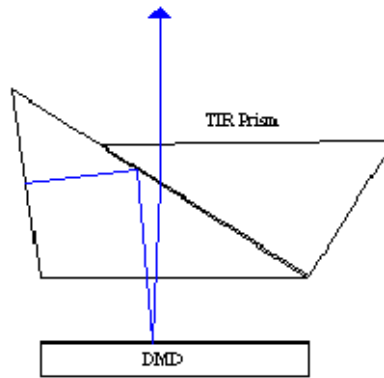


Figure 3.19: TIR prism setup

Translation Stages

The translation stages are very important to the overall accuracy and capability of the machine. For a traditional stereolithography machine only one stage is needed, the Z-stage. This stage travels in a vertical direction and makes it possible to make the layer thickness. In our MSL machine it is desirable to add the ability to move the vat in the X-Y plane. The first reason for this is that it makes it possible to make multiple parts in one build. Also, if the stages are very precise then a “stamping” method may be possible to create parts that are much larger than the size of the SLM but will still have the accuracy of one pixel size.

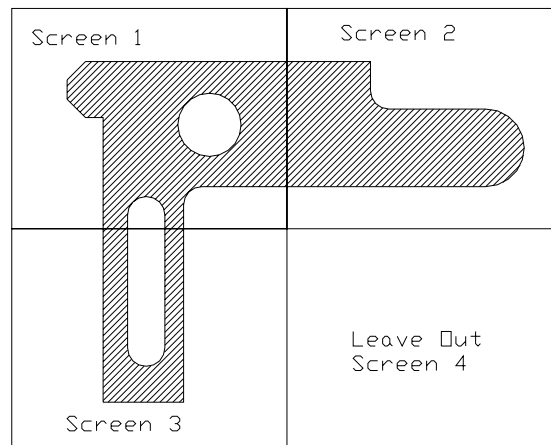


Figure 3.20: Example of screen stamping

The first step is determining how much travel each stage will need. Since this machine is mostly being designed to build micro-components the Z-stage was determined to need to have two inches of travel. For the X and Y stages a longer travel is needed in order to make multiple parts or larger parts. Taking into consideration possible upgrades of the machine in the future a travel of six inches was chosen for these stages.

The most important specifications that the stages need are the repeatability and the resolution. If a stage is sent to 10 mm then to some other number and back to 10 mm the error is known as the repeatability. This is important for creating accurate layer thicknesses over and over again. The resolution of a stage is the minimum value that the stage can move and detect. The desired

repeatability is $.5\text{ }\mu\text{m}$ and the desired resolution is $.1\text{ }\mu\text{m}$. These very precise stages are needed to create very uniform layer thicknesses and also to align the stamps to make larger parts.

Vat Setups

Due to the high viscosities of the photopolymers the vat and build platform design was crucial. The previous research all built very small components with small cross sections. Our machine was designed to build parts over a hundred times larger than any other MSL system. This required some serious thought to the regular style build process.

One possible vat design is a simple open top build. This system is what is on standard stereolithography machines. The part is lowered deep into the vat to receive a fresh coat of polymer. Then the build platform is raised to one layer thickness below the surface of the polymer. The viscous polymer bulges above the surface where the part is located. In a standard macroscale machine a blade is swept across the surface to level the polymer out. In these machines the typical layer thickness is from $50\text{ }\mu\text{m}$ to $150\text{ }\mu\text{m}$. We would like to make layers as thin as $10\text{ }\mu\text{m}$, this makes it very difficult to use a blade for sweeping. Not only would the blade need to be manufactured with this type of accuracy it would also need to be on a translation system that is extremely precise. The alternative to sweeping a blade is to wait for gravity to eventually level the polymer out. This is done on the other MSL machines from other institutions. However, the wait time would be exponentially larger than theirs because our parts are over one hundred times larger. This forced us to think of a different way to create our layer thickness.

A novel inverted build design was thought up. This process would involve building the part through a transparent piece of fused silica. This optical grade glass would be placed on the bottom of the vat and the platform would be lowered to the exact layer thickness desired. Traditional scanning methods used in stereolithography machines cannot use this process because the laser beam enters the glass at a different angle for each point in the plane. Due to refraction of the beam it would create large errors in the part. In the integral process the entire layer can be projected incident to the glass so that there is no refraction of the light. This is why this method of creating a layer has never been done before.

There are some difficulties that need to be overcome with building through a piece of glass. The most important is that it is necessary to have a non-stick surface on the inside of the vat. This could be accomplished by coating the glass in some material with a non-stick material such as Teflon®. However, this material must be transparent and not be damaged by the UV light. Also the interface between these materials must be strong so that they don't separate during the build process. After some research in this area multiple solutions were found to be available.

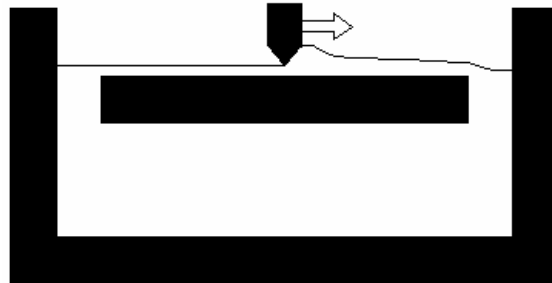


Figure 3.21: Sweeping method for making layer

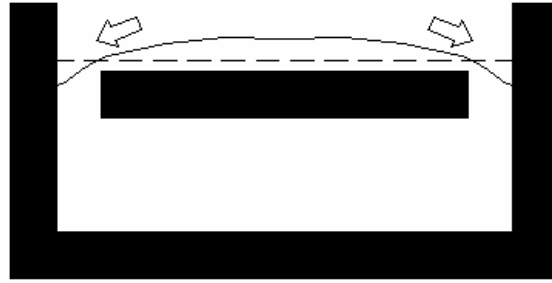


Figure 3.22: Waiting method for creating layer thickness

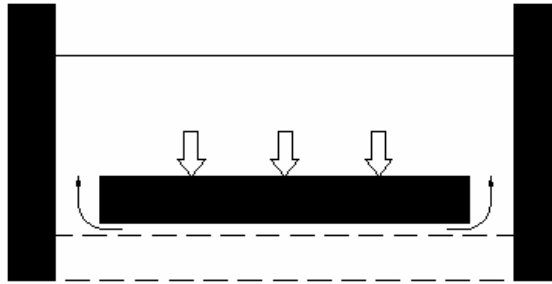


Figure 3.23: Inverted build layer creation

Vat Bottom Coating

The first option for coating the fused silica bottom is a new proprietary material from DuPont called Teflon® AF. This is an optically clear form of Teflon® that can be used to coat optics and other components. It was developed to be used in a contact lithography process in wafer fabrication. This material is very new and DuPont has only released the rights to distribute the product to one company called Random Technologies. A small amount was obtained directly from DuPont after the appropriate non-disclosure agreements had been signed. After the Teflon® AF arrived we tried to spin coat some samples. Due to our limited experience in spin coating and using the Teflon® AF we were unsuccessful in getting good samples. When we contacted Random Technologies we were told that they had the ability to make the part for us but they were too busy to stop production for one part. They gave us some alternatives that would probably work for our research needs.

There is another similar product to Teflon® AF that also has the low surface release energies needed for the inverted MSL process. It is also a fluoropolymer that is made by a company in Japan. This material is called Cytop™ and has a very low critical surface tension value of 19 dyne/cm. Also, it has a 95% transmittance to UV light at 355 nm. This material is an direct alternative to Teflon® AF and would be applied to the fused silica by spin, spray, or dip coating.

One alternative to the Teflon® AF is a Teflon® tape that is optically clear. These tapes are very inexpensive and can be bought on rolls that are 6 inches wide. The downside of the Teflon® tape is the installation onto the fused silica. It is difficult to attach a piece of tape that large with no air bubbles or particles underneath it. However, due to its cost the process can be tried over and over until successful. Also, as the tape wears down over time, or for any other reason it can simply be removed and replaced.

There are also various types of Teflon® films that may be able to be attached to the fused silica. These films have better surface quality than the tapes but need to be adhered to the fused silica.

This can be done with certain adhesives. Also it is possible to place the film on the glass and raise the temperature so the Teflon® melts and adheres to the glass when cooled. This can lead to fluctuations in the flatness of the optical piece. Overall a coating is the best solution even though it is the most costly.

Photopolymers

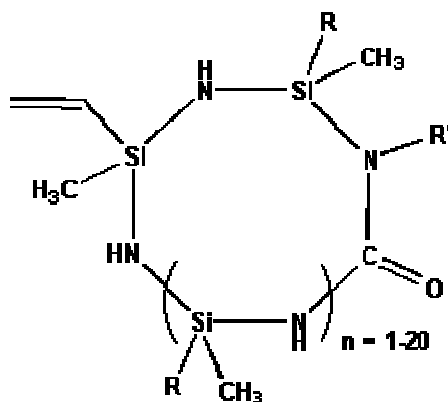
There are many different polymers that are readily available for traditional stereolithography. These polymers work equally as well in the microscale. The polymers vary in mechanical properties, but most are very viscous and require special solvents for cleaning the parts. After the photopolymer is exposed to the UV light it solidifies into its “green” state. After cleaning, post processing is done to these parts, usually combining heat and more UV light. Depending on the post processing the mechanical properties of the material can be altered slightly for the application.

For this research an older, less expensive commercial polymer was implemented. This polymer is DSM Somos 7120 which is used in commercial macro scale stereolithography machines. It is designed specifically for UV pulsed laser beams at 355 nm.

Table 3.1: DSM Somos Properties

Viscosity	Density	E _c	D _p	E ₅ 5 mil layer	E ₁₀ 10 mil layer	Tensile Strength	Max Temp.
~700 cps	~1.13g/ cm ³	8 mJ/cm ²	.123 mm	23 mJ/cm ²	64 mJ/cm ²	63 Mpa	97 °C

Another material that will be used in the MSL system is a polymer-derived-ceramic. Silicon carbon-nitride (SiCN) is a class of recently developed amorphous polymer-derived ceramics that remain mechanically stable and oxidation resistant at both high temperatures (exceeding 1500 °C) as well as in corrosive environments. SiCN ceramics can be conveniently synthesized from inorganic polymer. Commercially available Ceraset™ SN is a typical precursor for SiCN ceramics. Ceraset™ is a pale yellow liquid in oligomeric form, with density of about 0.96 gm/cm³.



R = H, CH=CH₂

Figure 3.24: The structure of Ceraset™ SN

The photoinitiator we used is IRGACURE 651 from Ciba Specialty Chemicals. It is a highly efficient photoinitiator which is used to begin the photopolymerization of chemical pre-polymers, e.g. unsaturated polyesters or acrylates, in combination with mono- or multifunctional monomers. These then cross-link and create a solid form of the polymer-derived-ceramic. As with other stereolithography materials there is post processing in an oven.

Computer Control System

The high accuracy of the MSL process requires a control system that is efficient and accurate. We used a computer with LabVIEW to run the processes of the machine. This program will make the build process automated which is necessary due to the length of time it will take to build one part, up to 24 hours. The computer program will control the stages, active layer mask, laser, external shutter, and other data acquisition functions. A schematic of the various components that are run through the PC can be seen in Figure 3.25.

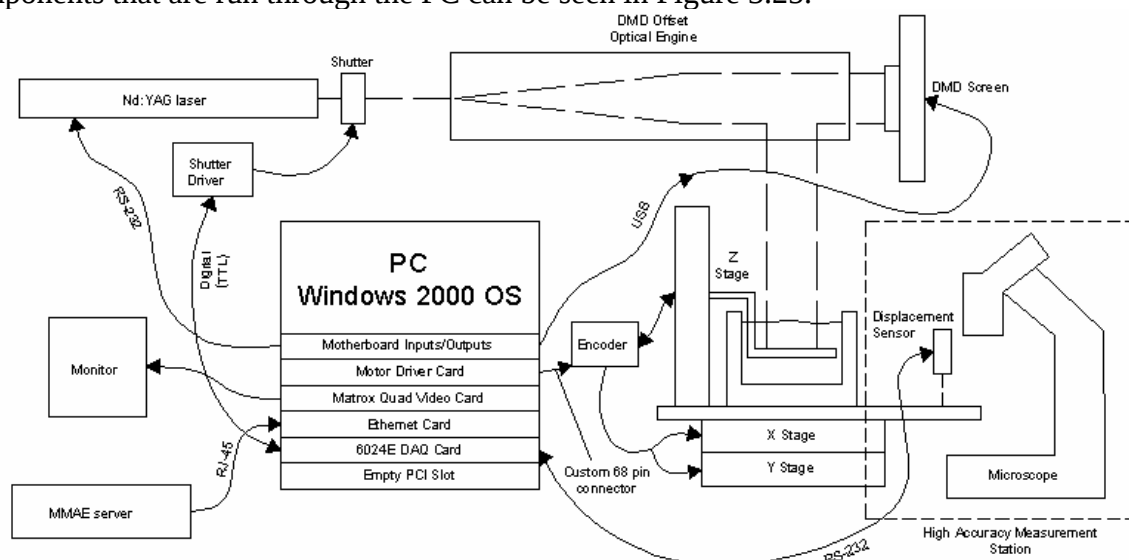


Figure 3.25: Computer communication for MSL system

Labview is a visual programming language. It allows many different functions to be done through one common interface. This allows communication to and from all the various components easily. Some hardware had Labview drivers already written and were available from the vendor. Other components, like the laser, had to have drivers written in order to control them. All the functions of the entire system could be run through a common visual interface (Figure 3.26).



Figure 3.26: MSL software main menu

These three function buttons brought up other visual interfaces to control different portions of the MSL machine. The laser initialization screen sets the laser up to the default configuration. It also allows the user to change any settings that may be necessary to improve the performance of the laser. The stage initialization will establish communication with the motor driver and the stages. Also, the stage settings can be adjusted to optimize the system.

The most important and useful VI is the manual test operations screen (Figure 3.27). Here each part of the MSL system can be controlled individually. This allows for experiments and adjustments to be made. Also, a number of preset experiments can be run from this screen by clicking the buttons on the right. Each of these automated experiments will run a series of commands to the various components. In the future of the project there will be a simple build button that will create an entire part automatically.

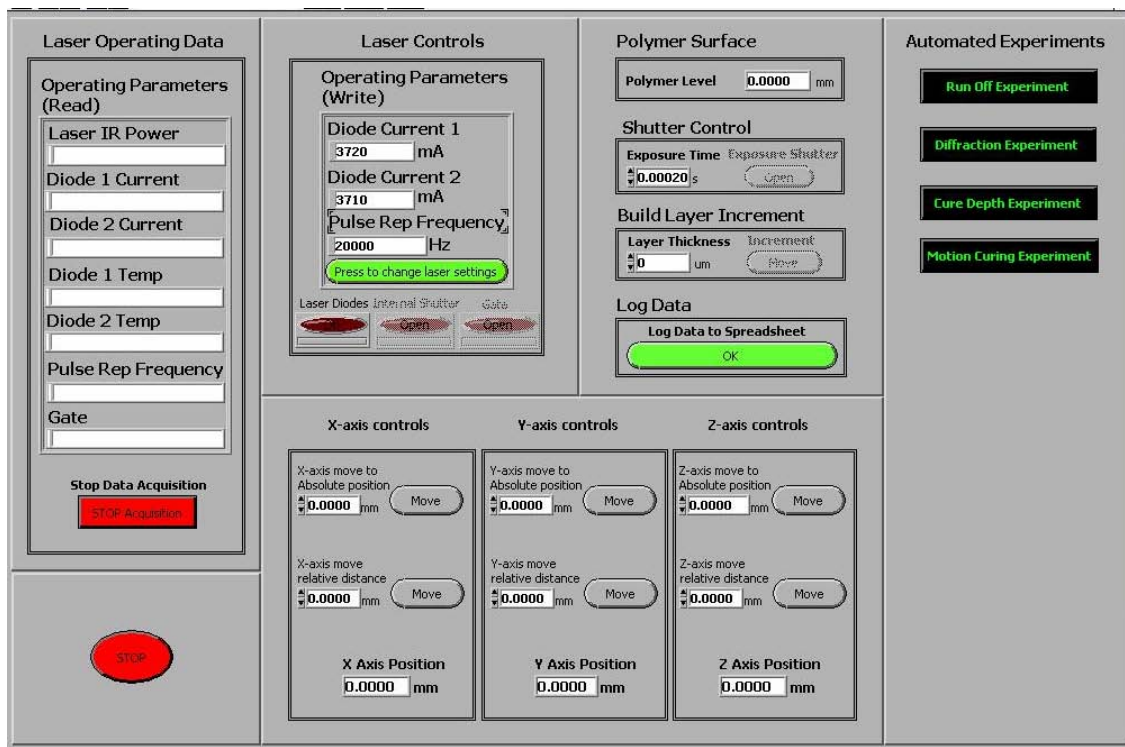


Figure 3.27: Manual test operations GUI

Final Design Setup

The first main component that was decided on was the Nd:YAG solid state UV laser made by Spectra Physics. This laser has a power of around 4 watts at a 20 kHz repetition rate. The wavelength is 355 nm with a TEM₀₀ spatial mode. The beam waist diameter is 466 μm with a full angle divergence of 1.1 mrad. The power supply for this laser runs off 110 VAC and a small external chiller provides cool water to the laser head.

Due to the UV wavelength all of the optical components are made from fused silica. The small lens is a plano-convex lens with a focal length of 10 mm and a 5 mm diameter. The small lens was mounted in a five axis holder from Newport. This allowed for the following adjustments:

± 1.6 mm in the vertical direction, ± 1.6 mm in the lateral direction, ± 3.2 mm in the axial direction, $\pm 5^\circ$ rotation about the vertical, and $\pm 5^\circ$ rotation about the lateral direction. The small lens mount was placed on a post and post holder from Melles Griot which was attached to the rail system.

The large lens is a bi-convex lens with a focal length of 1/2 and a 50 mm diameter. The large lens was mounted in a two axis holder from Siskiyou Design. This allowed for .187 inches of movement in the vertical and lateral directions. All axial adjustments for separating the lenses were done on the small lens. A custom mounting block was made to mount the large lens holder to the rail slide. This rail was obtained from Newport and has graduations of one millimeter. An overall rail length of 48 inches was used for the various optical component mounts.

The SLM chosen was the DMD due to its ability to be used in the UV range. The DMD is a XGA device that is 1024x768 pixels. The pixels are $13.7 \mu\text{m}$ with a tilt of $\pm 12^\circ$ and have a gap of $1.1 \mu\text{m}$. This means a fill factor of above 85%. The maximum illumination energy is rated at 10 watts per square centimeter. The DMD is mounted on a circuit board made by Productivity Systems. This board is controlled through a USB connection to the PC. The DMD can be driven at frame rates up to 100 Hz.

The offset optics light engine was made entirely by Brilliant Technologies. This device is protected under an agreement with the manufacturer. Details of the optical system cannot be discussed. The beam enters this system at a height of 5.65 inches off the optical table and the one to one image from the DMD is projected 3.36 inches away from the optical engine.

The three translation stages were made by Newport. The z-stage has a travel of 50 mm and the X and Y stages both have a 150 mm of travel. The resolution of the stages is $.1 \mu\text{m}$. The Z stage is also mounted on a custom angle bracket provided by Newport. These stages are controlled through the motor controller/encoder via a custom 68 pin connection. A custom PCI card is used to interface this motion controller with the computer.

An inverted vat design was built. A shallow plexiglass top was made with two inch high walls. Also, a square inset was machined for the fused silica to sit. A bottom half, made from aluminum, was also machined with the same inset. A gasket material was used between the fused silica glass and the top and bottom of the vat to make a seal. A five inch square fused silica window was used for the vat bottom.

The vat bottom was coated with Teflon® FEP. This coating was applied to the fused silica using a spray process. This procedure was performed by Micro-Surface Corporation. The Teflon® AF coating is $1/2 \mu\text{m}$ thick and optically clear. Due to the low surface energies the parts do not adhere to the bottom of the vat. The major polymer that will be used in the system is the DSM Somos 7120. The Ceraset will only be used for simple test parts. The final computer program was entirely written in Labview software and effectively combines all the various components of the project.

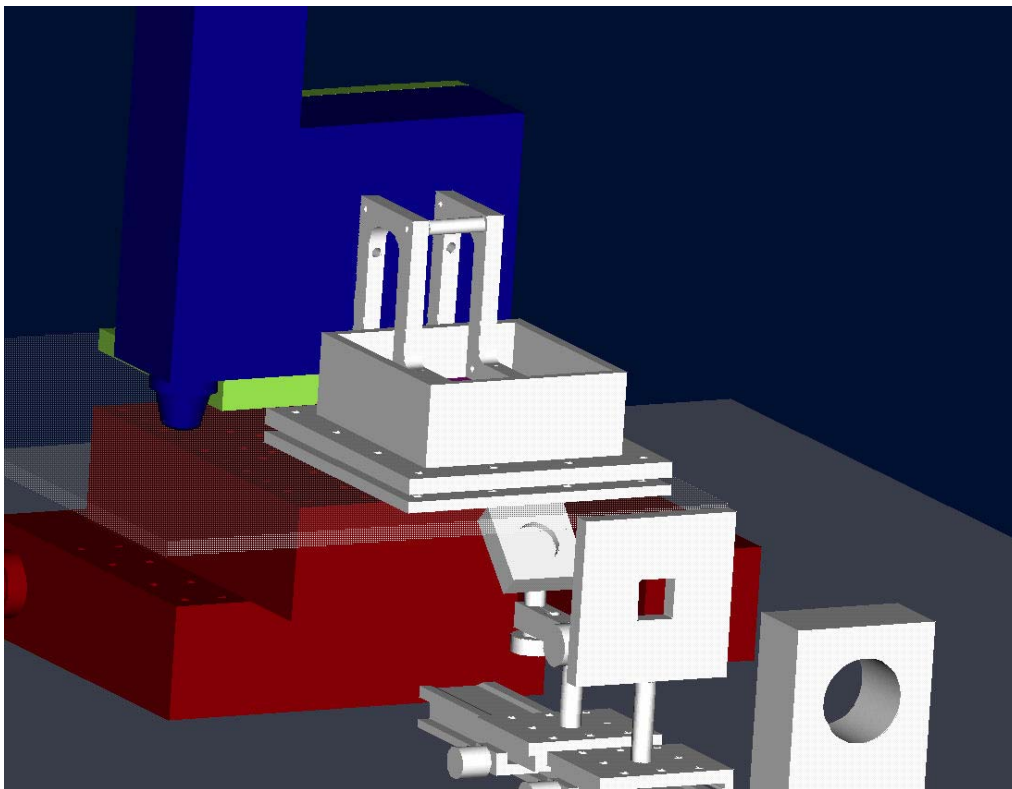


Figure 3.28: ProE Drawing of Final Design

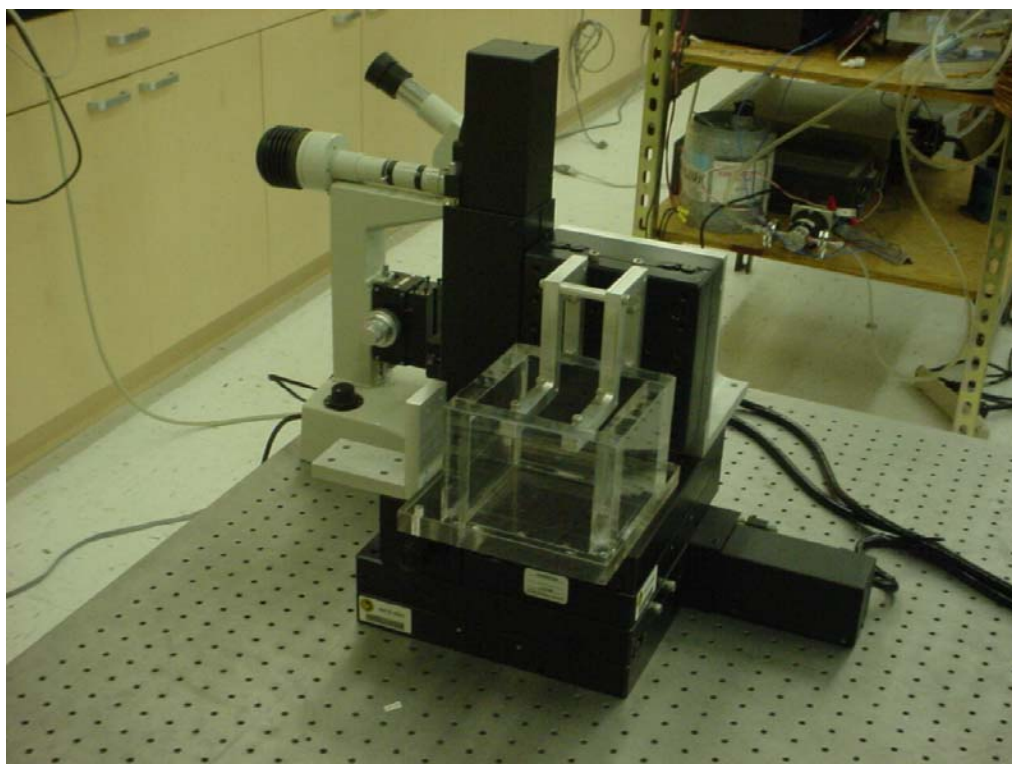


Figure 3.29: Photo of Final Design

3.4 RESULTS

Random Shrinkage

One of the major problems with the stereolithography process is that there is a phase change of the liquid photopolymer to a solid. This happens slightly different each time. When the parts are removed from the liquid and post processed they shrink. Due to the phase change there is some randomness to how much the parts may shrink. We call this the random shrinkage. On macro scale builds this can be ignored, however in the microscale it can become a manufacturing problem.

Before we even spent the time or money to build a MSL system it was first necessary to determine the random shrinkage of the ceraset material. If the randomness caused more error than the accuracy we required then another method would need to be used. For our random shrinkage tests samples were made using a mold. The liquid ceraset was poured into the mold and heated until it solidified. Then the samples were given two micro-indentations. The distance between these points was measured using high resolution translation stages.

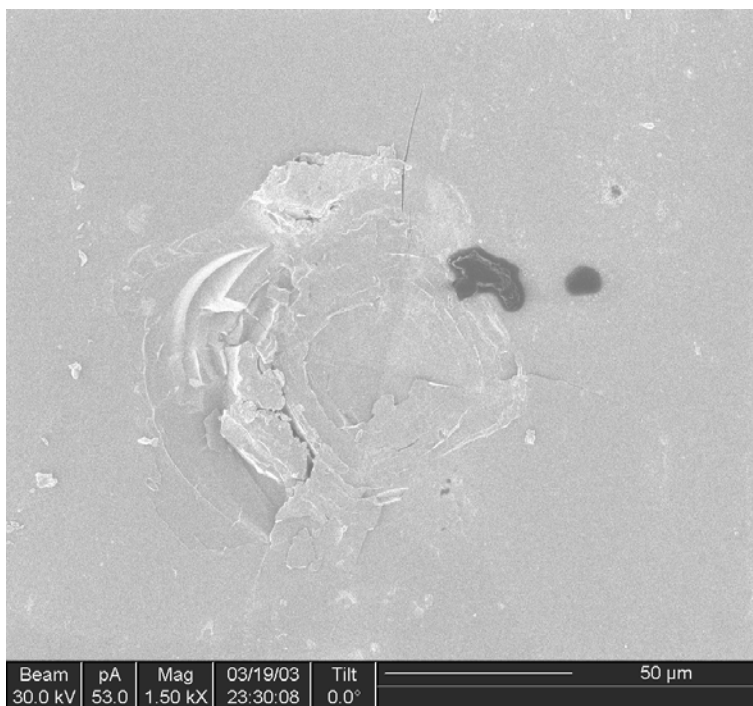


Figure 3.30: SEM Photo of Indentation Mark

The parts were then pyrolyzed (converted to full ceramic) by heating them to a high temperature. The samples were then taken and the distance between the points was measured again. The ratio of the two measurements gave the percent shrinkage. Ten samples were made using this same process and a statistical analysis was performed on the shrinkage numbers.

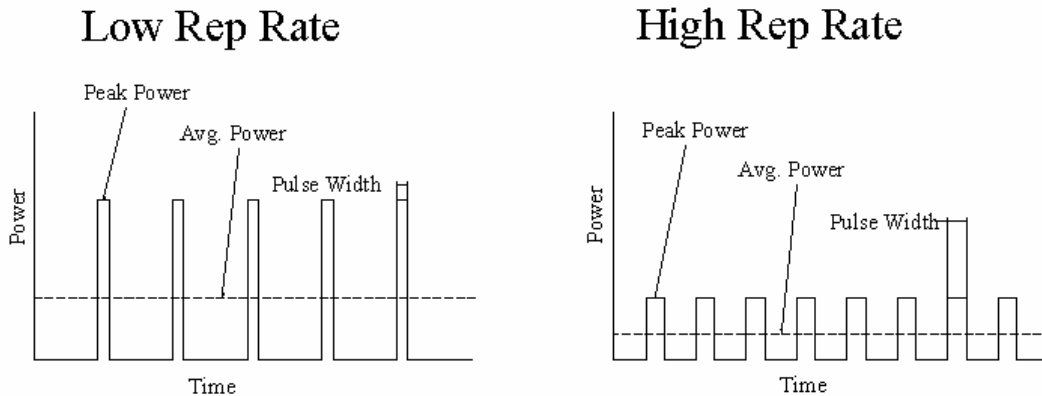
Table 3.2: Random shrinkage numbers

Sample #	Distance Before Pyrolysis	Distance After Pyrolysis	Standard Deviation	Shrinkage
1	3.3145	2.457043	0.002016377	25.87%
2	2.4548	1.827009	0.017719138	25.57%
3	2.805	2.075755	0.008849075	26%
4	3.1352	2.315948	0.072958572	26.13%
5	3.1358	2.326047	0.007959479	25.82%
6	3.1974	2.39437	0.007027	25.12%
7	2.9541	2.197384	0.000175735	25.62%
8	3.0618	2.2527	0.0075251	26.46%
9	2.5091	1.857329	0.011874	25.98%
10	3.2178	2.363596	0.009912	26.55%

We calculated the sample mean shrinkage is 25.912%, the standard deviation is 0.0042255. For $v = n-1 = 9$, the $t_{0.025,9} = 2.262$, calculating the two-side confidence limits, $t_{0.025,9} \times S/n^{1/2} = 0.0030225$. Hence, the shrinkage is $25.912\% \pm 0.30225\%$ with a confidence of 95%. The random portion of the shrinkage is within the values needed for an MSL system.

Laser Power Curve

The solid state laser can alter its output power by varying the repetition rate. Figure 3.31 shows how the q-switch timing can alter the average output power of the laser. When the repetition rate is increased then the energy in the laser cavity does not have enough time to get to a fully excited state. Therefore, if the repetition rate is increased the peak power, and therefore average power, will be decreased. This is valuable information for a MSL system because the ability to finely adjust the power can help determine necessary exposure times.

**Figure 3.31:** Low vs. High repetition rate

It was necessary to calibrate the laser power to the repetition rate so that this information could be used for various exposure time experiments. An external power meter was set up in the path of the laser beam and the repetition rate was changed to different values. The corresponding values from the power meter were logged and a simple calibration curve was created. This

process would need to be done every so often because the laser power slightly changes over time. Also, when the power drops beyond a certain specification the crystals in the tripler module must be translated.

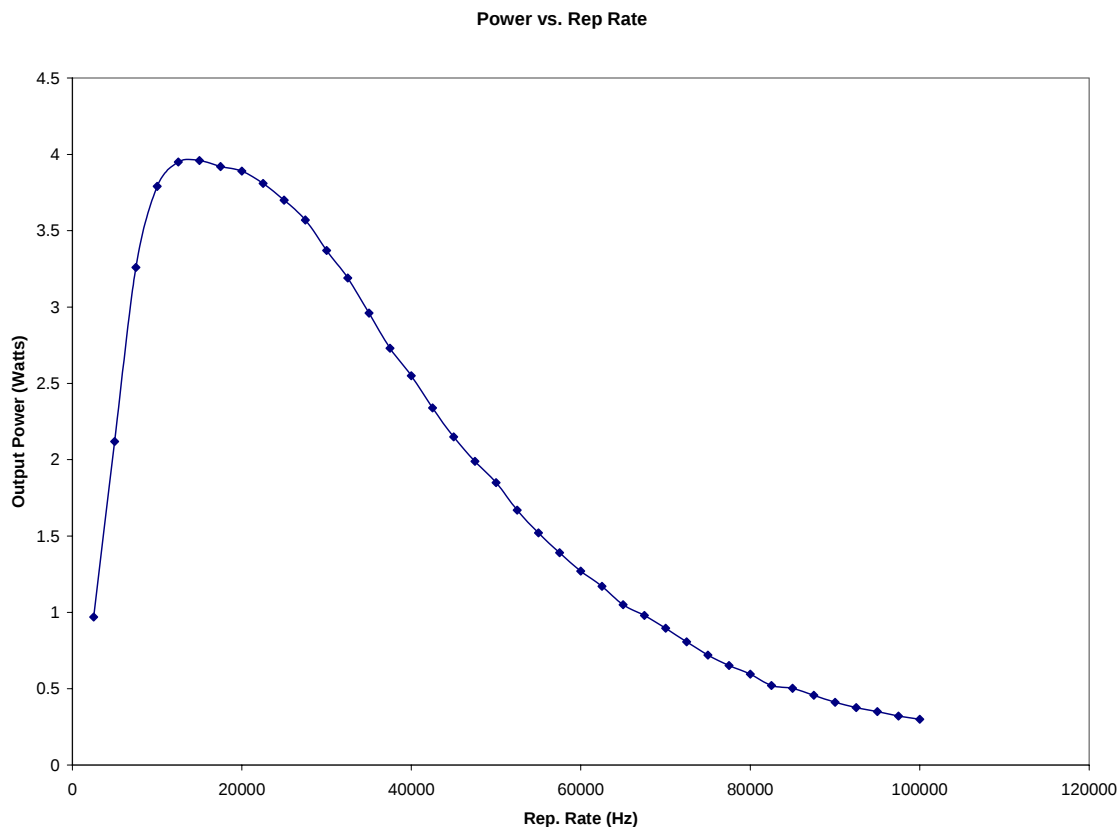


Figure 3.32: Laser power as a function of Repetition Rate

Another important piece of information that was gathered from the external power meter was a calibration between the internal power meter and the actual power. On this particular laser there are two internal power meters, one for the IR region and one in the tripler module that measures the UV power. With this laser the UV internal power meter did not register the correct values. The actual values will be needed to adjust the exposure time according to the current laser output. Another calibration was performed to determine if the internal power meter was related to the actual power through some predictable relationship. When the external power and internal power readings were graphed it was clear that there was a linear relationship (Figure 2.33). This makes it possible to read the power from the laser head and multiply by a constant. This makes an external power meter not necessary for the build process.

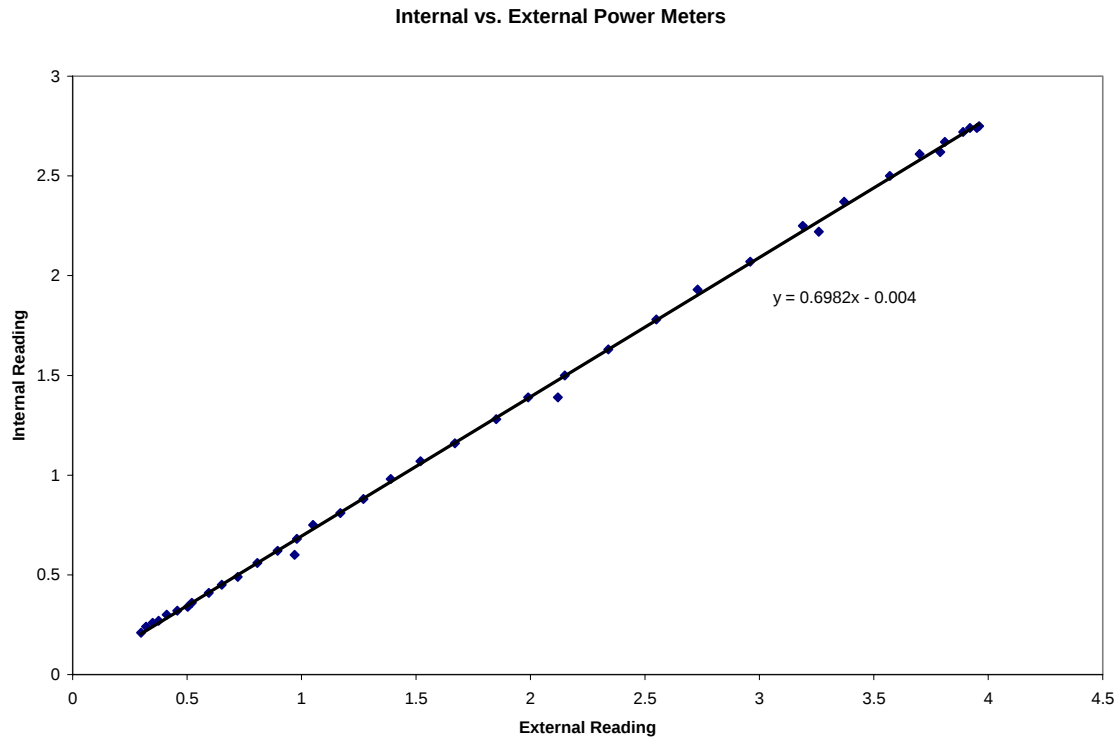


Figure 3.33: Relationship between internal and external power meters

Ceraset Cure Depths

In order for the MSL system to successfully create parts from the ceraset the curing properties must be known. Unlike the 7120 these properties are not published or sent with the polymer. The cure depth value can be altered by the concentration of photoinitiator in the polymer. In order to determine the exposure time needed to make a given layer thickness an experiment was needed to find the cure depths.

The cure depth experiment was performed while the laser was at full power. A small cup was filled with ceraset and a thin glass slide was slid on top of the cup. The laser beam was passed through a .25 inch diameter aperture and exposed the ceraset through the glass slide. After exposure the sample was removed and washed with acetone and then left overnight to fully cure. Three samples were made in this manner for each exposure time.

A measurement system was set up using the motorized translation stages and a distance sensor that uses a small laser beam and CCD to determine very small distances. We called this sensor the displacement sensor. For most applications it can be used to measure very small displacements of a solid or liquid surface. The measurement range of the sensor is 3.175mm and the resolution is around 1 μm . Combining the displacement sensor with our translation stages allowed for a high resolution measuring system.

The samples were placed on translation stages and incremented by 1mm. Each increment a measurement was taken by the displacement sensor and the height at that point was logged. This

created a data sheet that would be able to create a surface plot of each sample. By numerically evaluating these values a cure depth would be determined for the different exposure times.

Although this process should work there were a few significant problems with this procedure. After all the samples were made and the heights evaluated it was found that the standard deviation of the data points was on the same order of magnitude as the cure depth values. We were not sure if the data was bad or the samples really did vary that much so another experiment was done to determine if our experimental method was sound.

The Somos 7120 is sent with data values the readily tell what the cure depth should be for a given input energy. It was decided to run a control group to test whether the cure depth measuring was an error in the measurements. The same experiment as the ceraset was performed using the 7120. These cure depths could then be compared to the analytical values to determine if the measurement process was sound. Unfortunately even when these samples were created and measured similar results were found with the standard deviations. This led to the conclusion that the optical properties of the photopolymers were probably the reason that the displacement sensor was not giving the correct values. Occasionally the sensor would output the correct value but the data was so random that an exact cure depth could not be determined.

The data did lead to a chart that is similar to what is expected from the cure depth. A logarithmic graph can be seen through the data points of the ceraset cure depths. However due to the errors in the measurement the bias of these values cannot be determined. In other words the shape of the graph will remain the same but it will be shifted by some unknown amount. See Figure 3.34.

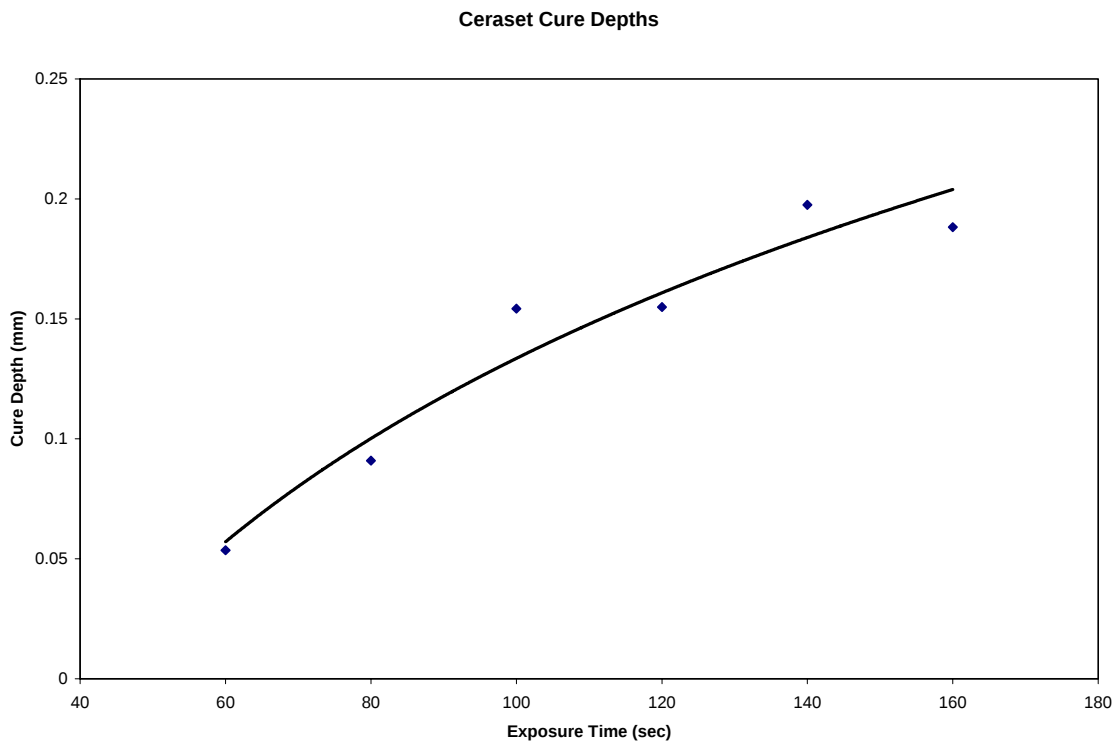


Figure 3.34: Ceraset Cure Depths

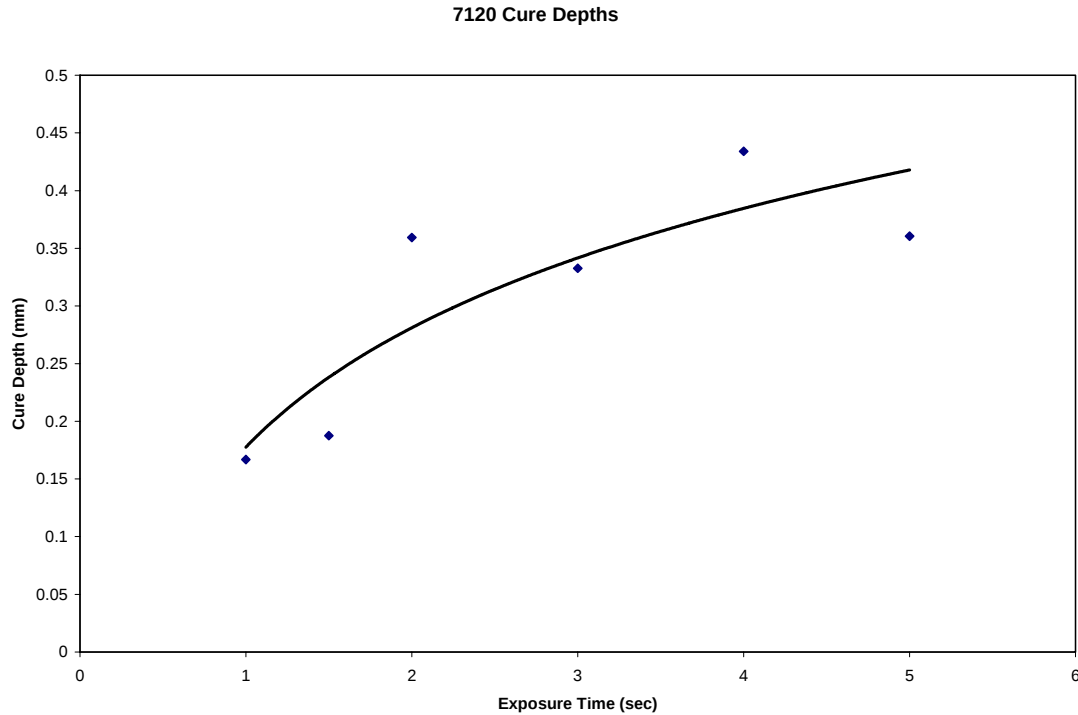


Figure 3.35: Commercial 7120 Cure Depths

Also, during these experiments it was found that the ceraset did not have good properties after exposure to the UV light. The ceraset did cross-link to an extent; however it took a very long time to create any layer. While the commercial polymer takes less than one second the ceraset took over two minutes. Perhaps the most troubling piece of information was that the ceraset could be cleaned by stirring the sample in acetone, however the ceraset did not fully cure into a solid form. After exposure it was in more of a gel state. This is a severe problem if the ceraset is to be used in an inverted build design. A lot more work is needed researching the properties of the ceraset for it to be used in an MSL system.

Single Layer Builds

The simplest part to make is a single layer using only one pattern of light. These parts can be useful and often simple structures only require a single layer. To make our first parts a mask needed to be made that had a part. This mask could be made in traditional mask making fashion on a piece of soda lime glass. A photomask vendor was contacted and a “static” mask was made with nine sections. This static mask would be used throughout the experiment processes in making all of the sample parts. In the future of this project the static mask will be replaced by the dynamic mask, the SLM.

The static mask is four inches square with a three inch square of usable space. The critical dimension of the mask (the smallest feature) is 10 μm . The mask contains nine separate sections that are each 1 inch by 1 inch. Four of these sections hold the four layers needed to build a spray nozzle atomizer. The other sections contain; support structure, a heat exchanger, gears, simulated pixel arrays, and other sample parts.

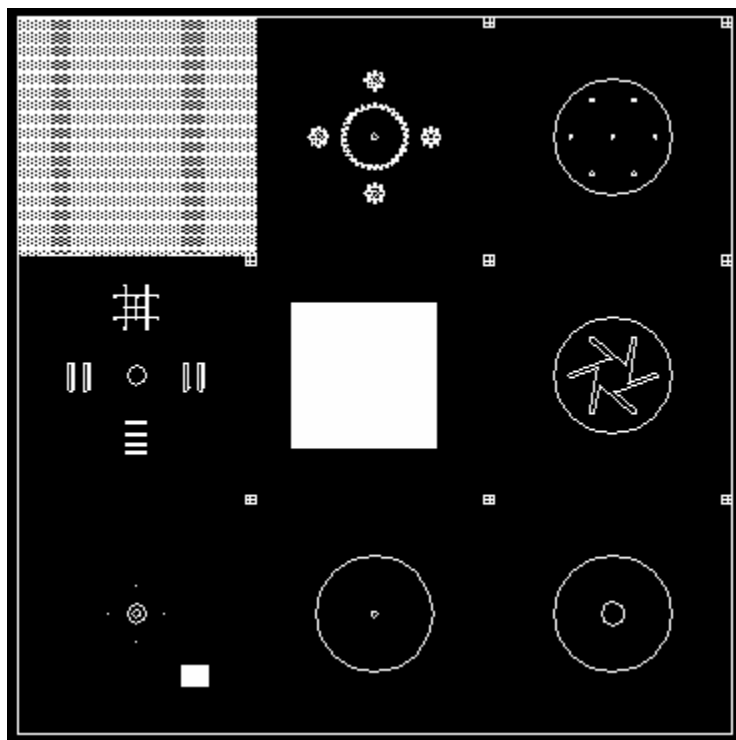


Figure 3.36: Static mask design

The one layer parts were made in a similar fashion to the cure depth experiments. The laser beam was projected through the mask and carried that image through a glass slide that sat on top of a small container of polymer. The sample parts were made from ceraset.

The ceraset single layer parts are the only components that can currently be built directly by the UV light. Attempts at building multiple layer parts with ceraset were unsuccessful due to the gel-like state of the ceraset after laser exposure. The single layer parts (Figure 3.37 & 3.38) were exposed for 120 seconds and then rinsed in acetone. They were then left to fully dry overnight.

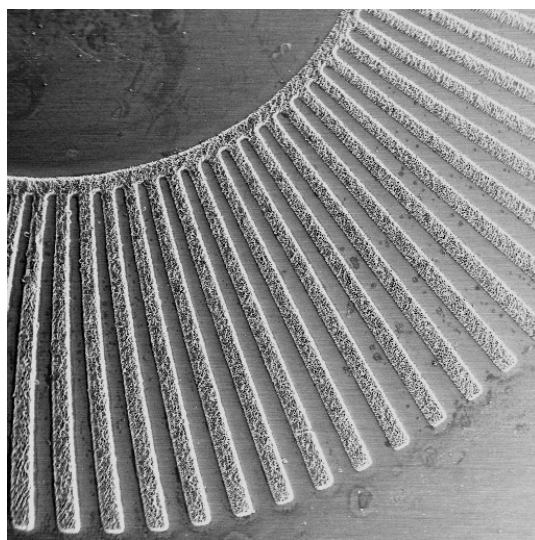


Figure 3.37: Ceraset single layer sample part

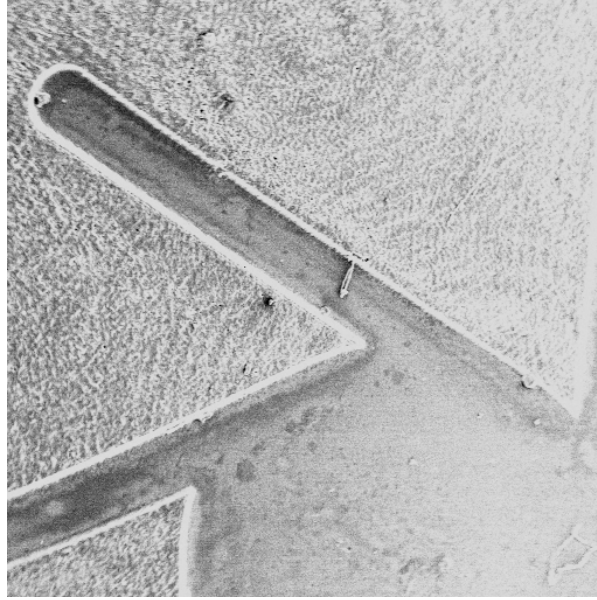


Figure 3.38: Ceraset single layer from a spray nozzle

The commercial polymer parts were also made in the same way as the ceraset components. After trying a few different exposure times it was clear that these polymers are very sensitive to the laser power and the exposure time. This leads to the concept that the actual laser power after all the losses through the optical system must need to be known. This value is difficult to measure due to the limited space between the translation stages and the laser path. An experimental method was determined to be a better way to get the precise exposure time necessary for a given layer thickness. This method is discussed in detail in the inverted build section.

Multiple Layer Parts (Standard Build)

After successful attempts at building single layer parts it was logical to try to make multiple layer parts. We began this process in a standard build style using the gravity leveling technique to create the layer thickness. The commercial polymer was placed in the vat and the build platform was lowered into the polymer. The platform was then manually raised until just below the polymer surface. The exact thickness of the first layer was not known which made choosing an exposure time to be difficult. This led to some samples in which the first layer did not fully adhere to the platform surface. This is a major obstacle that will not be necessary in the inverted build design.

After successfully determining the first layer exposure time by trial and error a few sample parts were fabricated. These parts had layer thicknesses of 250 μm . Even at these very thick layers a wait time of 480 seconds was needed between layers for the polymer to become level due to gravity. The static mask was placed just above the vat and the laser exposed for 1 second per layer. A total of 6 layers were built and the results were examined under a scanning electron microscope.

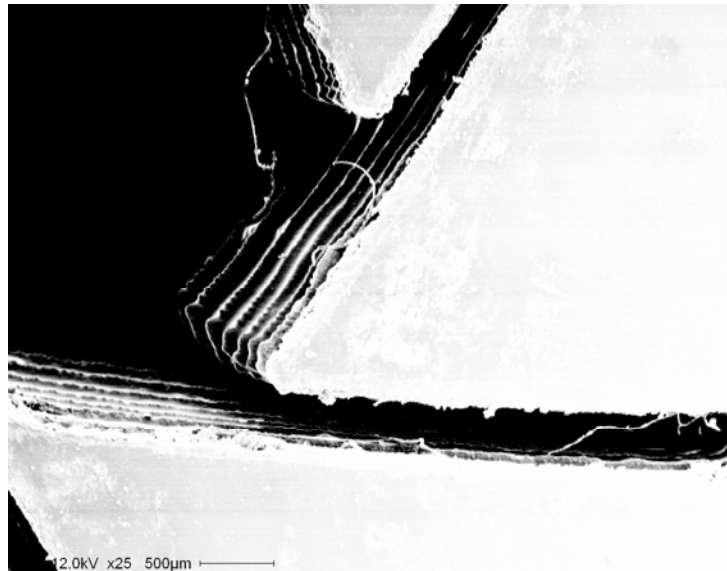


Figure 3.39: Multiple layer part standard build

The resulting parts were very promising. The vertical wall of the part appears very straight. Also, the layers appear uniform across the entire part surface. The parts were cleaned in a solvent called tripropylene glycol monomethyl ether also known as TPM. These parts were not post processed because we weren't interested in improving their mechanical properties. After these successful multi-layer parts we knew that the project would be a success.

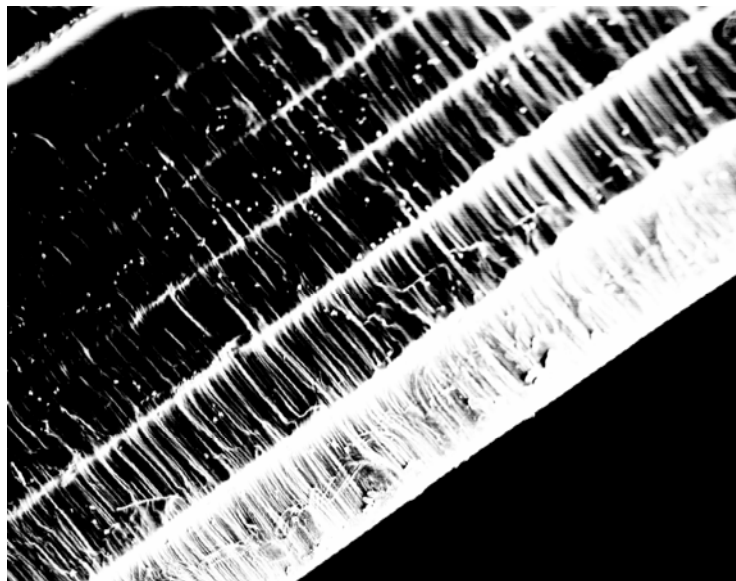


Figure 3.40: Cross section view of layers using standard build

Only single mask parts were made in this fashion. It was found, after doing some experiments, that the mask alignment was very difficult with the limited pieces of hardware that we had available. We attempted to align the mask using two manual translation stages and a manual rotation stage. We used an optical microscope to align the marks on the static mask. This

proved to be very difficult and it was determined that the time spent on mask alignment would be better spent on getting the system ready for the inverted build design.

Adhesion Forces

It is important to determine the separation forces between the current layer and the Teflon® coated window. These forces may limit the MSL system from building certain geometries. For instance, a cantilever beam will only be allowed to have a certain length before the separation forces will break the beam off of the rest of the part.

From this data a calculation was done to determine the maximum size of a cantilever beam. This is used to determine how many pixels can be placed in a row before they will break from the adhesion force. The distributed load is found from the critical surface energy of the Teflon®, 20 dyne/cm². This was then used to find the maximum stress that is seen by the cantilever beam using equation.

$$\sigma_{\max} = \frac{M \times c}{I} \quad (3.2)$$

This maximum stress is compared to the tensile strength for the DSM Somos 7120 with a safety factor of 3. For a beam one pixel wide it was found that the maximum length is 17.8 mm. This means that there should not be a geometry limitation when building parts.

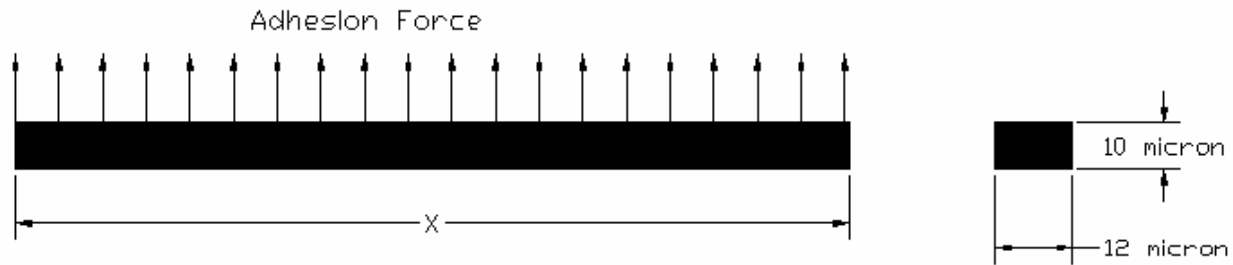


Figure 3.41: Adhesion force limitation on geometry

Multiple Layer Parts (Inverted Build)

The inverted build process is what makes this MSL machine different from any other at other educational institutions. For comparison with the standard build parts the same static mask was used to make the various test parts. In the inverted build the first layer exposure was determined by experimentation. The platform was lowered to a given position that was designated as the permanent start position for the build. The laser was exposed for various lengths of time until the samples did not adhere to the platform. The initial exposure time and start position could then be coded into the software.

After the first layer was built it was now possible to make layers any thickness that was desired. Again these exposure times depend on the laser power at the polymer level. This exact value is difficult to measure so the same experiment was done with subsequent layers as we did with the first layer. After this exposure time was determined for the given layer thickness it was possible to build some sample parts.

These parts could be built much faster than the standard build because the wait time between layers is less than one second. After a layer is built the platform is raised back into the polymer and then returned exactly one layer thickness away from the vat bottom. This made it possible for us to build our 10 layer part in less than 1 minute.

Picture of inverted build

Once again these parts were examined using an SEM and optical microscopes. As you can see from comparing Figures 3.37. The inverted build components are more defined and have cleaner and straighter edges. Due to time constraints sample parts were not made using the DMD and light engine.

Manufacturing Error

The most important part of the MSL process is that the desired accuracy is met. For the system to be considered successful the manufacturing error needs to be within the tolerance that is desired. After some parts were made using the static mask they needed to be compared to the designed size of the part. One section of the spray nozzle was chosen for the manufacturing error measurements.

First it is important to note the error not due to the MSL procedure. In this case the dimension for the channel length of the spray nozzle in the computer is 3018 μm . When the static mask is made it is not made to the exact size that it is drawn in the computer. The critical dimension of the static mask is 10 μm , however this dimension on the actual mask is 9.858 μm . The entire mask is 1.415% smaller than the designed value. This makes the channel of the spray nozzle on the static mask 2975.3 μm .

The DSM somos 7120 also has an average shrinkage number. The part is expected to shrink roughly .05%. This will make the expected channel length of the spray nozzle to be 2973.8 μm . When designing a part to be made using the DMD the scale factor of 1.0005 will be applied to the model before slicing the CAD model into layers. This will create bitmaps that should create the part as it is designed in the computer.

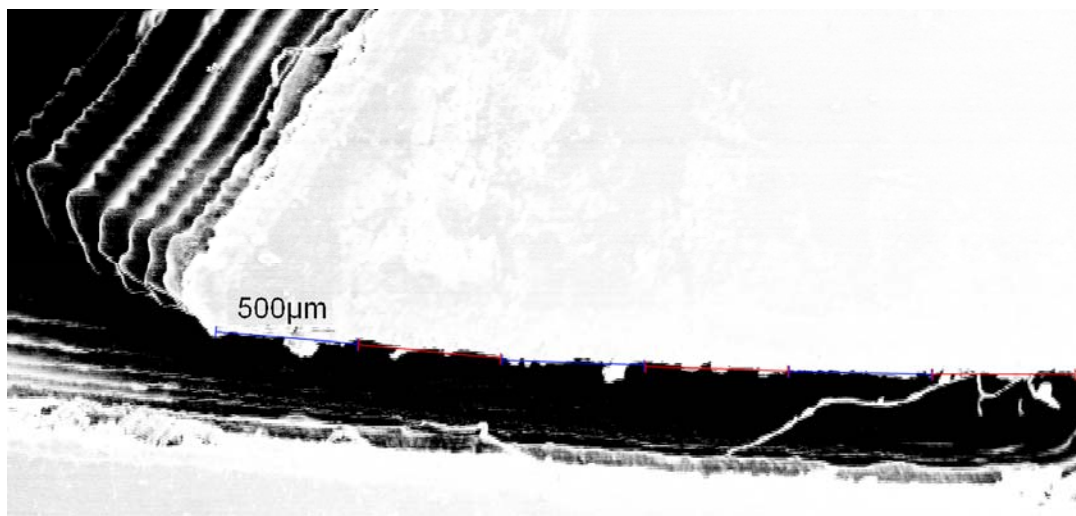


Figure 3.42: Measurement of spray nozzle channel

After the spray nozzle layer was created it was measured using an optical microscope and a gage. It was found that the part measured roughly 3000 μm . This makes the overall manufacturing error of less than 1%. For a part the size of the entire DMD this would mean an absolute error of only 135 μm . There is also human error in these measurements that may amplify the actual manufacturing error. The 1% value easily fits into the tolerances required for our applications.



Figure 3.43: Nozzle section compared to standard ruler

3.5 CONCLUSION

There is a high demand for a manufacturing method for high aspect ratio micro- components. The MSL system that was designed and built at UCF addresses these needs and is the first step in making this a valuable method for MEMS. With more work and optimization the MSL method could be ready for commercial use in a relatively short period of time. As more research in this area is completed it is clear that the cost of this technology will decrease.

Over the past few years the microdisplay industry has greatly increased as part of the consumer electronics market. High definition television and internet based forms of entertainment place additional performance requirements on displays that are best provided by spatial light modulators such as the digital micromirror device (Kunzman, 2000). As these technologies expand out into the market it will be vital for the cost to be decreased in order to reach the middle income bracket. The \$3,000 price point is believed to be the point where middle-income consumers begin to look at front projection systems seriously (Chinnock, 2002). As the prices drop and the resolutions increase a desktop MSL system is more and more likely to become a reality due to the decrease costs of the SLM's.

For now more experiments need to be performed to optimize the MSL system that has been developed at UCF. First thing that should be done is to implement the DMD with the light engine into the system. This will require some alteration to the laser beam and a device may

need to be added to remove the coherency of the laser beam. This can be done with a rotating holographic diffuser. This will allow the light to properly and uniformly illuminate the surface of the DMD.

Also, it would be desirable to have different exposure times for different layer thicknesses. This would optimize the time needed for the build. The computer would recognize the cross section and determine the exposure time and layer thickness that would optimize the build process. This would allow for parts to be built in less time and therefore less money.

It may also be possible to replace the expensive laser with a less expensive light source such as a high power UV lamp. For this to be possible the divergence and diffraction issues would need to be addressed to determine if the resolution from the SLM would be lost. If a UV lamp solution was found it may be possible to build the entire optical system for around the same price as a consumer projection system. As one can see there are a lot of improvements that can be done to the current design to make the MSL system less expensive and perform better.

Perhaps the most interesting things that could be done with the MSL system that we currently have built is a new technique that has not been tried ever before. It is a mesh of the raster scanning method and the integral method. The major advantages of the scanning method are to be able to build very large parts, however the hatching of the areas can take a long time. The main advantage of the integral method is the improved resolution and the speed of the build process, however only small parts can be made. If these two methods were combined into one machine the results could be fantastic.

The scanning-integral method that should be explored involves using the X-Y stage setup to move the vat. The CAD model is sliced into bitmap form in which the bitmap resolution is much higher than the resolution of the SLM. The image is displayed on the DMD and “panned” in the computer across the DMD. At the same time the X-Y stage system is also being moved to mirror the movement of the image on the DMD. If the timing was extremely precise then a single pixel would be illuminated as it traveled across the screen and would be fully cured at the far side of the screen. This would basically create a one inch square laser beam that would perform the same as the scanning method. The laser beam would just change shape via the SLM to only illuminate the pixels that were wanted.

This novel approach has never been attempted and would entail very strict timing requirements. However, if successful it would lead to the possibility of making very large parts on the macro scale with the resolution and accuracy of an MSL system. Currently we have the ability to begin trying research in this area with our current hardware. This may be an important method to examine the feasibility of for the future of MSL.

An extensive literature was completed that lead to the success of this MSL system. Based on that work a novel inverted MSL system has been designed and built at UCF. Over the next months and years there will be many papers published regarding our work. And hopefully this project’s success will help lead to a MSL system to become available for commercial use.

4. A FLUID MANAGEMENT SYSTEM FOR A MULTIPLE NOZZLE ARRAY SPRAY COOLER

4.1 INTRODUCTION

The advantages of two-phase cooling systems make them ideal for use in high heat flux cooling for aerospace and space based application (Baysinger, 2004). The advantages like compactness, light overall weights, and high heat flux dissipation lend themselves very well to particular application such as high energy lasers and high power electronics systems.

Spray cooling and micro-channel cooling have presented themselves as the best solutions; however, spray cooling has advantages over micro-channels coolers such as isothermal surfaces temperatures, lower working fluid flow rates and smaller system size. It has been demonstrated (Chow, 1997) that spray cooling can remove heat fluxes as high as 1000 W/cm^2 for single spray nozzle over an area of less than 2 cm^2 .

However, many applications require cooled areas on the order of tens of square centimeters. Spray cooling over larger areas (20 cm^2) has been tested and it was found that flooding of the cooled surface occurs due to the lack of excess liquid drainage (Lin, 2004). This flooding decrease the heat flux by 30% to 34%. Therefore a fluid management system is needed to minimize the degradation in heat removal capabilities caused by flooding.

Design Problem

The goal of this research is to design a scaleable pressure atomized spray cooler capable of cooling large areas, greater than 16.8 cm^2 . The inherent problem with spray cooling large areas is the flooding of the cooled surface and the creation of unwanted temperature gradients across the cooled surfaces.

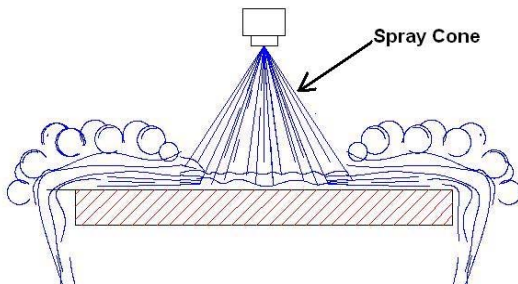


Figure 4.1: Single spray: horizontal

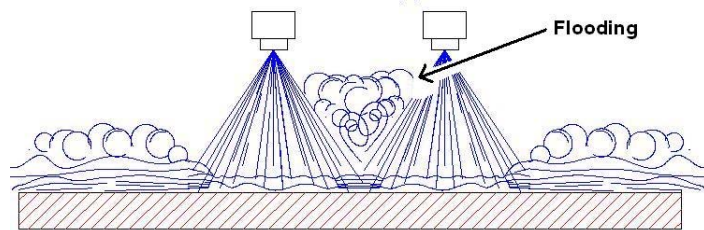


Figure 4.2: Multiple spray interaction: horizontal

The main driving mechanism in spray cooling is the thin film, which is created just above the cooled surface (Chen, 2002). In a single pressure atomized spray nozzle the liquid run off which

is about 90% of the input is pushed away from the cooled surface via the momentum of the spray (Figure 4.1), the other 10% is evaporated away (Chen 1995).

However, in large area spray cooling, multiple spray nozzles are used to cool the heated surface and their spray cones and run off liquid interact with each other (Chow, 1997). If the run off liquid is not removed from the cooled surface it will build up or flood the cooled surface thereby destroying the thin film need for effective spray cooling. Figure 4.2 shows a build up of liquid between the two spray cones. This would move the liquid/vapor phase change to a pool boiling regime which is unwanted and has low heat flux associated with it than spray cooling. The design problem presented in this paper is how can one control the flooding effect of multiple spray nozzles and maximize the overall heat flux.

Design & Solution

In view of our design goals we decided to construct and test 4 by 4 array of spray nozzle and create a system that could manage the excess fluid trapped between the adjacent spray cones. This system which we call the fluid management system utilizes an array of siphons by which the excess liquid is removed.

In a 2 by 2 array of spray nozzle it was observed that some of the excess fluid was being trapped between the spray cones, this can be seen in the top view of Figure 4.3 and in the side view in Figure 4.2. Additionally the spray cones interacted with each other; this can be seen in the darker areas of the spray cones in Figure 4.3. The spray cone interaction forced the excess liquid away along the cooled surface in the direction shown in Figure 4.3.

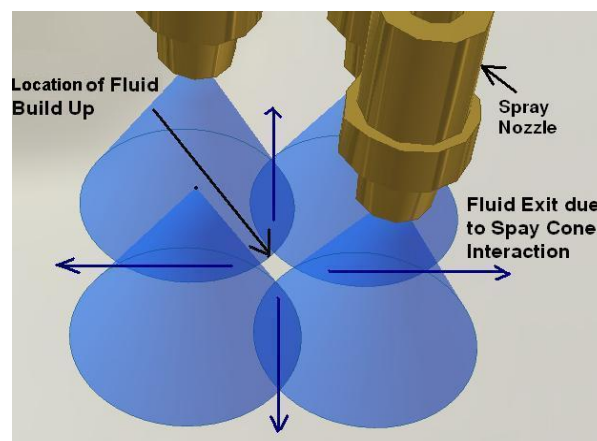


Figure 4.3: Fluid build up points

In order to gain control over the build up of the excess fluid and its exit direction two design features were implemented. First being the placement of siphon tubes at the flooding points and the second being an implementation of a grooved in the cooled copper plate. The grooves which are aligned in a grid pattern (Figure 4.6) were used to directed excess liquid created by the spray

cones interaction to the nearest by suction point or siphon. Both of these design features can be seen in Figure 4.4.

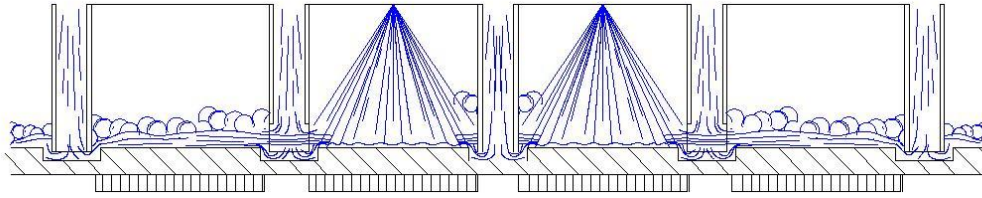


Figure 4.4: Fluid flow observed from the side of the spray cooler with the unmodified siphons, flooding between spray cones reduced

The siphons have an inner diameter of 2.3 mm and outer diameter of 3.1 mm. This siphon size was chosen out of convenience of its design, a large siphon diameter in theory would be able to pull more fluid away from the surface and thus improve drainage at the cooled surface. The grooves in the copper plate were 0.5 mm deep and 2.8mm wide. The siphon tubes and the grooves can be clearly seen in Figure 4.5 along with the spray cones. The 25 siphons are in close contact with the cooled surface in Design 1 (Figure 4.7) and 37 in Design 2 (Figure 4.12). The full size of the spray cooler array is 4 by 4 for a total of 16 Nozzle. The sixteen spray nozzles were distributed evenly so that their spray cones covered a square area of 16.8 cm^2 , this roughly give a cooled area of 1 cm^2 for a single nozzle. The selected spray nozzles were 1/8GG-FullJet 1 from spray systems co and were selected based on their even spray distribution and higher droplet velocities (Chen, 2002).

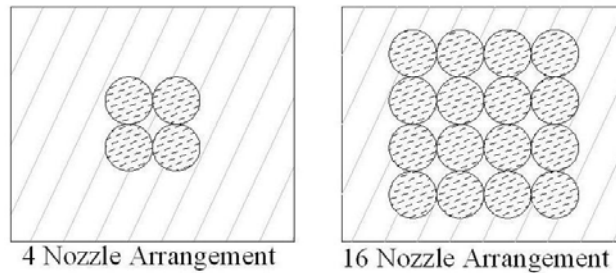


Figure 4.5: Spray Nozzle Configuration

The spray cooler array was tested in two configurations the first being a 4 nozzle and 16 nozzle arrangement Figure 4.6. The 4 nozzle arrangement was mainly used to visualize the fluid dynamics occurring at the cooled surface.

The fluid management system and the spray coolers nozzles were designed to be compact and scalable to any surface area. The spray cooling unit consists of two manifolds, one being the high pressure water manifold which feeds the spray nozzles, and the second manifold is the suction manifold which pulls suction from all the siphons. The spray nozzles are isolated from the suction manifold via an array of 16 copper tubes which pass through the suction manifold. This design can be seen in Figure 4.7.

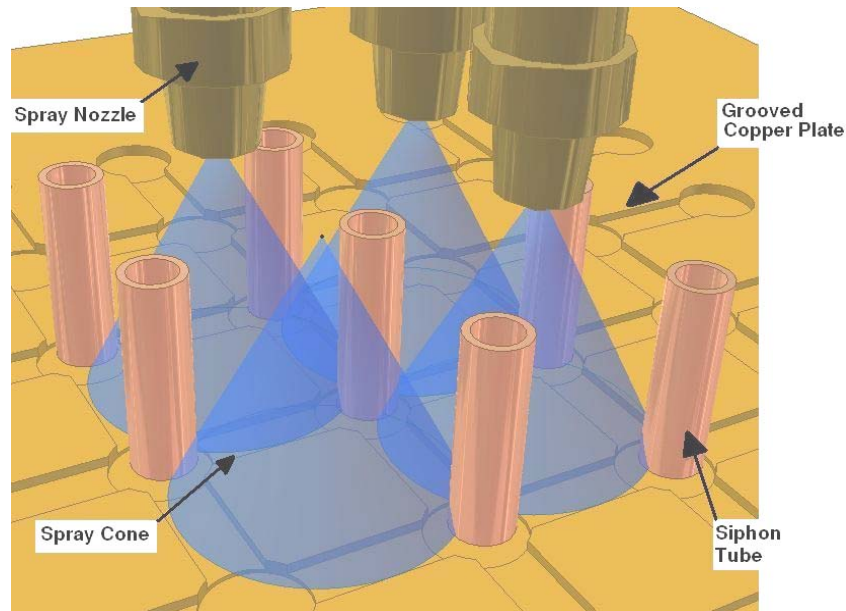


Figure 4.6: Overall siphon placement

The siphons in this design are 1 mm above the grooved plate. This was done so that the thin film thickness may be controlled via the use of varying the suction through the siphons. The spray nozzles are positioned 13mm above the cooled surface which was determined as the optimum distance. The suction manifold is evacuated via eight holes around the side of the suction manifold.

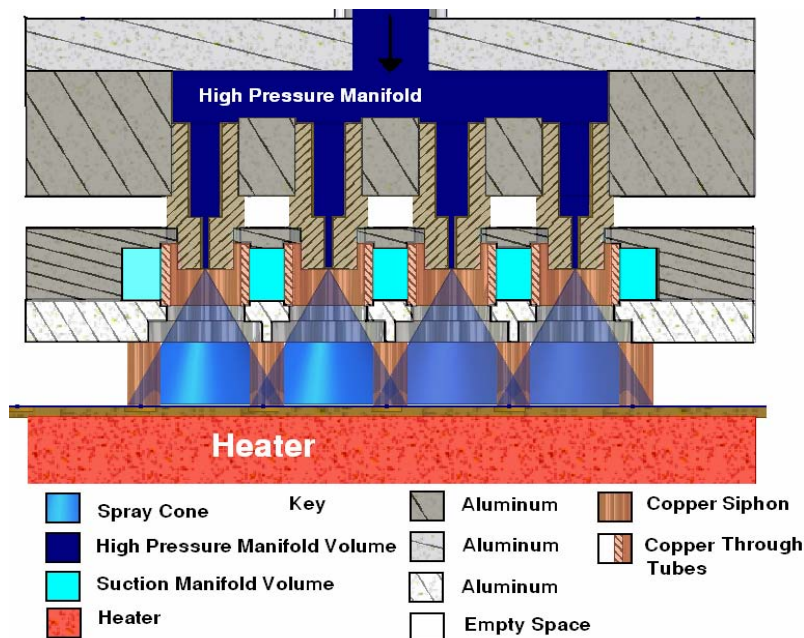


Figure 4.7: Overall spray cooler design

4.2 EXPERIMENT SETUP

The spray cooler and fluid management system were tested in a closed loop setup with the spray cooler in the horizontal position as shown in Figure 4.8. The loop starts at the main water reservoir then a gear pump capable of pumping 116 gallons/min at 90 psi was used in conjunction with a bypass valve to provide liquid water to the spray cooler at a range of 20 to 40 psi. The flow rates of the pump depend upon spray nozzle configuration shown in Figure 4.8.

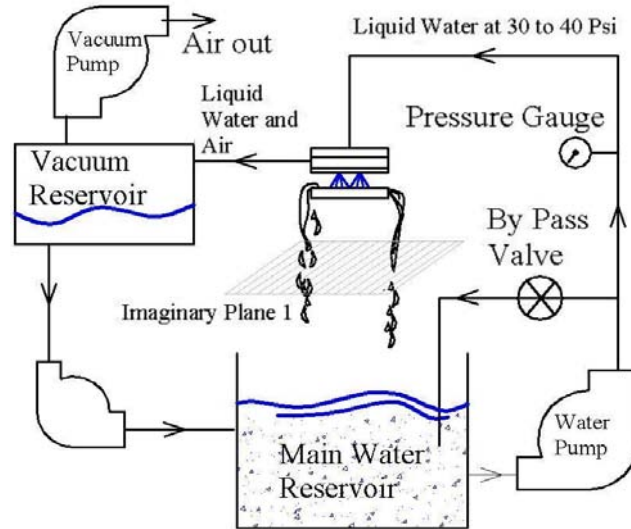


Figure 4.8: Experimental setup

The water is then passed through the spray cooler nozzle and atomized into a spray via the head pressure. Excess liquid that was not removed by the siphons was allowed to drain over the edge of the cooled plate back into the water reservoir. The vacuum reservoir was evacuated to 2 in-Hg via a single air drive vacuum (Level 1 Suction) and to 4 in-Hg via two vacuums (Level 2 Suction). The liquid that accumulated in the vacuum reservoir was then pumped back to the main water reservoir with a diaphragm pump. Flow rates for all experiments were measured by capturing the excess liquid in a graduated cylinder for 30 seconds at the imaginary plane shown in Figure 4.8. This gives an error in the flow rates recorded to be +/- 0.2 liter/min.

Fluid Dynamic Analysis & Setup

A fluid dynamic analysis was conducted so that the removal of the excess liquid from the cooled surface may be maximized. All of the experiments concerning visualizing of the fluids were done in the absence of heating and in a horizontal position. Suction effectiveness is defined by the following equation.

$$Eff_{Suc} = \frac{\dot{Volume}_{in} - \dot{Volume}_{edges}}{\dot{Volume}_{in}} * 100\% \quad (4.1)$$

The spray cooler has one input and two exits, one being over the edges and the other being the siphons. The input volume flow rate was measured with the suction system off. The volume flow rate over the side ($\text{Volume}_{\text{edge}}$) was then measured when the suction system was turned on.

The main steps taken to improve the suction efficiency were done by modifying the suction tube or siphons in the spray cooler. For visualization purposes a clear piece of grooved Plexiglas was used to replace the grooved copper plate (Figure 4.6). All of the observed flows were transferred to AutoCAD drawings for easy visualization.

Figure 4.4 shows the fluid flow observed from the side of the spray cooler in the 4 nozzle spray nozzle arrangement. The full 16 nozzle spray fluid flow could not be observed from the side of the spray cooler due to a lack of visibility created from excess misting. Figure 4.4 shows clearly that the fluid has to pass underneath the siphons to be removed from the cooled surface. Figure 4.4 also shows a small amount of flooding in between the four nozzles.

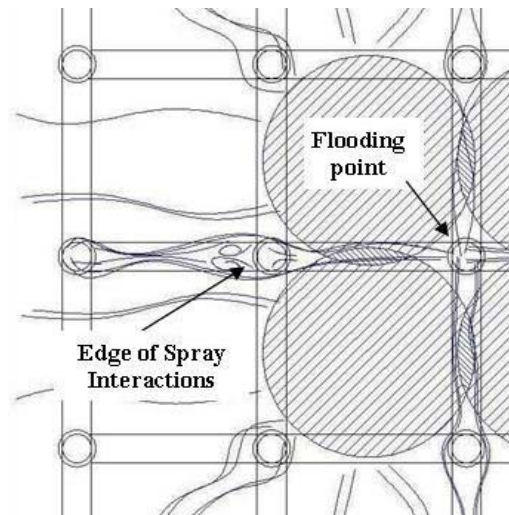


Figure 4.9: Observer flow from bottom of grooved plate: un-modified siphons used

Figure 4.9 is a bottom view of the grooved plate in the horizontal configuration, only the left portion of the plate is shown here since the flow pattern is symmetric. Figure 4.9 is based on the flow visualized through a duplicated grooved plate made out of Plexiglas. It should be noted that the fluid mainly exits via spray cone interaction lines as in Figure 4.3. Additionally, it was noticed that the grooves in the Plexiglas helped channel the fluid away from the spray cones

The siphons were placed at the intersection of the grooves on copper plate to catch the excess liquid, but due to the high fluid momentum a lot of the excess liquid flowed around the siphons then off the edge of the cooled plate (Figure 4.9). The poor efficiency results ($\text{Eff}_{\text{suc}} < 50\%$ for 4 nozzle array at 30PSI head pressure) obtained from this initial siphon design led to the design and testing of five different siphon nozzles seen in Figure. 4.10.

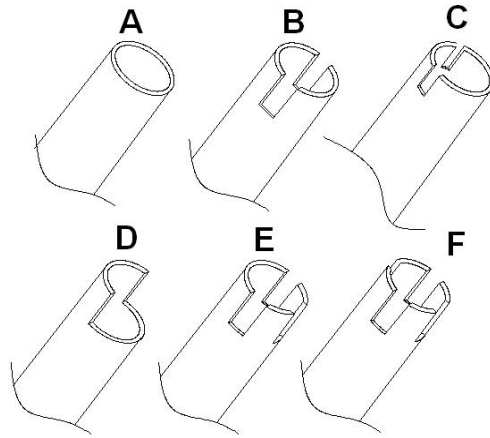


Figure 4.10: Modified siphon designs

The siphon designs are as follows: siphon B has two slits opposite to each other, siphon C has two slits at 60° to each other, siphon D is a 160° cut of the tube, and siphon E had three cuts each at 90° from each other, and siphon F has 4 equally spaced cuts. These siphons were then tested in several areas in the spray cooler and the flow around them was noted with the suction on and off. The results of these tests can be seen in Figure 4.11.

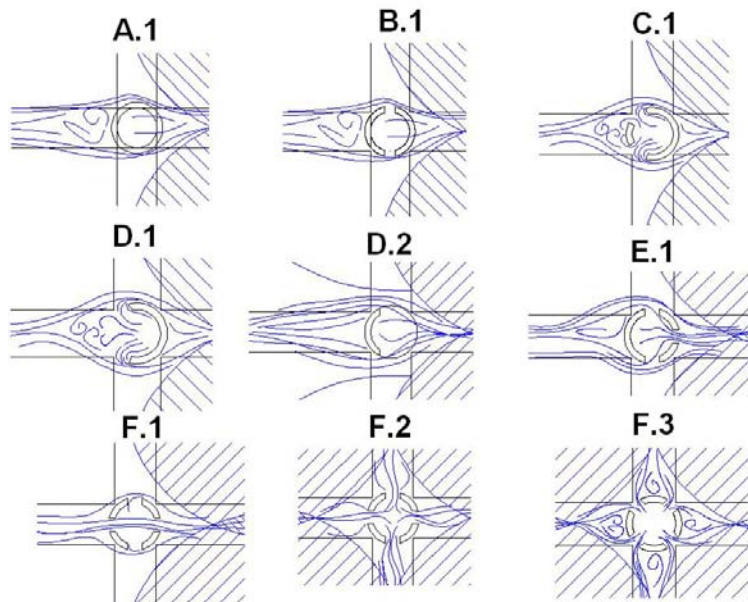


Figure 4.11: Fluid dynamics visualization around the base of the siphons

The bulk of the liquid leaving the plate flowed down the grooves in the plate thus around the unmodified siphons. The flow around the unmodified siphons (siphons A) can be compared to the flow of air over an infinitely long circular cylinder (Young, 1997).

Again, the effectiveness of siphon at removing excess liquid was determined by observing the flow through the clear grooved Plexiglas plate. Due to the space constraints the fluid velocity on

the grooved plate could not be measured around or in front of the siphons. All of the siphon designs were tested at the edges of the spray interactions (Figure 4.9). Only Design F was tested between the spray cones (Figure 4.9, and 4.10). The results of these tests are shown in Figure 4.11.

Siphon A.1 in Figure 4.11 was the least effective in removing liquid from the cooled surface. This is due to the liquid flowing around the siphon with a very small amount flowing under it and out. Siphon B.1 also showed a low effectiveness in removing the liquid. This can be accounted to the liquid flowing faster around the sides due to circular shape of the siphon. Siphon C.1 was based on the idea that behind the siphon a steady wake region was formed (Chen, 1995). This region would have a reduced pressure so it should be a logical place to extract the excess liquid. Testing this design confirmed that it was better at liquid removal. This was taken one step further in siphon D.1 which when tested was an improvement over siphon C.1. Siphon D was tested in different orientation the most notable being siphon D.2 which was not as effective as siphon D.1. Siphon E.1 was tested in the hope that the stagnation point in front of the siphon would force the liquid into it. But upon testing it showed the same effectiveness as siphon D.2. Siphon F.1, surprisingly, allowed liquid to flow through the siphon. Only siphons F.2 & F.3 were tested between the spray cones, because they had the most even slit distributions. Testing proved that siphon F.3 was superior to siphon F.2. These series of siphon tests lead to the final spray cooler design which is slightly different than the one shown in Figure 4.7 for it had additional siphons on the outside of the spray area. This can easily be seen in Figure 4.12.

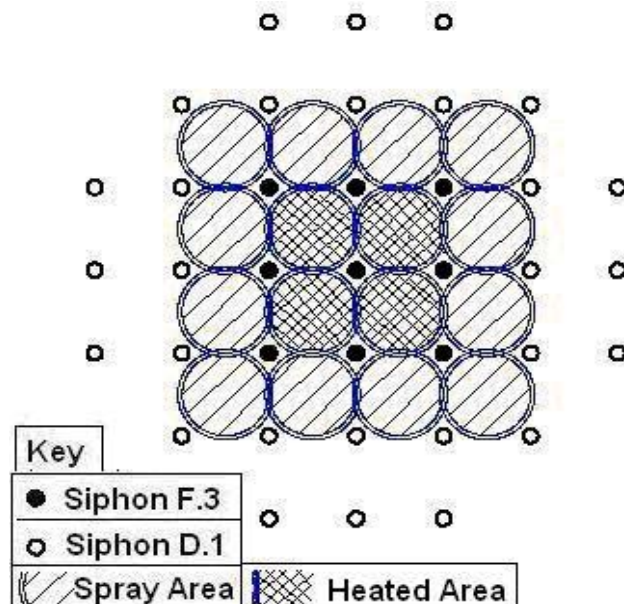


Figure 4.12: Siphons placement and siphon type
Refer to Figure 10 for type of siphon

Suction Effectiveness Testing

The suction system was tested at two levels of suction, the first being only one suction pump was used producing 2 in-Hg at the suction reservoir and 1 in-Hg in the suction manifold. For the second level of suction two suction pumps were used in series to produce 4 in-Hg at the suction reservoir and 2 in-Hg at the suction manifold. The method of determining suction effectiveness is described at the beginning of the fluid dynamic analysis section.

The suction system underwent one evolution from design 1 to design 2 before the actual thermal testing. Design 1 was the initial suction system design where only 25 siphon tubes were employed regulate flooding on the cooled surfaces; whereas, design 2 utilized 37 siphon tubes to drain the cooled surface. All suction efficiency data is the average of four trials.

Figure 4.13 shows that the suction effectiveness was greatly improved by the addition of the 12 extra siphons outside the spray cooled area. This led us to conclude that the bulk of the fluid is removed at the edges of the spray cooled area. Without the suction it was observed that the edges of the spray cooler (design 2) were completely flooded with liquid flows freely over the sides. With the suction on, regardless of the suction level, it was observed that the spray cones would become visible and liquid flowing over the sides would reduce to a trickle (Figure 4.13).

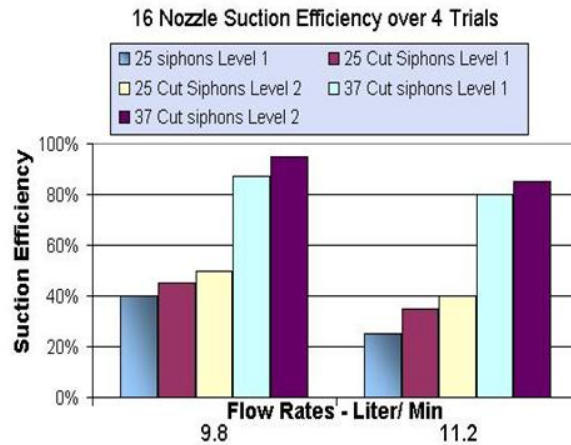
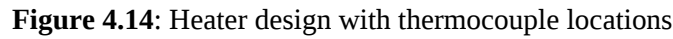


Figure 4.13: Suction Effectiveness for a 16 Nozzle Array: No Heat used during tests

Thermal Design & FEM Analysis

It was decided to heat only the areas of the 4 inner spray cones; the heated area can be seen on the previous page in Figure 4.12. This was done for several reasons: the design of the heater was much more simple and compact; the heat flux would be more uniform over a smaller area (4.41cm^2), and the electrical input power would be lower and more manageable.


$$\frac{E_{in}}{A_{top}} = \frac{3*400W}{2.1cm*2.1cm} = 272 \frac{W}{cm^2} \quad (4.2)$$

A Finite Element analysis was conducted on the heater block with the use of Cosmos Design Star 3.0. The FEM analysis was conducted to find out the effects of heat spreading at the cooled surface and to see if 4.8mm depth of the T.C is enough to accurately measure the heat flux.

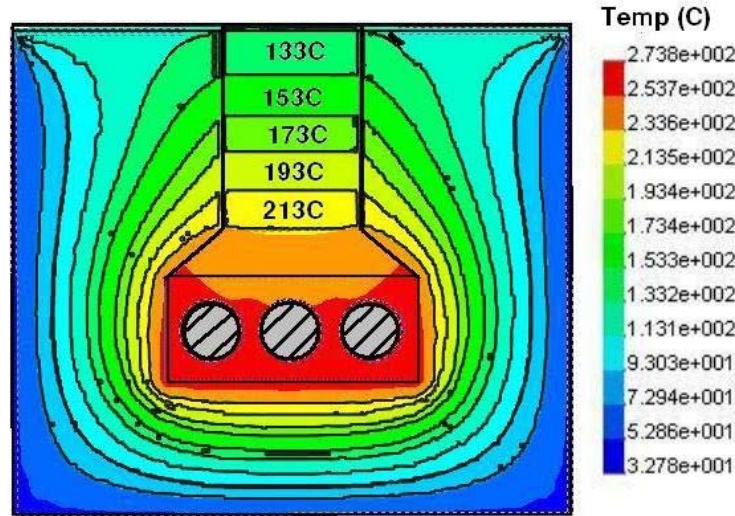


Figure 4.15: Sectional temperature profile of heater

The results of the FEM analysis were displayed in Figure. 4.15. Here the three heaters were supplied with 205W each for a total of 615 W, which translates to a theoretical heat flux of 140 W/cm² over the top area of the pedestal which is 4.41 cm². The top of the pedestal was maintained at a temperature of 127°C and the outer surfaces of the Durablanket insulation was cooled at 10 W/cm² and at a sink temp of 50°C, to simulate room temperature convective cooling. The thermal conductivity of the copper was taken as 393 W/m-K.

The temperature distribution along the pedestal heater is a uniform gradient, thus implying that the flux inside of the pedestal is uniform too. The FEM model shows a heat flux of 135 +/- 2 W/cm², inside the neck, whereas, the heat fluxes through the Durablanket insulation is less than 0.5 W/cm² around the neck. From this one can conclude that the FEM models is predicting a loss of,

$$E_{in} = E_{out} = E_{cooled} + E_{loss} \quad (4.3)$$

$$615W = 135 \frac{W}{cm^2} * 4.41cm^2 + E_{loss} \quad (4.4)$$

$$\frac{E_{loss}}{A_{out}} = \frac{61.5W}{302cm^2} = 0.2 \frac{W}{cm^2} \quad (4.5)$$

Solving for E_{loss}, the heat loss through the insulation is found to be 61.5 W at the outer heater shell is. Taken the outer surface area of the heater shell is to be 302cm², gives a heat loss of 0.2 W/cm² on the outer shell of the heater. This FEM model results was proven valid after the thermal results were found and compared to it, see the section title “Thermal Calculation & Results.”

4.3 THERMAL TESTING PROCEDURES

All thermal tests were conducted in the same manner. First the water in main water reservoir (Figure 4.8) was heated to 70°C, and then the spray coolers were turned on. Next the spray cooler were set to a flow rate of 9.8L/min, this setting was based upon the head pressure of the spray coolers. The error in the flow rates were estimated to be +/- 0.2 L/min which is based on the previous suction effectiveness testing data. Next the heaters were set to the desired heat flux, via an AC variac which range from 0 to 120 volts. The input voltage was read from a voltmeter with an accuracy of +/-0.05 volts. Thus the theoretical input heat flux was calculated by

$$E_{in,max} = \frac{Volt^2}{R * A_h} \quad (4.6)$$

Equation 4.7 represents the maximum heat flux attainable, Note this agrees with equation 4.3.

$$E_{in,max} = \frac{Volt^2}{R * A_h} = \frac{120V^2}{13\Omega * 4.1cm^2} = 270 \frac{W}{cm^2} \quad (4.7)$$

Once the flux level was set and a steady state temperature was achieved, the DAQ (National Instruments PCI-4351 hooked up to a TBX-68 board) would read all the channels at a rep rate of once a second. Secondly the DAQ measured temp. in a differential or floating mode and combined this with a cold junction ground temp. error is +/- 0.35°C. The T.C data was collected for 1 min with the suction off, 1 min. with the suction at level 1, and 30 sec. with the suction at level 2. This was repeated for all successive heat fluxes up to a maximum input of 169 W/cm². Once one round of testing was completed the flow rate was readjusted to 11.2 L/min and the experiment was repeated. Additionally, the heaters were tested only up to 150W/cm² which is about 60% of their total power. This was done to avoid damage to the heater.

The setup as seen in Figure 4.16 is very large because of the increase in flow rates. The water heater is not shown in this picture but it is located directly underneath the table. The water heater is a 6 gallon standard water heat, the thermostat was bypass and the heater coil was connected directly to the wall so that a water temp of 90 C could be achieved. All of the tubing in the pump line is 5/8in copper tubing; the drainage line for liquid return to the water heater is 1 in inner diameter copper tubing. There are 4 vacuum lines that are 3/8in vinyl tubing. All reservoirs were insulated; the water heater had built in insulation. The vacuum reservoir has 6 in thick home insulation around it and the drainage pot had bubble rap insulation around it. All of the copper tubing was cover by water heater insulation (1 inch thick). Keeping heat loss to a minimum helped maintain the working fluids temperature. The system is open so the vapor created at the heated surface escapes free atmosphere.

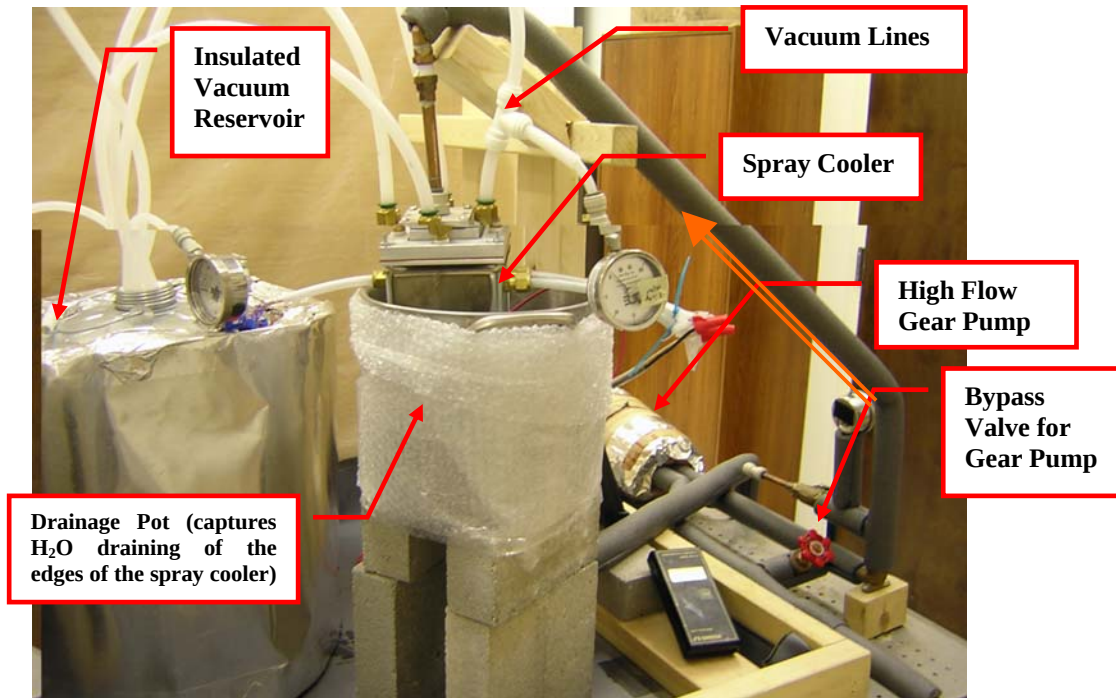


Figure 4.16: Experimental Setup for Multiple Nozzle Spray cooler

The vacuum chosen to provide suction for this experiment are two RIDIGED 2.5 hp shop vacuums. These vacuums attached in series provide two levels of suction for the experiment. During the experiment it was noticed that the vacuum performance was diminishing. This was due to the filter inside becoming saturated with water vapor. The problem was solved by removing the filter inside of each vacuum. The vacuums input line was then connected to the vacuum reservoir chamber as shown in Figure 4.17.



Figure 4.17: Vacuums

These vacuums were selected because they were easily obtainable, that way testing could start as soon as possible. Below Figure 4.18 shows the custom Lab View interface created for this test.

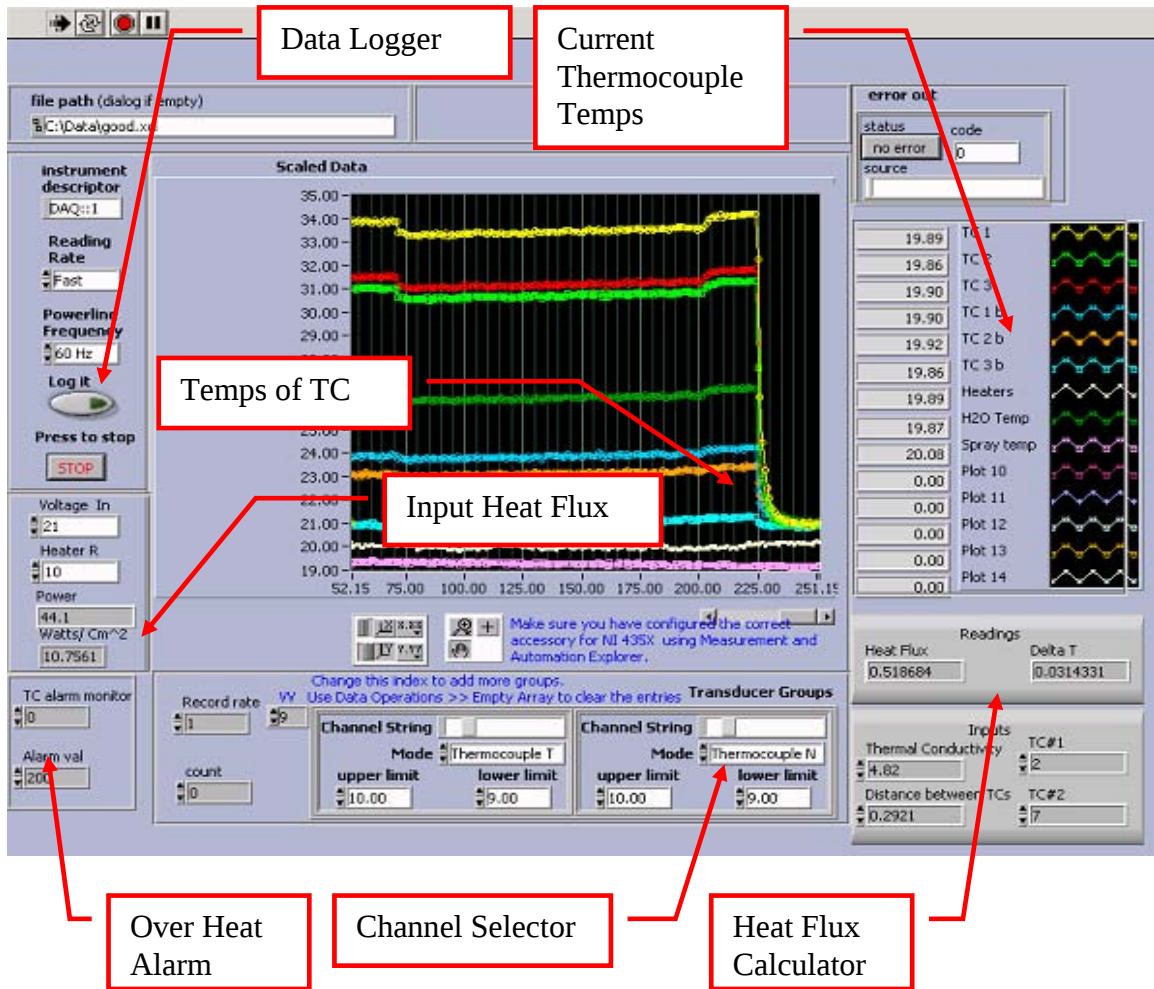


Figure 4.18: Custom Lab view Layout

The spray cooler has a secondary attachment that will keep the excess liquid from draining over the sides. This will allow the spray cooler to be tested inside of the chamber without loose fluid moving around inside of the chamber.

Thermal Calculations & Results

The T.C.s temperatures were used to calculate the experimental heat flux via equation 4.8.

$$q''_{n-m} = \frac{k_{inter} * (T_n - T_m)}{X_{n-m} / 10} \quad (4.8)$$

Where n, and m are the thermocouples used, and k_{inter} is the interpolated value of thermal conductivity which is given by

$$K_{inter} = \left(\frac{(401 - 393) \left(127 - \frac{(T_n + T_m)}{2} \right)}{100} + 393 \right) / 100 \quad (4.9)$$

Note: At 200°C, k for copper is 393 W/m-K and at 27°C; k is 401W/m-K.

The cooled surface temperature T_w was calculated via T_1 and the averaged heat flux.

$$T_w = T_1 - \frac{q_{ave}'' X_{1-w}}{k_{inter}} * 10 \quad (4.10)$$

Thermal Results

This section presents the collected experimental heat transfer data from the spray cooler design 2 with 37 siphons in the form of two plots. Both plots show average heat flux vs. $T_w - T_{sat}$ and are for water sprayed at 70 +/- 1°C, only their input volume flow rates per nozzle differ from Figure 4.19 to Figure 4.20.

Figure 4.19 shows the impact suction has on the heat flux. The flooded surface situation presented in Figure 4.19 by the line labeled “NO Suction”. From the “1 Vac” line (One Vacuum) can see an improvement from the pervious of on the average 20W/cm² at similar $T_w - T_{sat}$ temperatures above 10°C. Secondly, the “2 Vacs” line shows even greater improvement over the flooded situation. In Figure 4.19, the “2 Vacs” line shows an average of 30W/cm² for similar temps above 5°C for $T_w - T_{sat}$.

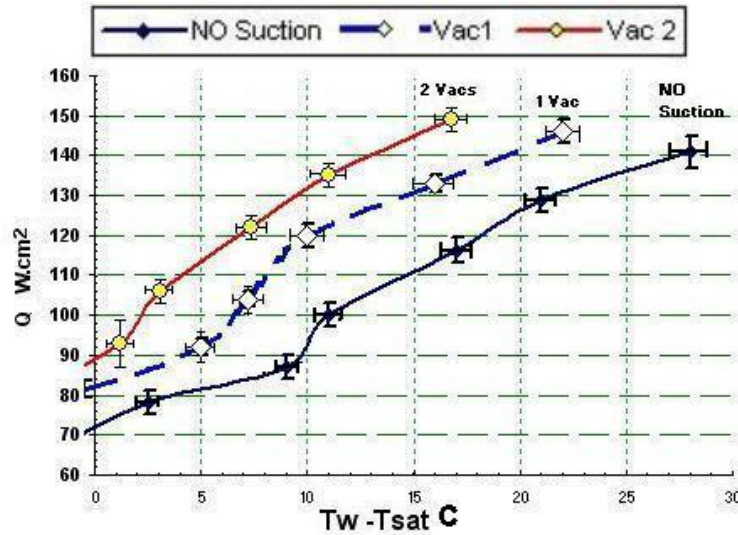


Figure 4.19: Q vs. $T_w - T_{sat}$ 20 PSI Head Pressure and Flow Rate of 9.8 Liters/min

Figure 4.20 shows the heat flux vs. $T_w - T_{sat}$ results for a higher input volume flow rate of 11.2L/min. The higher spray nozzle volume flow rate along with the increase in droplet velocity (Chizhov, 2004) due to higher head pressures help the heat flux on the average 10W/cm² over the results shown in Figure 4.19. Figure 4.20 shows that the suction improves the heat flux at a similar $T_w - T_{sat}$ temperature.

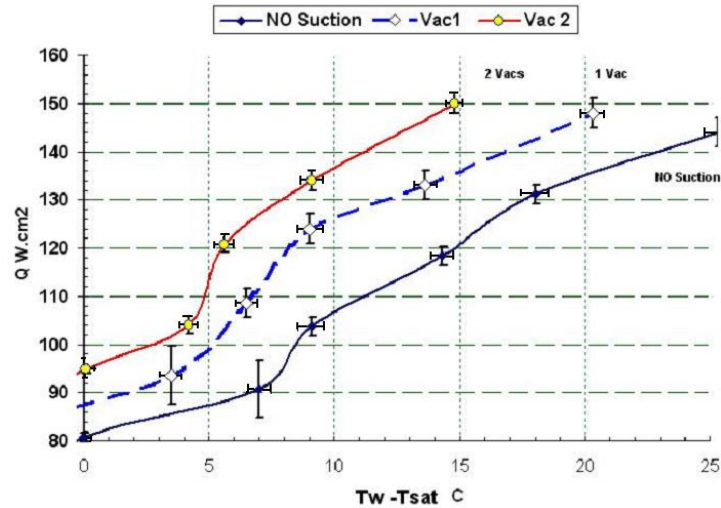


Figure 4.20: Q vs. $T_w - T_{sat}$ for 30 PSI head at Flow rate of 11.2 Liter/min

From both plots one can see that the higher the suction the greater the heat flux achieved. A greater heat fluxes could have been achieved if the smooth grooved copper plate was sanded. But this was not done because it was outside of the scope of this experiment.

Experimental results from the thermal testing show at the same input power loss was 96W, which translates to average 0.3 W/cm^2 , heat loss through the shell of the heater. The similarity between the FEM results and the experimental results validate the FEM analysis.

Uncertainties

The thermocouples on both sides of the heater were systematically 0.4°C different from each other. This causes a slight difference in the heat flux calculated from via T.C. 2b, it was also found that T.C. 2b was not in direct contact with the copper block as the other T.C.s were. This created a slight discrepancy in the data, and thus the heat fluxes calculated using T.C.2b were not used. The plotted averages of heat fluxes had a deviation of 2 to 3 W/cm^2 . The error in calculated surface temperature is mainly due to the error in the thermocouples ($\pm 0.35^\circ\text{C}$) and the error in the associated heat flux, which results in a surface temp error of ($\pm 1.5^\circ\text{C}$). Also the calculated surface temperature could have a systematic error due to the solder contact between the copper heater and the copper grooved plate, this error was estimated to be -1°C .

4.4. CHAMBER DESIGN

The results from the multiple nozzle spray cooler are extremely promising, however, heat transfer effects in microgravity and variable gravity are not known. Following this logic a flight chamber was designing to house and test the spray cooler in variable G and micro-gravity environments. A detailed description of the chamber design is depicted below.

The manufacture chamber was design mainly of 4020 Extrusion aluminum parts and an aluminum chamber. Dimensions of the chamber are as shown in Figure 4.21. The chamber was design with flight requirements in mind, which is, it will be able to with stand the flight loading required by NASA.

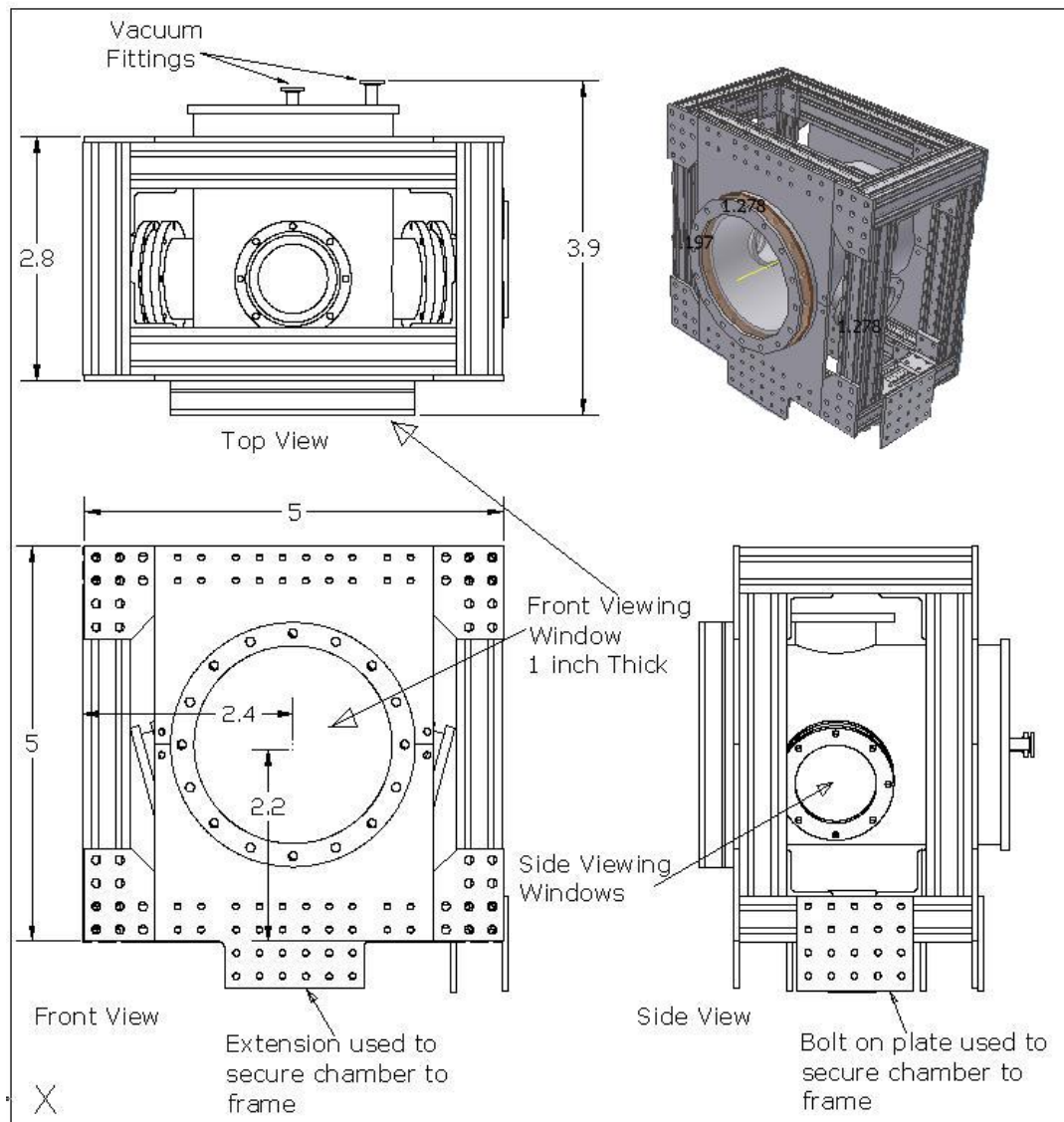


Figure 4.21: Dimension of manufactured Chamber

Figure 4.21 shows the main window which has an 8.5 inch in diameter viewing area and 3 smaller ports in the side of the chamber. The smaller ports such as the top and the left and right sides will be used for pipe feed through and will not be transparent. The back of the chamber has fitting for vacuum pumps and drainage. The chamber side are 1/4in thick 6061 T6 Aluminum and the front window is 1inch thick Plexi-glass. The chamber is rated at 300 psi via o-ring and hoop stress calculations. The chamber can also hold a vacuum and has a copper o-ring around the front window as an additional seal.

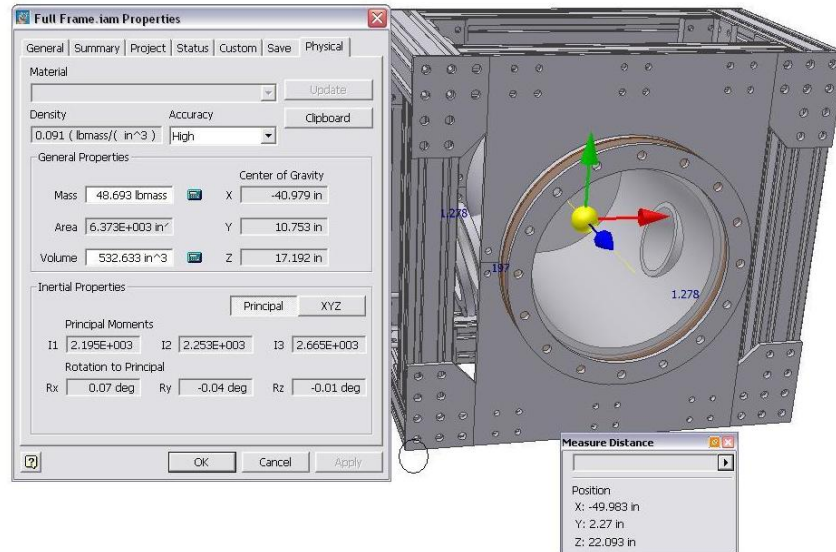


Figure 4.22: Physical Location of Center of Gravity

The chamber mass is 48 lbs in the CAD model and physical model mass is 52. The Center of gravity defined from the bottom left corner in red is CG= 8.5” of the height, 9.004” of the Length, 4.9” of the width. This chamber has two seals around the main window, one being a Buna-O-ring and the second being machined copper o-ring. Figure 4.23 has the bottom retaining plate removed so one can see the port window. The top of the chamber is resting on the desk in this picture.

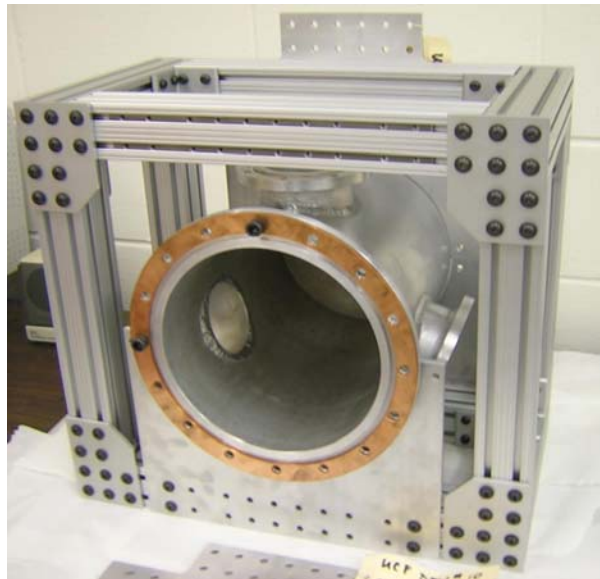


Figure 4.23: Front view of manufactured chamber

4.4 CONCLUSIONS

The fluid management system or suction system which regulated flooding on this 16 spray nozzle array improved heat transfer on the average of $30\text{W}/\text{cm}^2$ for similar values of superheat above 5°C . It was also concluded that increasing the amount of suction increased the heat flux

and thus the heat transfer coefficient. Suction effectiveness was helped greatly by adding extra siphons outside the spray area. Additionally suction effectiveness was also increased by adding small slits to the sides of the siphons. Thus, this design presents a solution to the problem of surface flooding in a multiple pressure atomized spray nozzle array.

5. THERMAL MANAGEMENT OF DIODE LASER ARRAYS

5.1 THERMAL MANAGEMENT AND SPRAY COOLING

The removal of waste heat from diode lasers is one of the major issues in utilizing high power diode laser arrays. In this section, we discuss the basic physics of cooling and spray cooling which we proposed and applied to a high power diode laser array.

Thermal conduction and convection (Frank, 1996)

Heat transfer is energy in motion due to a temperature difference. Here we discuss the simple heat transfer mechanisms we utilize in our experiments and simulations. There are three different heat transfer modes: conduction, convection and radiation. Since the radiation effect is negligible when the subject temperature and the ambient temperature are close, we will focus mainly on heat conduction and convection.

The rate equation for heat conduction is also known as Fourier's law

$$q_x'' = -k \frac{dT}{dx} \quad (5.1)$$

The heat flux q_x'' (W/m²) is the heat transfer rate in the x-direction per unit area perpendicular to the direction of transfer. It is proportional to the temperature gradient, dT/dx . The coefficient k is the thermal conductivity (W/m·K), a property of the material through which the heat transfer takes place. The minus sign indicates that the direction in which the heat flux is transferred toward a lower temperature. A more general form of the time independent heat conduction equation is

$$\bar{q} = -k\nabla T \quad (5.2)$$

Here, q is the vector form of heat flux. If we consider heat generation as well, the equation is

$$\nabla \cdot (k\nabla T) + \dot{q} = \rho \cdot c_p \frac{\partial T}{\partial t} \quad (5.3)$$

where $\dot{q}$ is the volume heat generation (W/m³), ρ is the density (kg/m³) and c_p is the specific heat (J/kg·K) of the material. At the steady state condition, the temperature distribution is independent to time. Therefore, Equation (5.3) can be simplified as Equation (5.4).

$$\nabla^2 T + \dot{q} = 0 \quad (5.4)$$

Heat convection is more complicated than heat conduction because the fluid mechanics has to be considered. However, a very simple equation can be written down as

$$q'' = h(T_s - T_\infty) \quad (5.5)$$

$$q = \int_{A_s} q'' dA_s \quad (5.6)$$

Here, q'' is the local heat flux and h is the heat transfer coefficient. T_s is the surface temperature and T_∞ is the coolant temperature.

Spray cooling

As mentioned in the previous chapter, traditional forced convection and micro/macro channel cooling have some inherent problems with the coolant flow rate and system complexity. Therefore, we developed an evaporated spray cooling (ESP) technology that can considerably reduce the required coolant flow rate and pump pressure drop (Sehmbey, 1997).

Evaporated spray cooling uses a nozzle to spray fine droplets on a surface at a temperature higher than the boiling point of the liquid. The liquid droplets evaporate and absorb heat corresponding to the latent heat of vaporization when striking the hot surface. However, when the droplets wet the hot surface with a layer of water, vapor forms a thin layer of bubbles which isolates the cold water from the hot surface. Under such condition, the heat transfer coefficient will be dramatically lowered. Therefore, one role of the impinging droplets is to break up the bubbles and blow away the water vapor to assure that the hot surface is continuously cooled efficiently.

This technique utilizes the phase change boiling/evaporation transition as the heat transfer mechanism; unlike micro/macro channels or cold-water heat exchangers that are limited by single-phase convection and the specific heat of the coolant. The use of spray cooling allows all the diodes in an array to be cooled in parallel. It is therefore readily scaled to cool any size array. Utilization of spray cooling also makes possible smaller, lighter thermal subsystems because a much lower coolant flow rate is needed. A comparison of the typical coolant flow rates for a 2 cm² diode array is given in Figure 5.1 (Huddle, 2000).

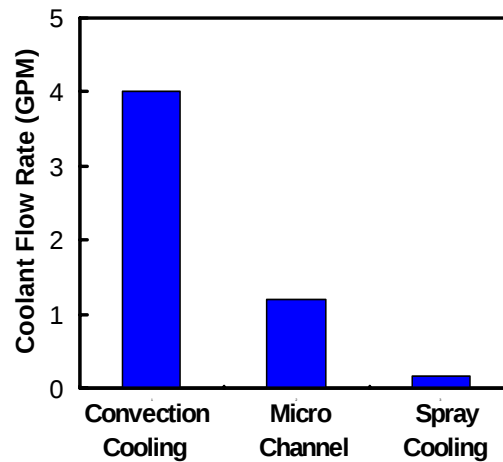


Figure 5.1: Coolant flow rate comparison of different cooling methods

To keep the emitters at a desired temperature, the pressure in the spray cooling chamber can be adjusted or with a wide variety of fluids other than water to be used. The combination of fluid and system pressure will determine the temperature of the diodes. A spray cooling experiment using water at a low chamber pressure was carried out using thick film resistors as the heat source; this proved that spray cooling at low chamber pressure can remove comparable heat fluxes at a low temperature and with similar heat transfer coefficients as can other techniques. Figure 5.2 represents the system pressure versus water boiling point temperature (David, 1996). This figure shows water boiling point is 100°C at 1 bar, or 1 ATM. Figure 5.3 shows the

experimental results of low pressure water and ammonia evaporated spray cooling. At 1 bar the water ESP can remove heat flux up to 600 W/cm^2 . However, the surface temperature is about 40°C higher than the water boiling point at 1 bar which is 100°C . In order to reach a lower surface temperature, a lower system pressure is needed. Since the diode laser performs better at lower temperature, we choose the system pressure of 0.02 bar which corresponds to water boiling temperature of about 17°C . Other liquids should be considered if an even lower temperature is preferred.

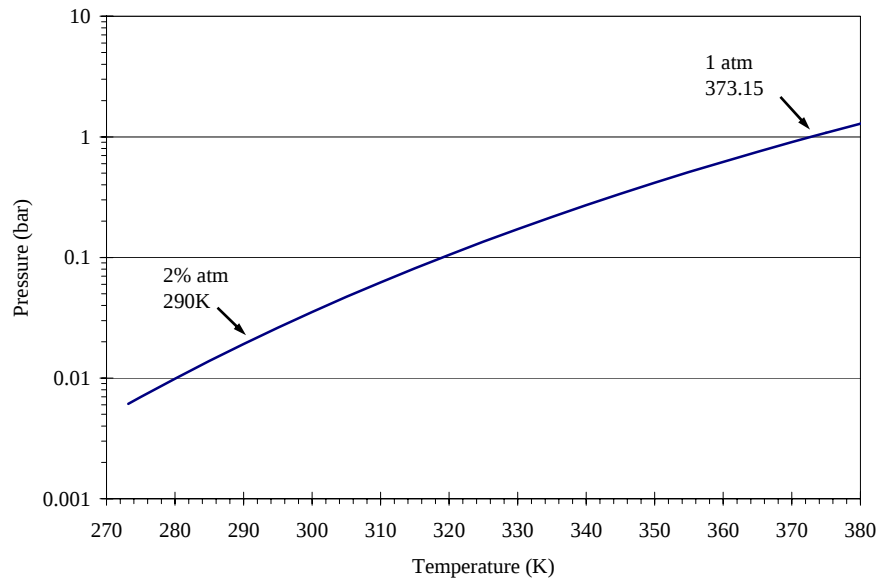


Figure 5.2: The system pressure versus water boiling temperature

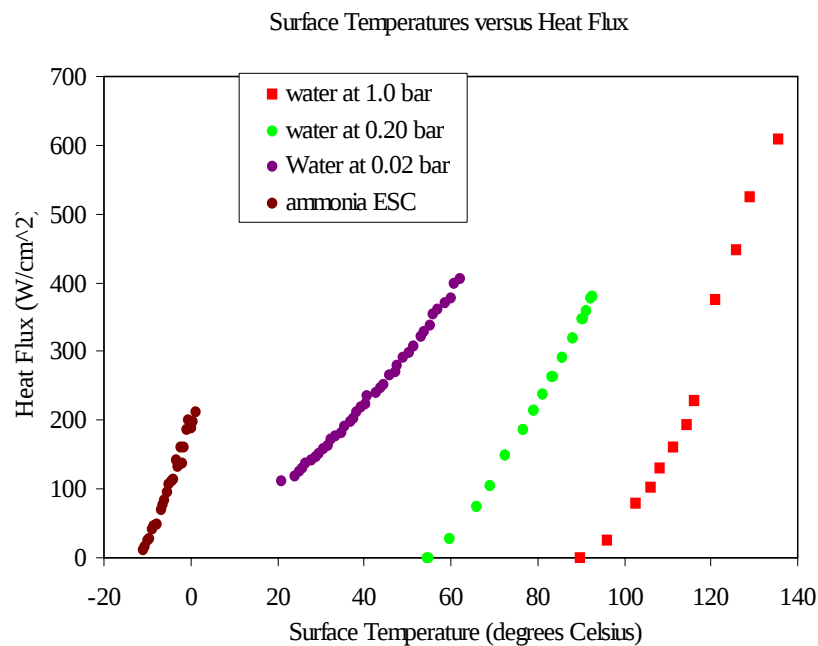


Figure 5.3: Low pressure water and ammonia spray cooling experimental results

Spray cooled diode laser array

A diode laser array was installed so that it could be cooled by spray cooling with a controlled chamber pressure to set the vaporization temperature of the water coolant. Four Coherent 808 nm B1-40C-19-30-A diode bars (with 19 emitters each) were used for the laser array. The output power of each diode bar is 40 W. The detailed specifications are listed in Table 5.1. The diode bars were packed in a traditional stack package and water was sprayed onto the back copper plate, as shown in Figure 5.4. Diode bars were separated by indium coated copper blocks which were 1.1 mm thick. The copper blocks were 1.5 mm deep and their length were the same as the diode bar or 1 cm. A 0.5 mm layer of BeO serves as the common substrate and electrical insulation layer for the array. A 1.2 mm thick copper layer separates the diode array and the spray cooling chamber. Since the system pressure is only 0.02 bar, we designed the chamber to maintain this pressure as shown in Figure 5.5.

The spray cooling chamber was made of stainless steel except the top cover plate which was made of copper and the high power diode array module was soldered to it. The spray water had to be boiled to remove any dissolved gas before being put into the chamber. Otherwise, the presence of such gas will prevent the chamber pressure from reaching the desired 0.02 bar. A gear pump forces water through a TG.3/.6 nozzle with a pressure of 48 psi or 33.1 nt/cm² to form the spray. Water droplets strike the back of the surface to which the diode laser array is soldered. They evaporate and absorb heat from the surface and hence from the diodes. Eventually, the water vapor condenses on the cooling coil and returns to the reservoir of the spray water to complete the cooling cycle.

Table 5.1: Coherent 808 nm B1-40C-19-30-A diode specification¹

Coherent 808nm B1-40C-19-30-A diode		Value	Unit
Optical Characteristics	Output power	40	W
	Center wavelength	808	nm
	Center wavelength Tolerance	±2.5	nm
	Wavelength temperature coefficient	0.28	nm/°C
	Spectral Width (FWHM)	<2.5	nm
	Array Length	10	mm
	Number of Emitters	19	
	Emitter Size	150 × 1	μm
	Emitter Spacing (center-to-center)	500	μm
	Slow Axis Divergence (FWHM) SA	<10°	Degrees
	Fast Axis Divergence (FWHM) FA	<35°	Degrees
	Polarization	TM	
Electrical Characteristics	Slope Efficiency	1.1	W/A
	Conversion Efficiency	>45	%
	Pulse Width		ms

¹ http://www.coherent-lasergroup.de/laserdioden/download/LaserDiodeBars/3.1_CCP.pdf

	Duty Cycle	100	%
	Threshold Current	<10	A
	Operating Current	45	A
	Operating Voltage	1.8	V
	Series Resistance	<0.005	Ω
Thermal Characteristics	Thermal Resistance	0.7	$^{\circ}\text{C}/\text{W}$
	Recommended Case Temperature	25	$^{\circ}\text{C}$
	Operating Temperature Range	15 to 30	$^{\circ}\text{C}$
	Storage Temperature Range	-40 to 60	$^{\circ}\text{C}$

Figure 5.6 shows another spray cooling chamber design. It uses two mist cooling nozzles instead of the cooling coil. The mist cooling spray nozzles spray droplets to cool the water vapor. Also, the heat exchanger is moved outside the chamber. Other parts in this chamber are identical to the design in Figure 5.5.

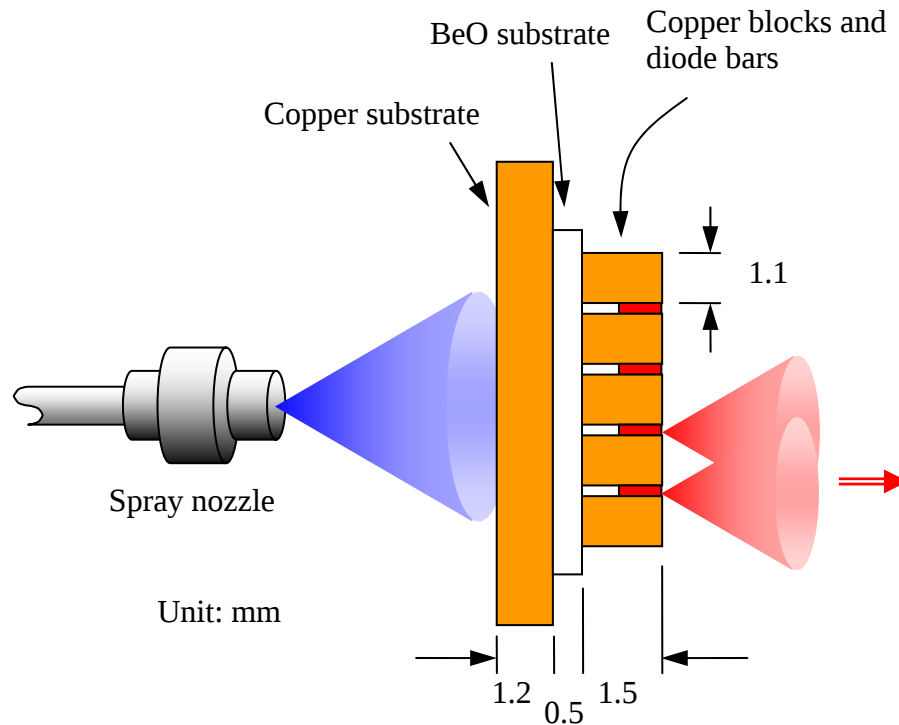


Figure 5.4: The arrangement of the diode array and the spray nozzle

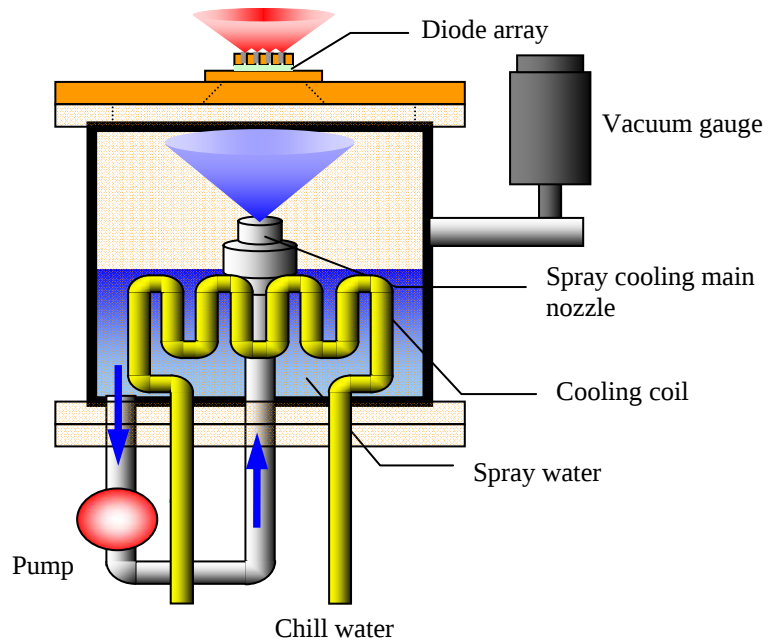


Figure 5.5: The spray cooling experimental set up with the cooling coil design

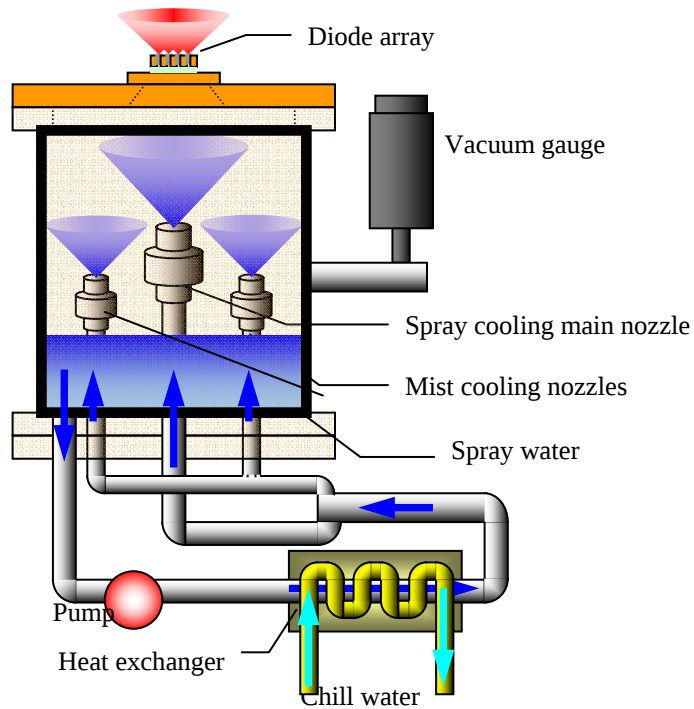


Figure 5.6: The spray cooling experimental set up with the mist cooling design

The power output versus the power input measured for the spray cooled diode array is shown in Figure 5.7. It shows that the efficiency can reach 46% when we increase the current to 45 A. The maximum optical output power was 165 W.

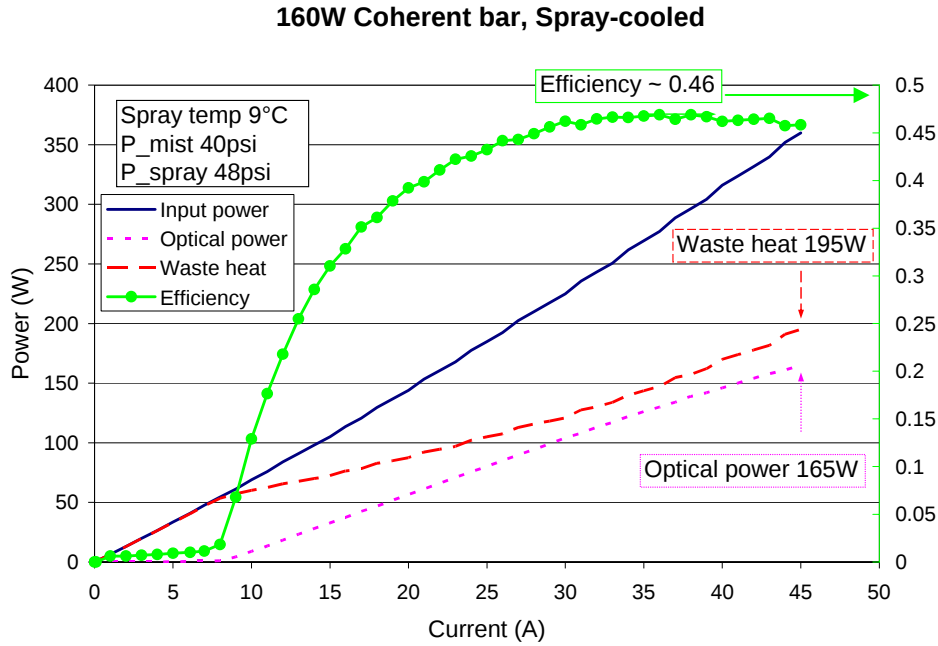


Figure 5.7: Diode laser array input and output power versus current

Using a lens to project an image of the individual emitters we can separately measure the output spectrum of any one. By slightly defocusing the image to make the neighboring emitters' images merge together, we can measure the average spectrum of several emitters at the same time. An Ocean Optics HR2000 spectrometer with high resolution centered at 800 nm was used to measure the output spectrum of several emitters at the right and left sides of the diode array. The results are shown in Figure 5.8. The center output wavelength is about 812.3 nm and the FWHM is less than 2 nm when the array is running at full power.

The wavelength temperature coefficient of the diode laser is given as 0.28 nm/°C in Table . The temperature versus wavelength chart is shown as Figure 5.9. We also measured the output wavelength of each emitter. In this figure, we list the longest and shortest wavelengths of the emitters when the total current is 8.5, 30 and 45 A. We can easily find the corresponding temperature of the emitters. The diode emitters are labeled from 1 to 19. The coding in Figure 5.9 is based on the position of the emitter location. For example, 30A10_4 indicates this emitter is at the 4th bar and it's the number 10 emitter on this bar when the driving current is 30 A. Since the diode bars have only 19 emitters, the 10th emitter is at the center of the bar. Also, because we have four bars in this array, 2nd and 3rd bars are in the middle of the array. Edge effects reduce the emitter temperature of the emitters on the outer edges of the array; consequently, we can expect emitter 10_2 and 10_3 to have higher temperatures as compared with 01_1 or 19_4. We observe that the 2nd bar has higher temperature than bars 1 or 4. When the driving current is higher, the temperature distribution range becomes significantly larger. The experimental data show this trend, but the data is not symmetrical, most likely due to solder layer non-uniformity.

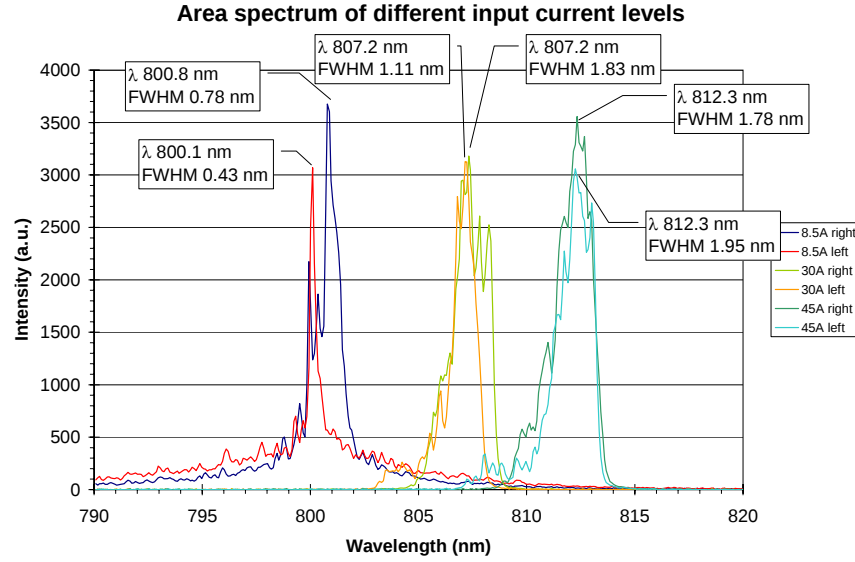


Figure 5.8: Average spectrum at different input current levels

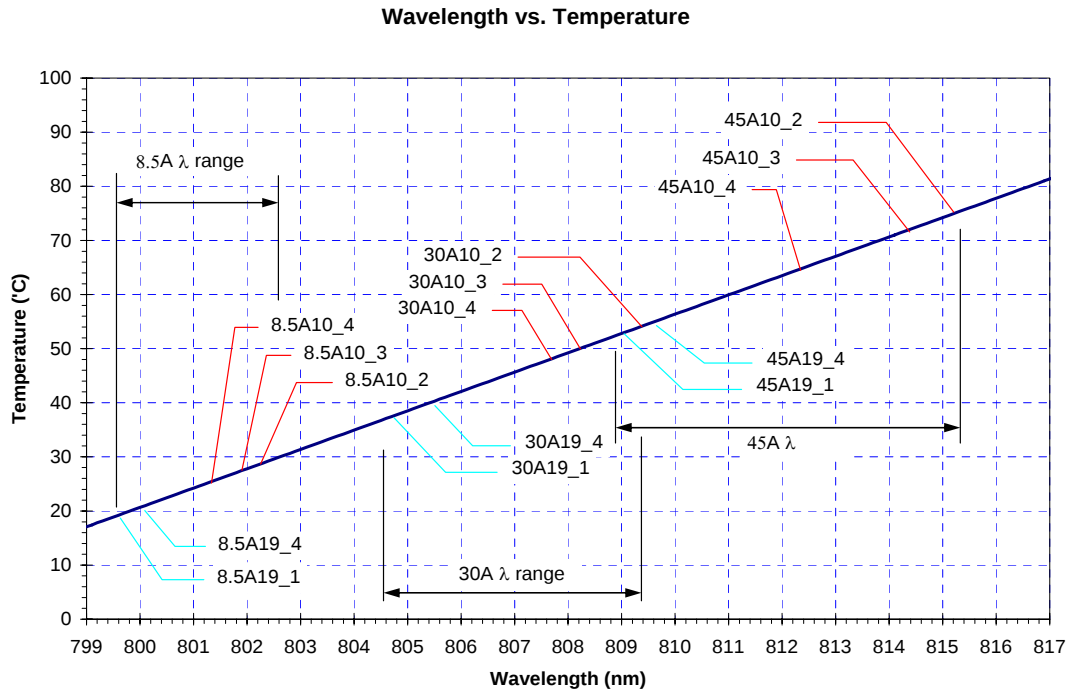


Figure 5.9: Wavelength and temperature of different emitters at different driving currents

The wavelength spread at full power is about 6.3 nm. This implies a temperature difference of about 20° C from place to place in the array. Further the central wavelength is higher than desired indicating that the average emitter temperatures are still too high. Not only does the high temperature and large gradient make the diode emitting wavelength mismatch the Nd:YAG peak absorption wavelength reducing the utility of the array for pumping this medium, it also reduces the emitter operating lifetime and efficiency. Consequently, further packaging improvements are necessary to reduce thermal resistance, lower operating temperature and increase uniformity.

Spray cooled stack packaging results and analysis

By putting the HR2000 fiber probe far away from the diode laser array, it can measure the average output wavelength of the whole array and this measurement at full power gives an average emitter temperature of 62°C. Figure 5.10 shows the relation between estimated average emitter temperature and the waste heat flux at different operating power levels. The total thermal resistance is about 0.082°C·cm²/W based on the bar area (0.65 mm by 1 cm). This includes the thermal resistance for conduction from the emitter to the back plate and the thermal resistance for spray cooling. In order to determine the thermal resistance of this package, the heat conduction at full laser power operation was simulated with finite element methods using the commercially available code called ALGOR. Figure 5.11 gives the computed temperature distribution. In the simulation, a coolant temperature of 10°C was assumed with a heat transfer coefficient of 180,000 W/m²·K. It can be seen that simulation results are consistent with data from the optical measurements in Figure 5.10 which also gives the temperature of the sprayed surface at different power levels. The temperature difference between the emitter and the back plate is about 27°C resulting in a thermal resistance for conduction in this package of about 0.047°C·cm²/W. These results demonstrate that it is very important to reduce the conduction resistance between the emitters and the back surface by reducing the thickness of the diode substrate and the copper plate (see Figure 5.4). If the distance were changed from the value in present package (2.7 mm) to 1.5 mm, the temperature difference between the emitter and the back plate could be reduced to 16°C. The thermal resistance for conduction becomes 0.021°C·cm²/W if the total thickness of the copper is 1.5 mm which includes the thickness of back plate and diode spacers. In such a case, the temperature of the emitter is estimated to be about 51°C which corresponds to operation at 808 nm, or very near the peak absorption wavelength for Nd:YAG. However, a thin structure like this might not have enough mechanical stiffness to sustain the pressure difference of the spray cooling chamber.

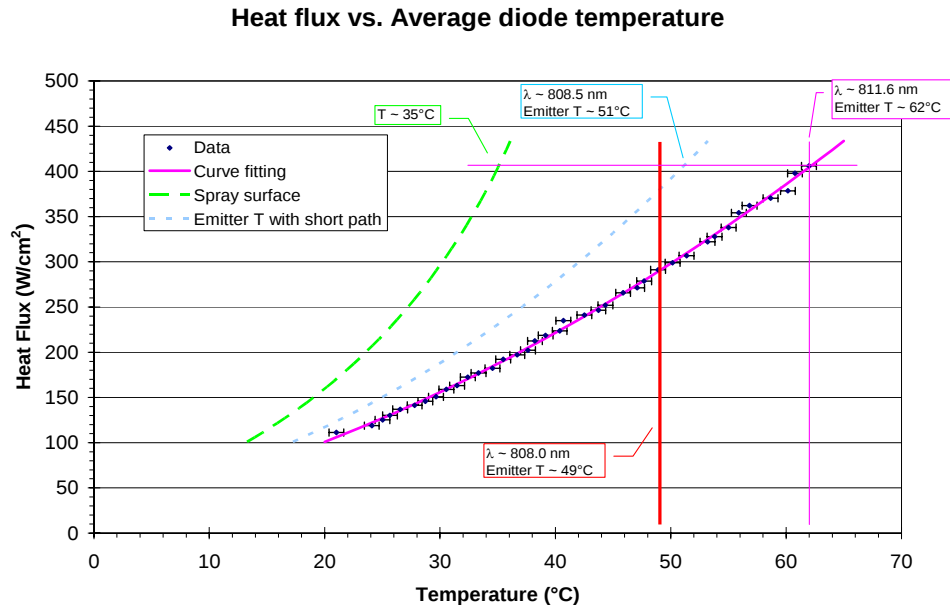


Figure 5.10: Heat flux versus average diode temperature

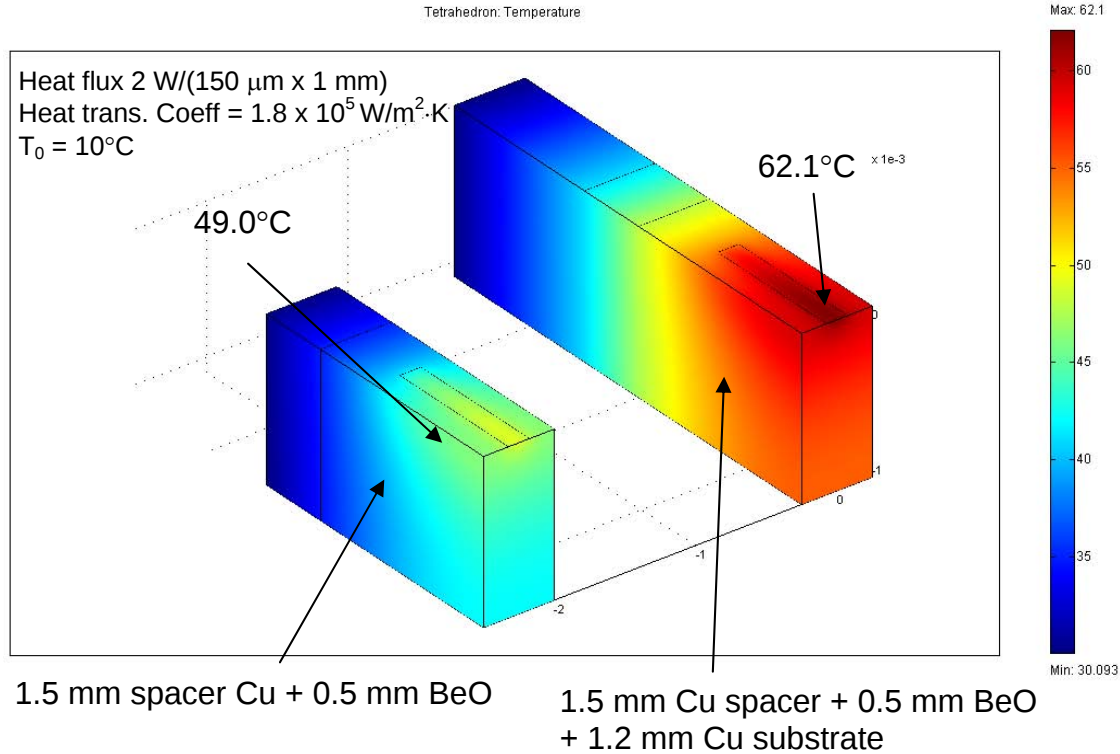


Figure 5.11: FEA result of temperature distribution of diode array stack with 2.7 mm and 1.5 mm thick copper spacers

5.2 BEAM CONTROL PRISM PACKAGING DESIGN (BCPP)

In previous section, we demonstrated that spray cooling is an effective and efficient way to cool a diode laser array generating a heat flux of 400 W/cm^2 or even higher. In a traditional stack package of diode bars there is a relatively long cooling path resulting in unwanted thermal resistance. Even though spray cooling can provide very a high heat transfer coefficient a long cooling path can keep the emitter temperature too high. Therefore, a new packaging method is needed. The most straightforward packaging arrangement is to put the p-side of the diode emitters closer to the coolant and lying parallel to the cooling surface. In such a circumstance extracting the light from the diode requires such optical devices as BCPs.

Basic design concept

An improved package design for diode laser arrays has to result in lower emitter temperatures and good optical output. We categorize BCPP design issues into four different groups. They are optical, thermal, mechanical and electrical considerations and their interactions with each other. These issues make the BCPP design complex. Optical issues are directly related to the performance of the diode array. The concepts of the BCP have been discussed thoroughly in the previous chapter. A good BCP needs to have no absorption at the diode wavelength, large acceptance angles, low aberration, small size, short back focal length, minimal positioning accuracy requirements, and a large angular tuning range.

Thermal issues focus on reducing the emitter temperature. For a good design, lower emitter temperature, faster thermal response time, and good temperature uniformity along the diode bar

are the targets to be achieved. Mechanical issues are really related to how to make the final product. These issues are very practical concerns. High packing density, ease of manufacture, simple assembly procedure, better positioning tolerance, and fewer parts are considered in improved mechanical design. Electrical issues are conceptually simple. A serial electrical connection of diode bars is preferred for the diode laser array because the voltage drop across a typical diode bar is only a few volts, but the current can be quite high, 40 A or more. If the electric connection is in parallel, this requires a very high current, but low voltage power supply which is difficult to obtain.

Beyond these four packaging issues cost is also a major consideration for a practical design. The smaller the BCP is, the higher the packing density can be. A smaller BCP, however, requires higher positioning accuracy and it's also harder to make. On the contrary, if we use a larger BCP, the part of BCP below the optical axis (see Figure 5.12) is larger. This can mean a longer cooling path in the package. We can remove some part of the BCP to solve this problem but it will require extra work and add to the cost of the package. The substrate design determines the thermal properties of the whole system. The chosen material should have high thermal conductivity and should be easily machined. Also, the p-electrode of the diode bar will be directly attached to the substrate through the p electrode and so a proper method of attachment should be chosen. Thermal resistance in the interface between substrate, electrodes and diode bar is also critical.

The thermal resistance, temperature uniformity and electronic connection are the three major issues in selecting a package design. Lower thermal resistance requires shortening the cooling path from heat source to the coolant and/or use of a high thermal conductivity material as the substrate. This low thermal resistance provides a faster thermal response time. In other words, we can have better control of the diode temperature and consequently, of the output wavelength. Good temperature uniformity gives a better wavelength distribution along the diode bar. In most cases, the end emitters in a bar have lower temperatures than the emitters at the center because they lack emitters on one side. Therefore, we may have to change the substrate design to keep the edge emitters warmer.

Since BCPs can redirect the light into the desired direction, we can put the p-side of the diode laser much closer to the coolant to achieve lower thermal resistance. In the following example, we consider using a commercially available 1 mm diameter BK7 half-rod with a high reflection coating on the flat surface as our BCP. AR coatings might be needed on the curved surfaces. Figure 5.12 shows the basic BCPP we suggest. The substrate can be either a piece of metal or an electrical insulator. If we choose a metallic substrate, another layer of insulator, e.g., GaAs, must be deposited on top of the substrate in order to enable serial electrical connection. The p-side electrode is deposited or placed directly on the insulator layer of a metallic substrate or directly on the insulator substrate. Since there will be about 40 A passing through the electrodes, heat will be generated in these electrodes. However, the p-electrode is immediately on top of the substrate and the waste heat so generated can be removed directly. Therefore, the thickness of the p-side electrode is not too critical. However, the n-side electrode has a longer cooling path. Hence, we must make the cross section of n-side electrode large to reduce Ohmic heat generation.

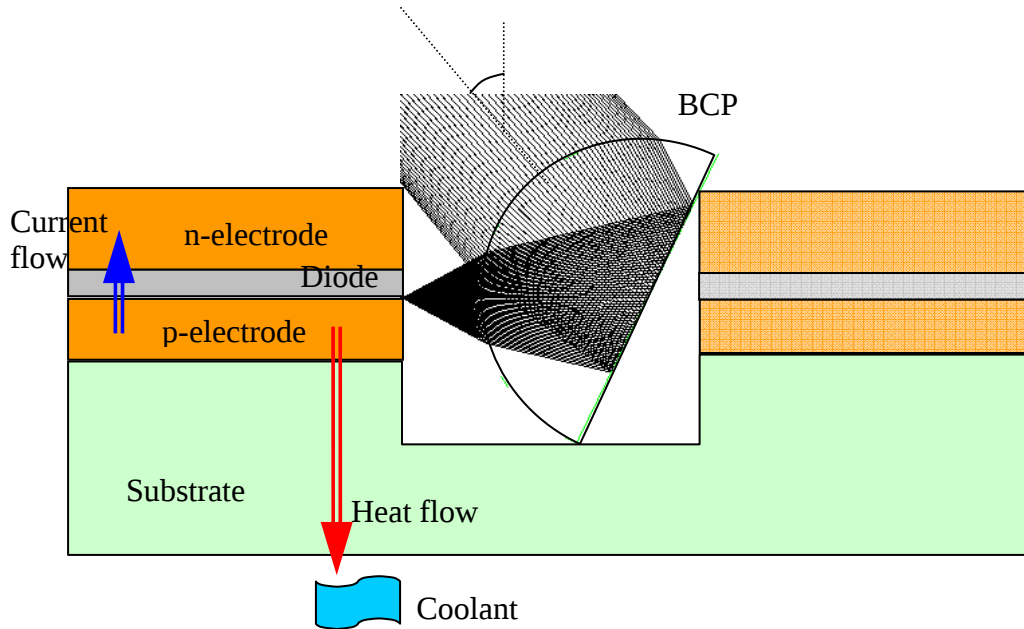


Figure 5.12: The Basic structure of the BCPP

The BCP has a finite size and the diode laser has to be on the optical axis. To reduce the emitter temperature, the p-side of the diode bar has to be placed closer to the coolant. Consequently, the diode has to be lifted higher or the BCP has to sit in a groove. The larger the BCP, the further will be the vertical distance between the p-side and the lower part of BCP. In other words, assuming the bottom of the substrate is flat, the distance between the p-side of the diode and the coolant will be larger if the BCP is large.

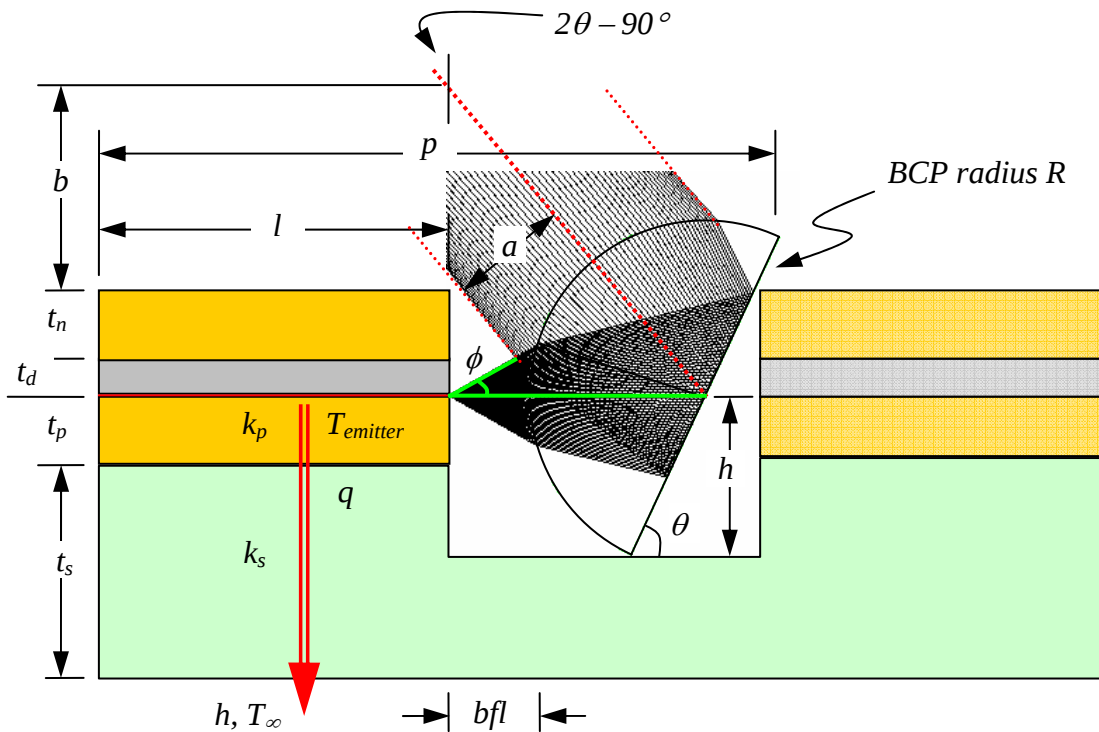


Figure 5.13: Parameters of the BCPP

Optical issues

Figure 5.13 defines several parameters for the BCPP design. We can write down several equations to represent the relationship among these parameters.

1. Channel depth equation

$$\sin \theta = \frac{h}{R} \quad (5.7)$$

2. Pitch equation

$$p = l + bfl + R + R \cos \theta \quad (5.8)$$

3. N-electrode non-blocking condition

$$b + t_n + t_d = (bfl + R) \cdot \tan 2\theta \quad (5.9)$$

$$b \cos 2\theta < a \quad (5.10)$$

4. Output beam radius equation

$$a(R, n, \phi) \quad (5.11)$$

Equation (5.11) cannot be represented in a simple form. For a typical 60° full divergence angle, we can use numerical methods to obtain the beam radius a versus the folded-ball BCP index of refraction as shown in Figure 5.14 for the case where the light source is placed at the paraxial focal point. A larger divergence angle does not have a solution when the position of light source is farther than this because the marginal rays might not enter the BCP. The folded-ball radius is set as 0.5 mm in this figure. For convenience and simplicity, we can consider a , the output beam radius defined in Figure 5.13, is 0.3 mm for the following calculation.

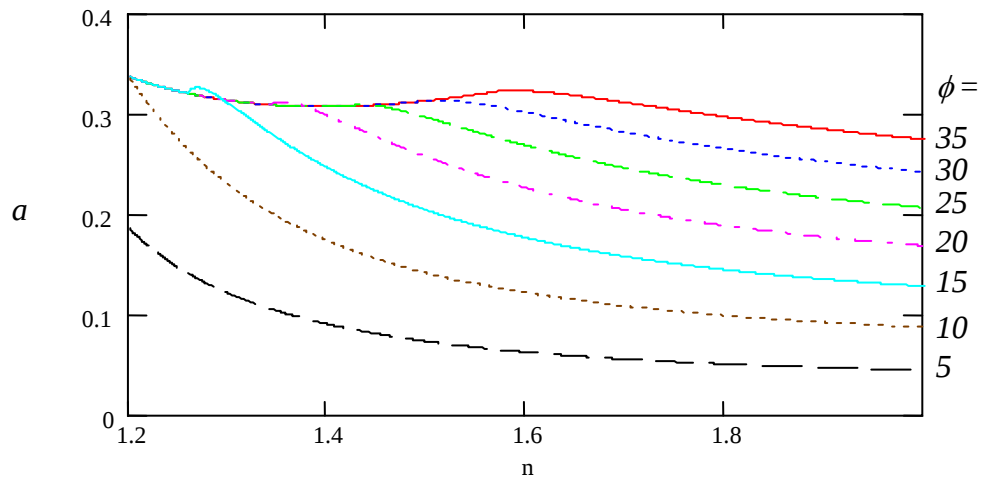


Figure 5.14: Output beam radius, a , versus BCP index of refraction for different ϕ angles

5. BCP focal length

Equation (5.12) is the paraxial focal length of a cylindrical lens.

$$f = \frac{nR}{n-1} \quad (5.12)$$

$$bfl = f \frac{2-n}{n} = \frac{2-n}{n-1} \cdot R \quad (5.13)$$

Equations (5.7 – 5.13) provide the guidelines for the BCPP design. A preferred BCPP should have small p , $T_{emitter}$, and satisfy equation (5.10). If we consider a beam angle of 90° or the BCP tilt angle $\theta = 45^\circ$, Equation (5.10) will automatically be satisfied. Obviously, we should choose smaller BCPs to reduce the pitch and the cooling path. However, a smaller BCP will require higher positioning precision. Also, a smaller BCP is much harder to make. In addition, if a smaller BCP is used, the angular tuning range will be limited by Equation (5.10).

Electric connection

Figure 5.13 shows the basic design concept and the structure of the BCPP without considering the electrical connection. As discussed above serial electrical connection is preferred. Several ways have been proposed to achieve the desired series connection. There is also another concern about the current directions on the electrodes as shown in Figure 5.15. We call a design in which the current direction in the p-electrode is the same as in the n-electrode a case of parallel current. If the direction is opposite we call the case one of anti-parallel current.

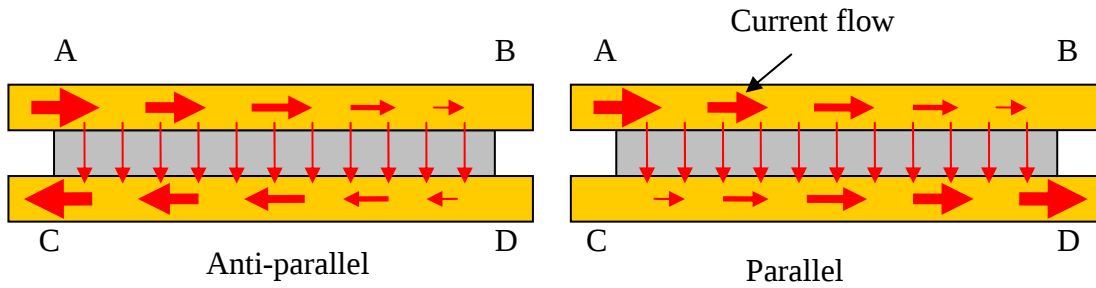


Figure 5.15: Current directions in the p- and n-electrodes

The current direction will influence the temperature distribution in the diode bar because of Ohmic heating in the electrodes. Under anti-parallel conditions, Sides A and C have higher current densities, higher ohm heating, and higher temperature compared with sides B and D (See Figure 5.15). Consequently, the parallel current case is preferred. Figure 5.16 shows a large array with patterned p-electrodes, known as a “vertical” connection. This design provides a simple way to connect the diode bars and extend the size of the array. However, the current direction in the p-electrode and the n-electrode are anti-parallel. Figure 5.17 is the modified version of Figure 5.16 and known as the “horizontal” connection.

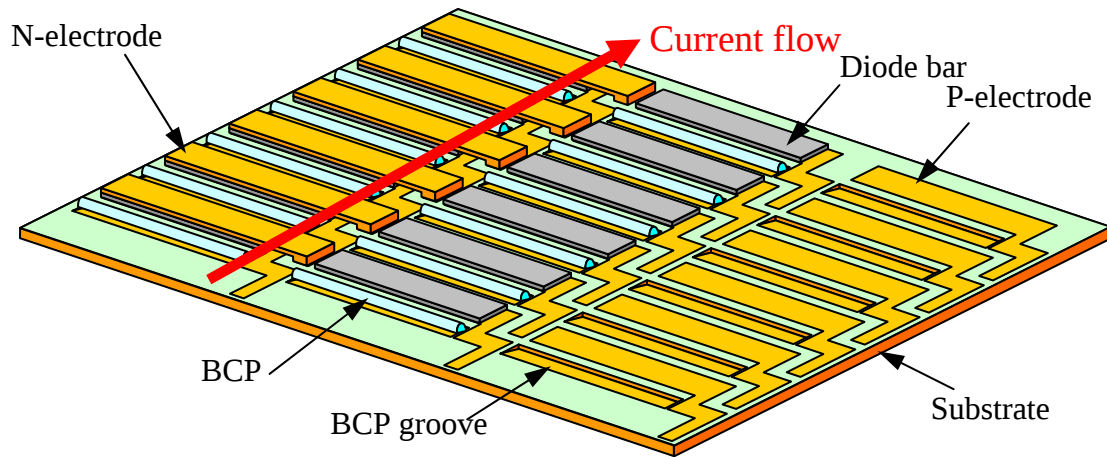


Figure 5.16: Packing example for larger array with “vertical” electrical connections

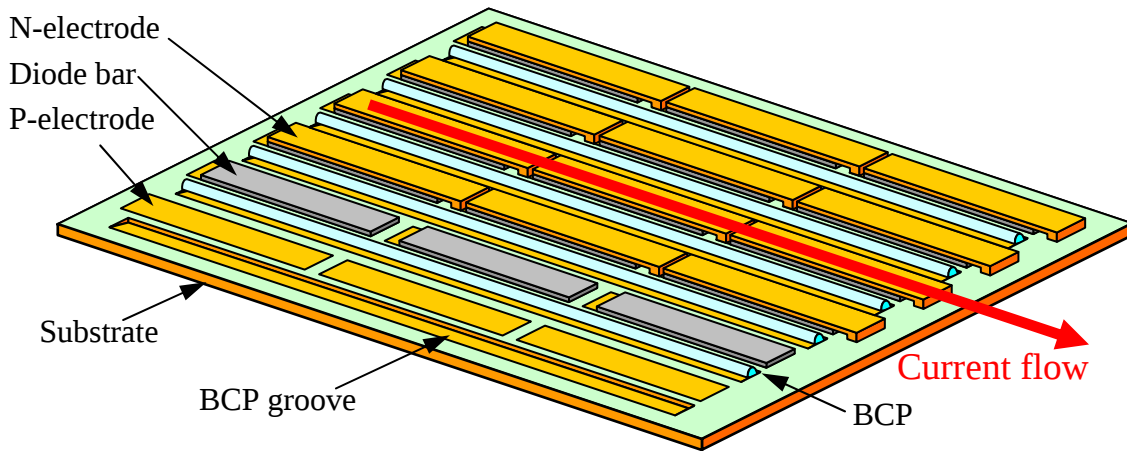


Figure 5.17: Packing example for larger array with “horizontal” electrical connections

In the case shown in Figure 5.17 the electrode shapes are relatively simple. The spacing between bars can be shortened compared to that required in Figure 5.16. This design can also use long BCPs which allow several diode bars to share a single BCP. Also, the current direction is parallel. However, as earlier Equations show, a longer BCP requires a higher angular positioning precision. Also, between each row, a serial connection is harder to achieve. Therefore, this design is not suitable for a small size array which only contains a few diode bars. Another approach is shown in Figure 5.18. It is also the design used in our experiments. This design uses the n-electrode to achieve the serial connection. Its n-electrode is complicated and needs two different designs which are mirror images of each other. The higher end of the n-electrode ensures it doesn't short circuit. Essentially, this design is also a modification of Figure 5.16. It is suitable for any size of the array. Also, the current directions on the p-electrode and the n-electrode are parallel.

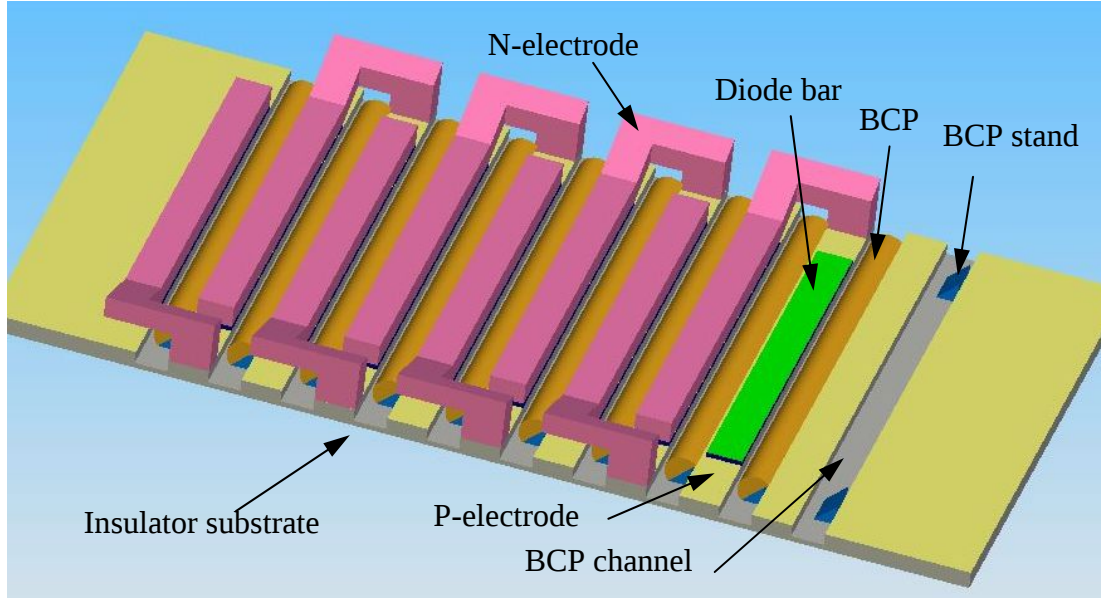


Figure 5.18: BCPP with shaped n-electrodes

Thermal issues

Thermal issues are the major concern of BCPP design. 1D temperature distribution can be computed using Fourier's law and Newton's law of cooling. The system as shown in Figure 5.14 can be described as having thermal resistance given by

$$R_{th} = \frac{1}{h \cdot A} + \sum_i \frac{t_i}{k_i \cdot A} = \frac{1}{h \cdot A} + \frac{t_p}{k_p \cdot A} + \frac{t_s}{k_s \cdot A} \quad (5.14)$$

$$T_{emitter} = T_{\infty} + R_{th} \cdot q \quad (5.15)$$

k are the thermal conductivities of the materials. h is the heat transfer coefficient between the coolant and the contacting surface. A is the area of the interface. t_i is the thickness of the material and q is the heat flow of the heat source. We know that choosing materials with high thermal conductivities and using a very thin substrate will reduce the emitter temperatures. Also, a higher heat transfer coefficient cooling method should be employed. For the conductive p-electrode, the most convenient material with high electric and thermal conductivity is copper. An electrically insulating material with a high thermal conductivity is not common. The most desirable insulating material is diamond. It has the unbeatable thermal conductivity of 2300 W/m·K. However, its hardness and availability present major problems. Also, attaching copper or some other material on the diamond is not an easy task. Another commonly used material is BeO which has the thermal conductivity of 272 W/m·K. It is easier to obtain and all of the peripheral techniques for handling BeO are well-developed. The direct bond copper is a mature technique to adhere the patterned copper film on top of the BeO substrate. Also, BeO is relatively easy to machine. The only draw back is its toxicity, but many companies have facilities to safely machine this material.

A diode emitter generates waste heat of about 2 W, and the emitter area is about 10^{-7} m^2 . Therefore, an effective cooling method is necessary. In our experiment, we use evaporated low pressure water spray cooling for our test diode array as mentioned before. When dealing with the real system, we should consider the 3D case and use the heat diffusion equation:

$$k\nabla^2 T + \dot{q} = \rho \cdot c_p \cdot \frac{\partial T}{\partial t} \quad (5.16)$$

For a complicate geometry like the BCPP, the only way to solve the problem is through numerical methods. We used FemLab®, a finite element analysis software program, to simulate the temperature distribution of the BCPP.

Mechanical issues

Mechanical issues are the most practical and trivial part of the design. The thermal expansion coefficients of the conduction layer and the insulator layer are usually different. This mismatch might cause thermal stress or even distort the substrate. Also, the total thickness of the substrate needs to be strong enough to sustain the pressure difference due to the active cooling. For example, the low pressure water, evaporative spray cooling we used must sustain a pressure difference of about 1 atm between the diode array and the spray cooling chamber. Keeping the BCP precisely positioned is also an important mechanical issue. Some calculation shows the translational and the angular position tolerances. In a practical design a passive means to align the BCP is essential.

Experimental design

We consider the folded-ball BCP design with BCPs made of BK7 glass ($n=1.51$ at 808 nm) with radius of curvature 0.5 mm. We chose the folded ball BCP because it is simple and inexpensive to prepare. The output beam angle was selected as 90° and so the BCP tilt angle is 45° . Assume the beam angle tolerance $\Delta\Theta$ is equal to 3° and we can obtain the translational and the angular positioning tolerance as

Table 5.2: BCP positioning tolerance

$\Delta X =$	67 μm
$\Delta Z =$	43 μm
$\Delta\alpha =$	4.3 mrad
$\Delta\beta =$	26 mrad
$\Delta\gamma =$	6.7 mrad

The substrate design was most critical. A BeO substrate with directly bonded copper to serve as the p-electrode was chosen because of it is inexpensive to produce using relatively mature methods. Copper films can be directly bonded to a BeO surface through a high temperature reaction. However, the expansion coefficients of the copper and BeO are quite different. This difference induces stress in the substrate when the substrate cools down to room temperature. Therefore, an additional direct bond copper film needs to be added on the opposite side of the BeO substrate to balance the stress. Hence, the final substrate with the p-electrode directly bonded to the BeO became a copper-BeO-copper sandwich. A diamond saw was used to cut through the top n-electrode copper layer to achieve electrical isolation between the copper stripes serving as the p-electrodes. Figure 5.19 shows the top and side view of the substrate.

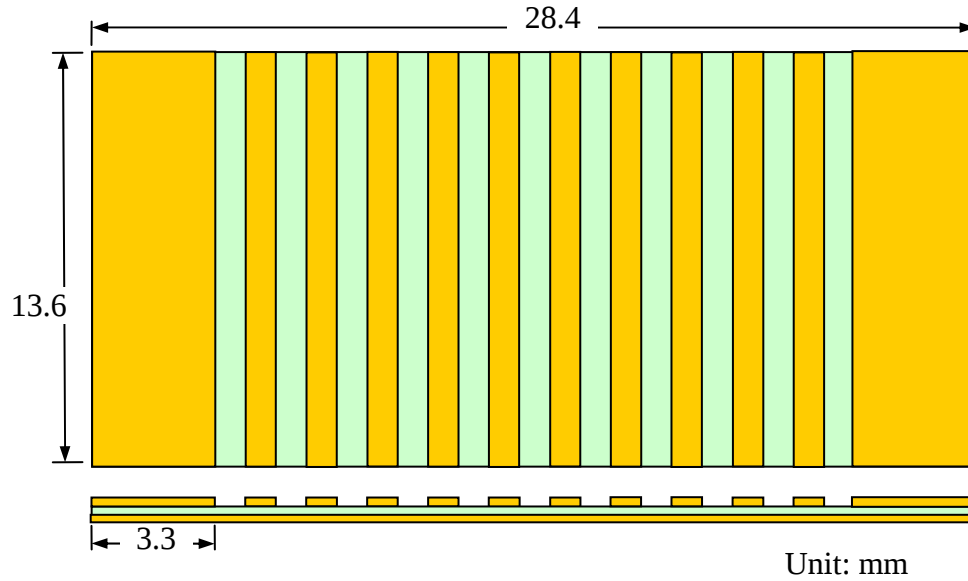


Figure 5.19: Sandwiched copper-BeO-copper BCPP substrate design

There are two methods to fix the BCP in place. One is using a BCP stand and the other is using dual BCPs design. The BCP stand is a triangular piece of glass or any material that is sturdy, stable and easy to machine. Figure 5.20 shows the detailed dimensions of the BCPP substrate employing the BCP stand concept.

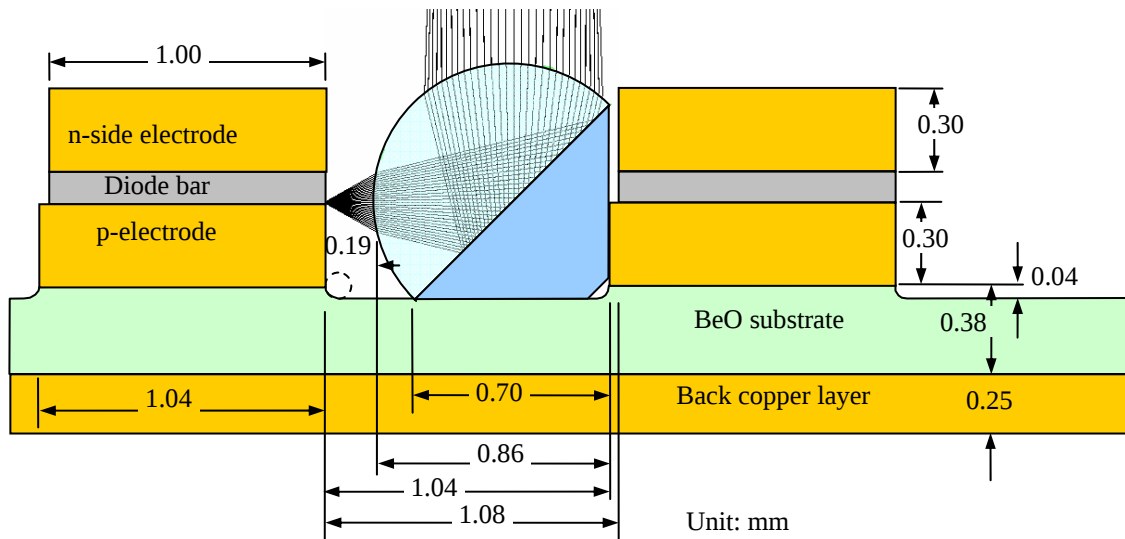


Figure 5.20: Detail dimensions of BCPP substrate employing BCP stand concept.

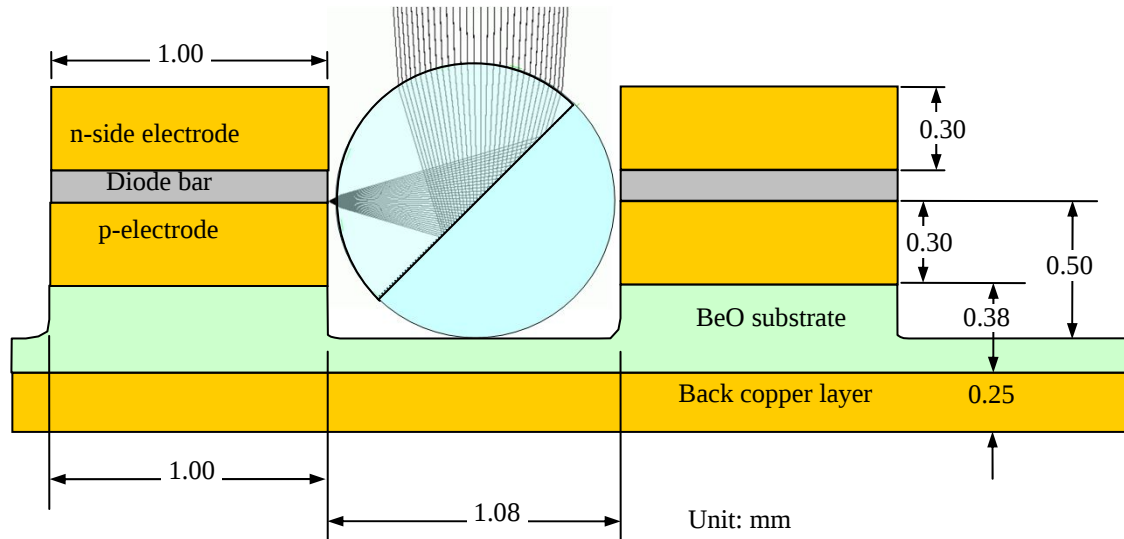
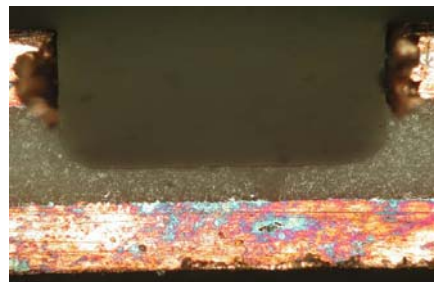
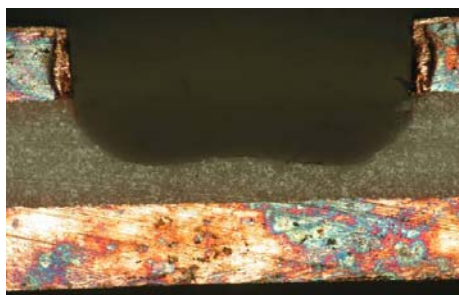


Figure 5.21: Detail dimensions of BCPP substrate employing a dual BCP

A BCP stand design gives a fixed beam angle and a more compact design, and this design has the potential to achieve a shorter cooling distance between the diode and the coolant. The dual BCP or the glued-BCP design is accomplished by gluing the BCP to another piece of a half rod which has the same dimension of the BCP. This makes the whole piece become a complete rod. Figure 5.21 shows the detailed dimensions of the glued-BCP design. Because of the symmetry, this feature gives more freedom for angle tuning as long as the HR coating on the flat surface remains reflective. Because the rod-shaped BCP has only one point in contact with the bottom of the BCP channel the flatness of the channel is not as critical as in the design employing a BCP stand. Currently, it is difficult to use a diamond saw to produce a flat bottomed cut for the BCP channel. Figure 5.22 shows microscope photos of several typical cross sectional views of the BCP channels produced using different size diamond saws.



(a) An ideal BCP channel



(b) 1 mm wide saw cutting result



(c) 0.3 mm wide saw cutting result

Figure 5.22: BCP channel side views of different cutting methods

Figure 5.22 (a) shows a cross section of a cut made using a new 1 mm thick diamond saw. The edge and the bottom of the channel are flat and smooth. However, a diamond saw tends to wear out and makes a channel with cross section as shown in Figure 5.22 (b). Using a narrower diamond saw will leave several lumps at the bottom of the channel as shown in Figure 5.22 (c). The lumpy bottom might make the BCP stand have an undesirable tilt. However the dual BCP design is quite insensitive to the uneven channel bottom. As long as the contact points between the dual BCP and the channel bottom have the same height, the beam output remains the same. Since copper is a relatively soft material for machining, the diamond saw pushes and distorts the copper edges of the channel when cutting. Such distortion causes all p-electrode surfaces to become concave.

Figure 5.23 shows the surface profile of a p-electrode after the cutting process. The height difference between the center and the edge of the p-electrode can be as large as 12 μm . An uneven surface potentially can cause a solder void between the electrode and the diode bar and/or an extra thickness of solder. Consequently, all surfaces were polished after the cutting process.

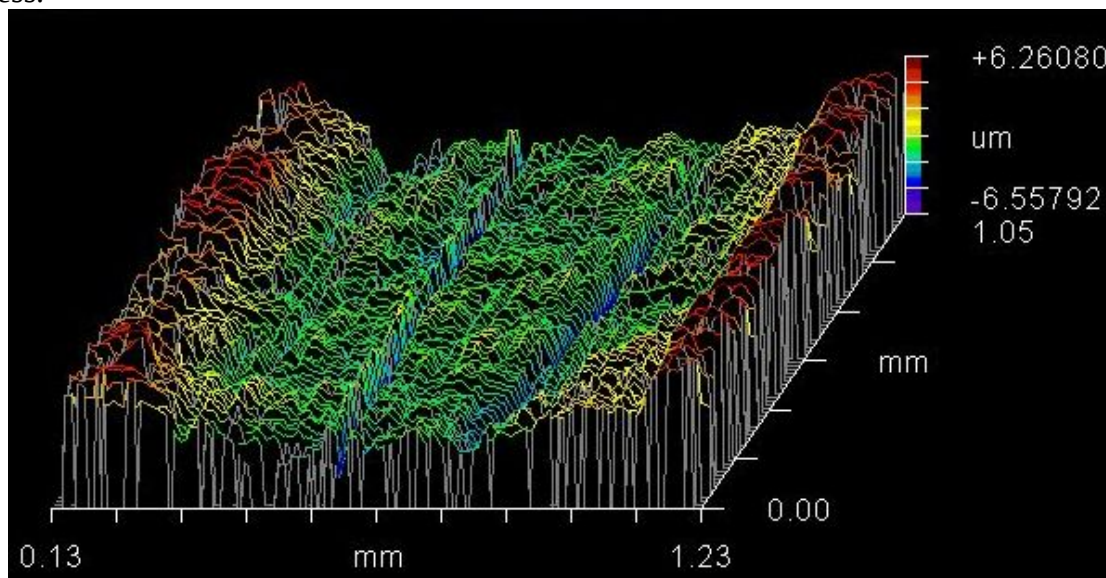


Figure 5.23: The n-electrode surface plot after cutting the channel

The N-electrodes have two designs because of the alternating connection directions along the array as shown in Figure 5.18. Figure 5.24 shows one of the designs. The two designs are mutual mirror image of each other. The end structure of the n-electrode ensures the electrodes don't short each other electrically.

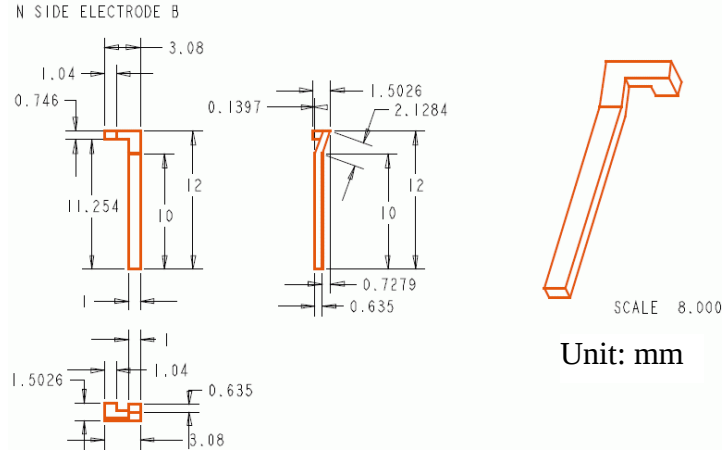


Figure 5.24: The n-electrode design

The assembly of the studied array was carried out by Decade Optical Systems. The first step was to assemble the diode bars and the n-electrodes on top of the substrate. Thin film solder sheets were inserted between the diode bars, p-electrodes and n-electrodes. The whole system was placed on a special mount and placed in an oven that melts the solder films and completes the soldering process. We did not cool the array directly through the back copper layer of the substrate. Instead the array was soldered to a heat exchanger which can be a traditional forced convection heat exchanger or a spray cooling chamber. The soldering step is similar to soldering the diode bars and the n-electrodes. In our experiment, we used low pressure water spray cooling as the cooling method. Its heat transfer coefficient was about $1.5\text{-}2.0 \times 10^5 \text{ W/m}^2\text{K}$ as discussed in Chapter 0. However, an extra 1 mm thick copper layer was necessary to provide mechanical stiffness and to sustain the pressure difference between the spray cooling chamber and the external environment. This extra copper layer increased the resulting emitter temperature slightly.

Placing a BCP in its channel in place is a major issue. For the BCP stand design, the positioning relies on the precision of the parts. It will only have a very limited range of fine tuning before curing the epoxy used to hold the BCP to the stand. On the contrary, the dual BCP design has a relatively large tuning range and the dual BCP is very easy to tune by rotating the rod slightly. This fine tuning procedure can be done two ways. One way is to drive the diode bars in the LED mode and the output light projected to a screen to monitor the beam angles while adjusting them as desired. The other way is to look directly into the BCPs under a microscope. If the BCP is well-aligned, the observer should see the image of the p-electrode and the diode bar p-edge at the center of the BCP. After the fine tuning process the BCPs were secured by UV curing the epoxy glue used to hold the BCP.

BCPP temperature distribution FEA results

We used FEMLAB finite element analysis to calculate the temperature distribution. Our calculation was based on the numerical values of dimensions and material parameters given in the diode bar manufacturer's data sheets Lasertel LT1200-40W and LT1300-60W shown in Table 5.3. The simplest way to calculate the temperature distribution of the package was to set the heat source exactly on top of the p-electrode. Because the first BCPP diode array utilized micro channel cooling, a constant temperature boundary condition was used at the bottom surfaces in the following simulations. All other surfaces were selected as insulating surfaces.

Figure 5.25 shows that the single emitter temperature would be about 38°C. Figure 5.26 shows the temperature distribution along a diode bar. The edge emitters have relatively lower temperatures compared to the emitters in the middle of the bar. A similar effect can be observed when an entire BCPP array FEA is calculated. The two end-most bars show slightly lower temperatures compared to the central bars. Figure 5.27 shows the temperatures of different bars in the array. Since there are channels between bars the data points are discrete. As presented in Figure 5.28, the edge emitters' temperatures are about 8°C lower than that of the central emitters when the waste heat is 60W. This is because the edge emitters lack heat sources to keep them warm. Similarly, the edge effect makes the edge bars about 1.5°C cooler than the central bars as shown in Figure 5.29.

Table 5.3: LaserTel LT1200-40W and 1300-60W diode laser specifications

LaserTel diode laser bars		LT1200-40W*	LT1300-60W**	Unit
Optical Characteristics	Output power	40	60	W
	Center wavelength	808	808	nm
	Center wavelength Tolerance	±5	±5	nm
	Wavelength temperature coefficient	0.3	0.3	nm/°C
	Spectral Width (FWHM)	2.5	2.5	nm
	Array Length	10	10	mm
	Number of Emitters	19	24	
	Emitter Size	150 × 1	200 × 1	μm
	Emitter Spacing (center-to-center)	500	200	μm
	Slow Axis Divergence (FWHM) SA	10°	10°	Degrees
	Fast Axis Divergence (FWHM) FA	35°	35°	Degrees
	Polarization	TE or TM	TE or TM	
Electrical Characteristics	Slope Efficiency	1.1	1.16	W/A
	Conversion Efficiency	50	49.19	%
	Pulse Width			ms
	Duty Cycle			%
	Threshold Current	8	12.35	A
	Operating Current	43	64.2	A
	Operating Voltage	1.9	1.9	V
Thermal Characteristics	Series Resistance	0.004	0.00349	Ω
	Thermal Resistance	-	-	°C/W
	Recommended Case Temperature			°C
	Operating Temperature Range			°C
	Storage Temperature Range	-40 to 85	-40 to 85	°C

*http://www.lasertel.com/media/LT1200_40W.pdf
**Datasheet from Lasertel for LT1300 S/N AAB2639

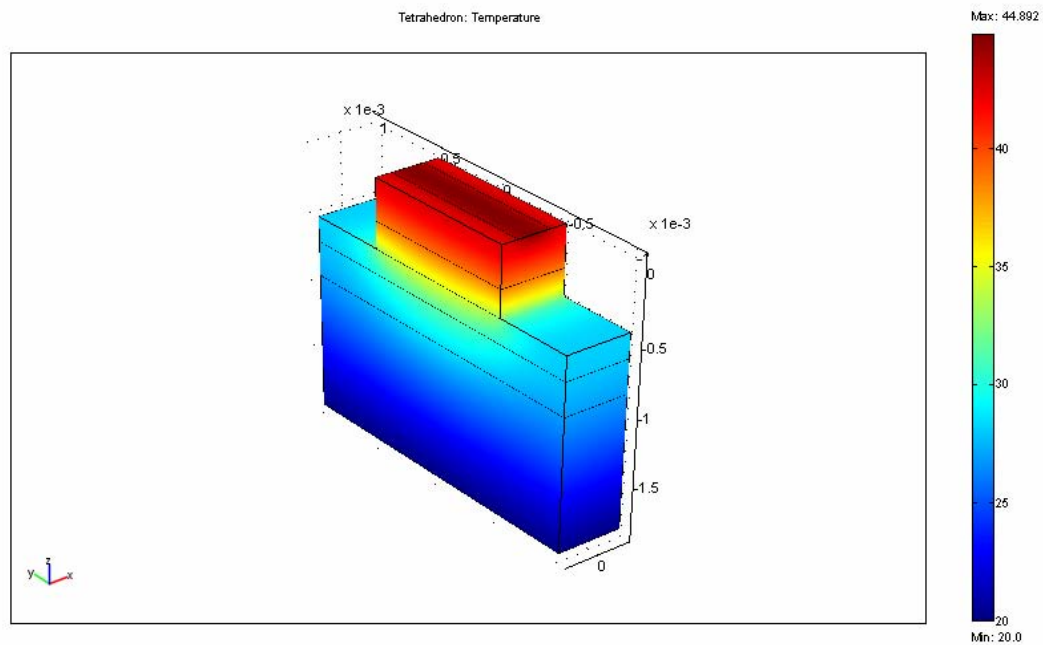


Figure 5.25: FEA computed temperature distribution of a single emitter in a 60W bar

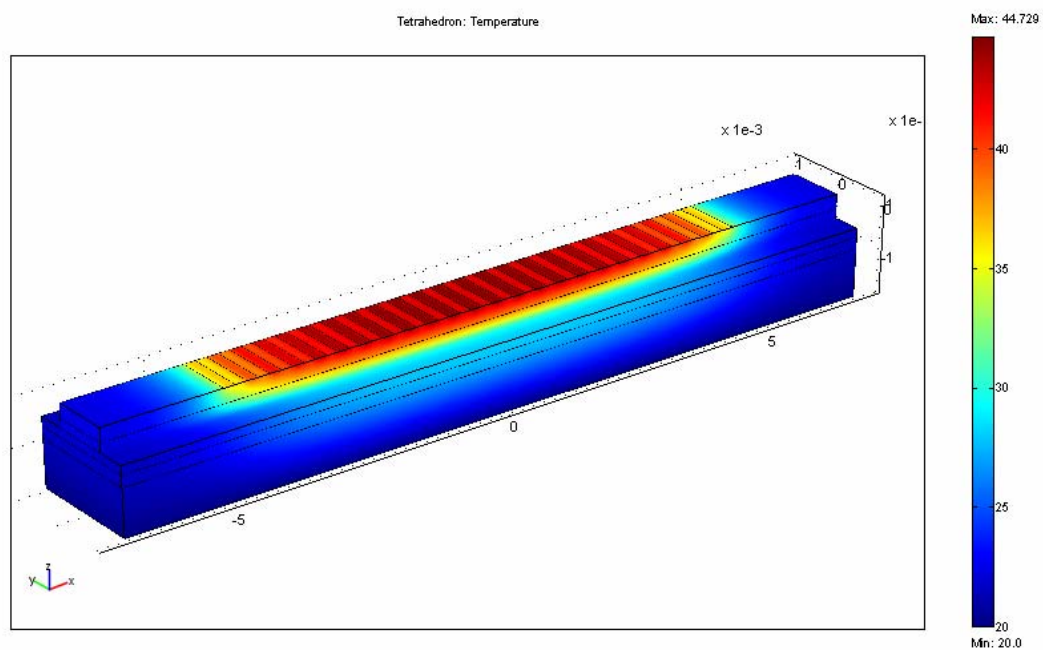


Figure 5.26: FEA computed temperature distribution of a single 60 W bar

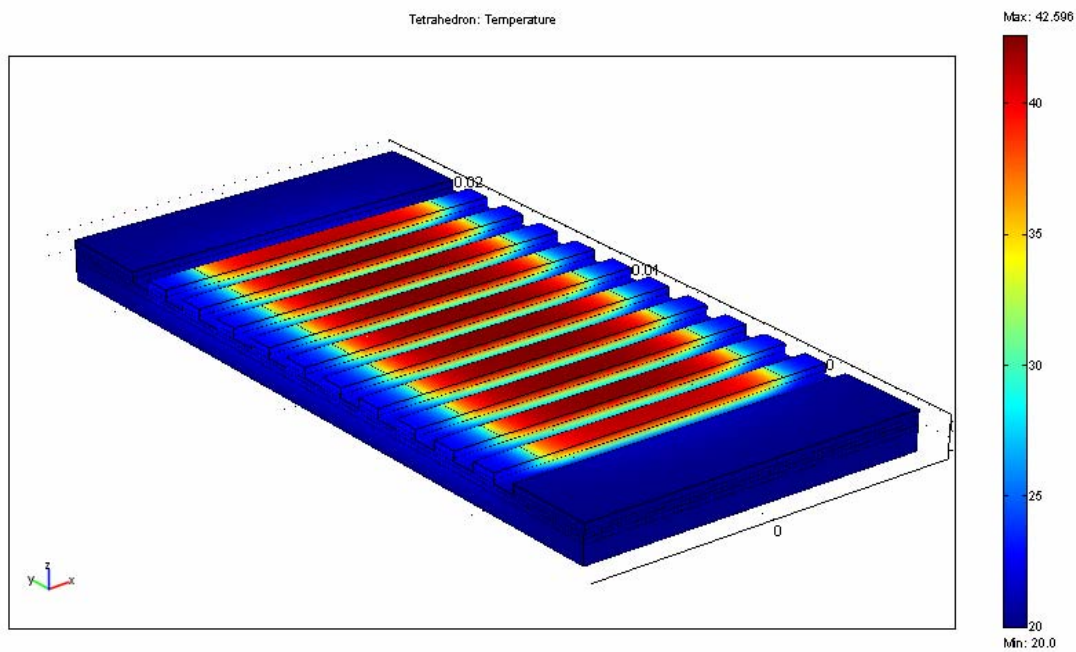


Figure 5.27: FEA computed temperature distribution of an entire BCPP array built with of 10 60W bars

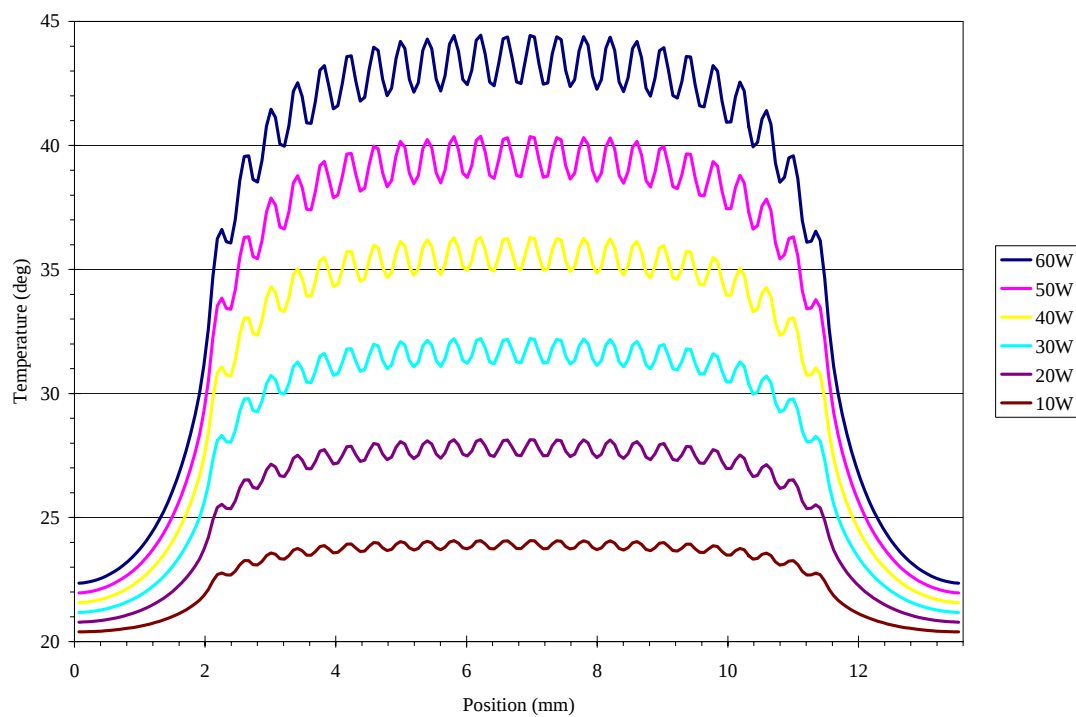


Figure 5.28: FEA computed temperature distribution of a single diode bar for different waste heat levels

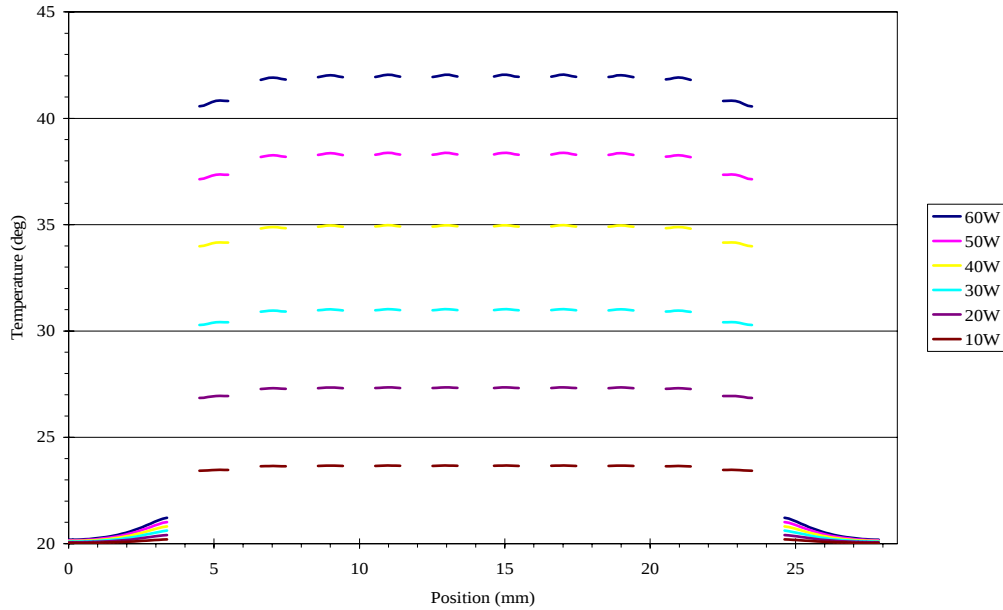


Figure 5.29: FEA computed temperature distribution in the emitter plane of the BCPP array for different waste heat levels

Experimental results

Figure 5.30 shows the top view of the BCPP array. This array is assembled by Decade Optical System, and is the very first BCPP array ever made. The BCPs are not aligned properly; hence, the laser output of each bar is directed in different directions, yet it shows the ability to aim all the beams to a certain spot if the BCPs tilt angles had been controlled properly when the array was assembled.

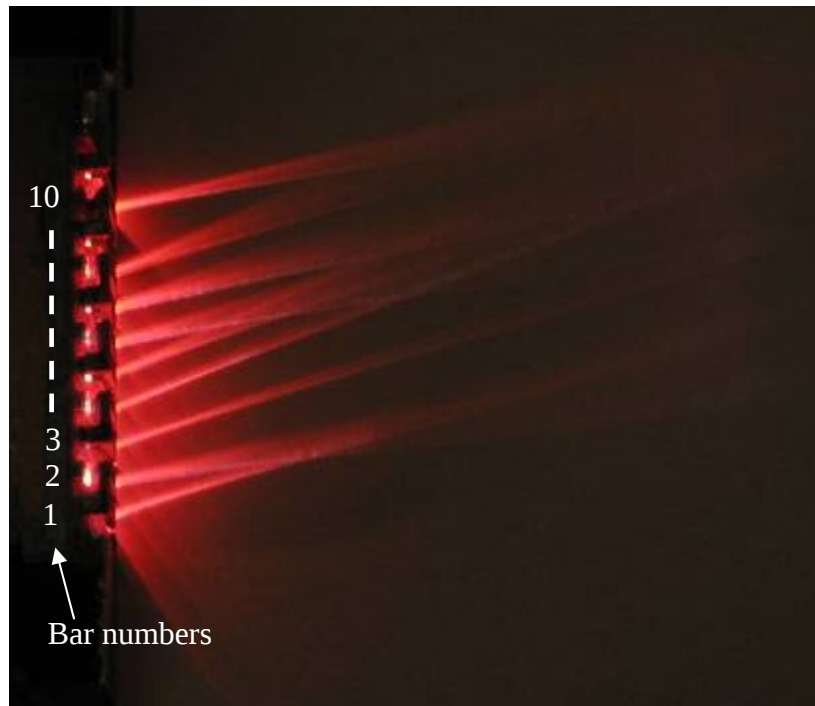


Figure 5.30: Top view of the BCPP array

Figure 5.31 shows the front view of the BCPP array and the numbering system for the bars and emitters. Since most of the beams are directed to oblique angles, the camera can take pictures at an angle to avoid the laser beams. The picture shows that bar 9 failed. Emitters 2 and 3 on bar 4 were not emitting. In bar 8, emitter 11 did not emit and emitter 12 has obviously lower strength than other emitters on this bar.

There is no good and direct way to measure the emitter temperatures because the emitters cannot be directly observed. Even though an IR image cannot provide a precise temperature measurement, it still provides a qualitative temperature distribution of the BCPP array. Figure **Figure** 5.32 show the IR image of the BCPP array. The image is taken with a FLIR PM290 IR camera which is placed at an angle to avoid direct illumination of the camera sensor by the laser beams. Comparing with Figure 5.31 and Figure 5.32, a hot spot at bar 9 implies a short circuit which makes the bar fail. The positions of the failed emitters on bar 4 and bar 8 are indicated by green arrows. Both emitters show higher temperature than the other emitters implying that these two still have similar resistance compared with other emitters, but that the power passing through them is mostly converted into heat instead of light. The IR image also shows the edge effect of the bars and the array. Bar 3 and bar 4 showed higher temperatures roughly intermediate between emitter 2 and 8. There are three possible reasons that can cause this effect. One possibility is that the solder layers between the diode bar and electrodes are in poor contact in the corresponding regions. This allows more current to go through and increases the local temperature. Another possibility is the solder layer between the common substrate and the micro channel heat exchanger is thicker at this part. A thicker solder layer implies a longer cooling path and higher thermal resistance which increases the temperature of the emitters at this region. The other possibility is that the hot spot on bar # 4 generates extra heat and heats the neighboring area to a higher temperature.

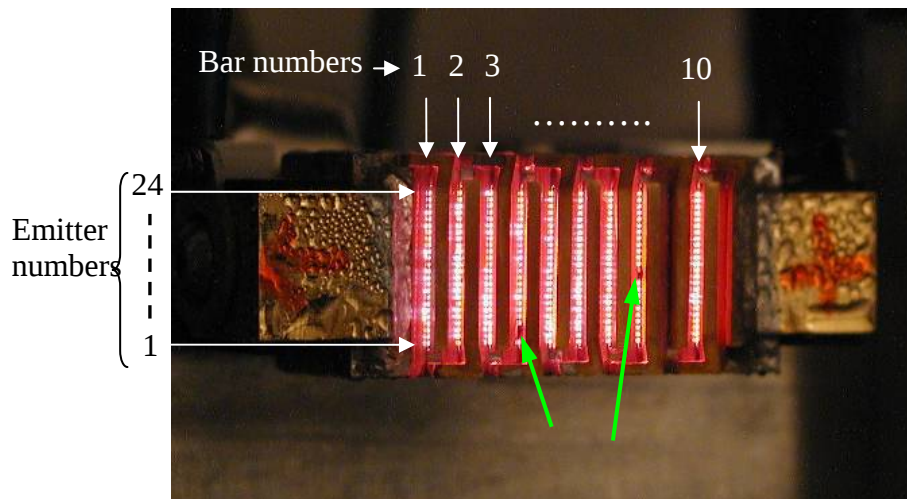


Figure 5.31: Front view of the BCPP array

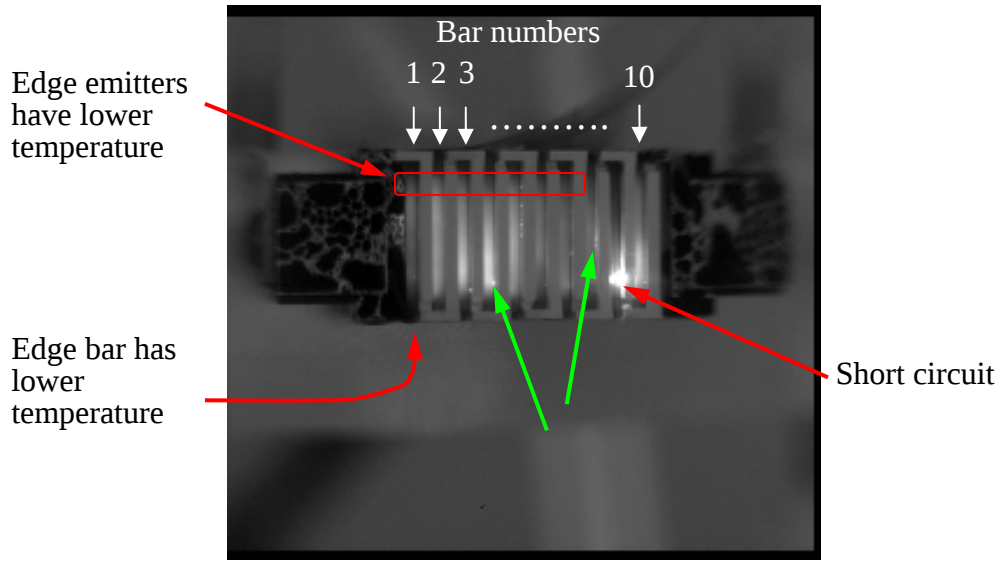


Figure 5.32: IR image of the BCPP array

A quick check of the total output spectrum was performed by using an Ocean Optics HR2000 spectrometer. The probe of the spectrometer was placed about 1 m away from the diode array to make sure the output beams were properly mixed so that the probe could receive light from most of the emitters. The experimental result is shown in Figure 5.33. When the total current is about 25 A, the peak wavelength is about 808 nm, and the FWHM is about 4 nm. However, the spectra show some long wavelength components which imply some emitters might be working at a temperature about 32°C higher than the average emitter temperature.

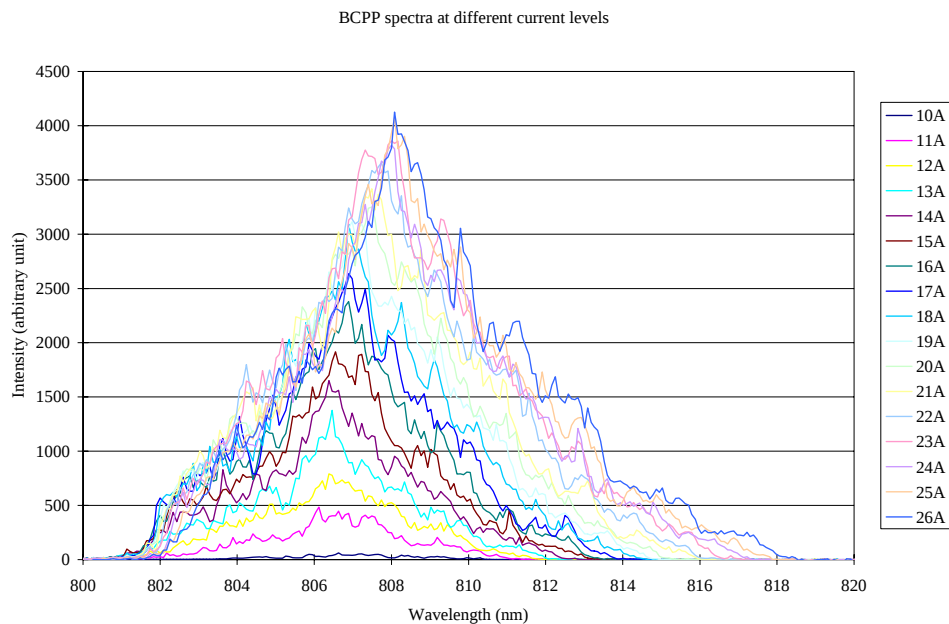


Figure 5.33: Output spectrum of the BCPP diode laser array for different current levels

Figure 5.34 shows the P-I curve and the efficiency of this BCPP array. The result shows the typical diode laser output behaviors. However, we suspected some emitters were working at an unusual high temperature as indicated by the spectral outputs shown in Figure 5.33. Hence, the maximum current applied to this array was 26 A.

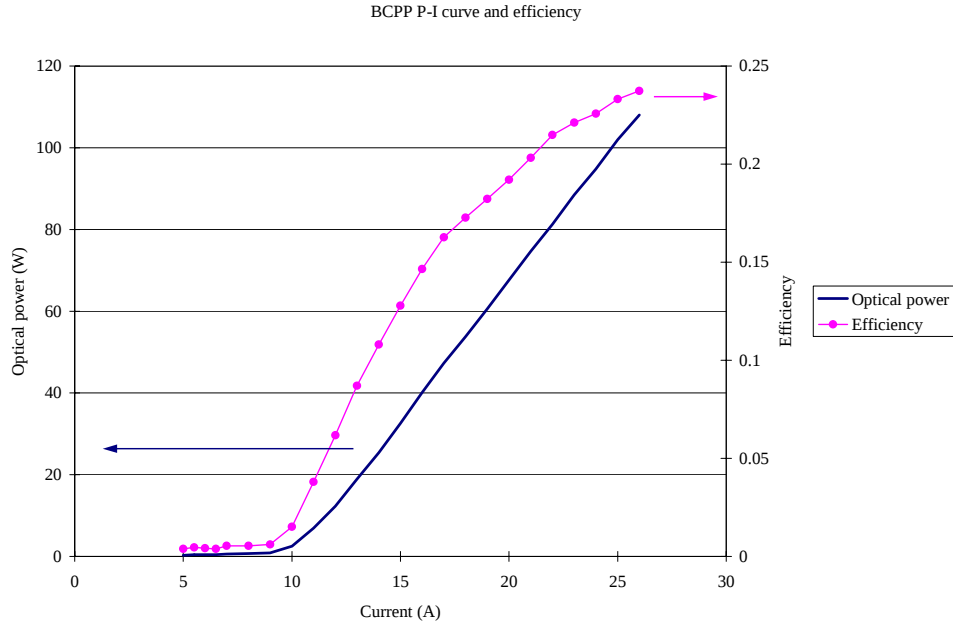


Figure 5.34: P-I curve and efficiency of the BCPP array

The emitter temperature can be calculated from its output spectrum and the temperature coefficient given in Table 5.3. By measuring the central wavelength of each emitter we can obtain the temperature distribution on the whole diode array. The wavelength measurement experiment is sketched in Figure 5.35. The BCPP diode array was mounted on an X-Y stage. A 45° tilted microscope slide was used to redirect a small portion of the output beam horizontally. An imaging lens projected images of individual emitters on the image plane. We then used the Ocean Optics HR2000 spectrometer to measure the spectrum of each emitter. Because the output power of the array was so high, the reflected laser light from the first microscope slide was still strong enough to saturate and even damage the CCD of the spectrometer. In order to reduce the light intensity, another microscope slide was placed before the image plane. By moving the BCPP diode array position we can select which emitter image is projected on the fiber probe. In this manner all of the emitters' output spectra can be measured. Hence, the temperature distribution of the array could be calculated. Figure 5.36 shows the calculated temperature distribution of bars 1, 2, 3, 5, 6 and 7 when the applied current was 25 A where the trend of this calculated result is seen to be consistent with the IR image result. The edge effect is clear. The average temperature difference between the central emitters and the edge emitters is about 11.2°C while the simulation in Figure 5.28 shows a temperature difference about 8°C for the ideal case. Bar 5 (the central bar) and bar 1 (an edge bar) differ by about 14.5°C. The temperature of bar in the region of emitter 2 to 8 was higher than the average.

The experimental results show that this first attempt at packaging an array using BCPs is not as good as the simulation suggests. This is caused by the fact that the current assembly procedure

requires that the whole array be heated up to 200°C for soldering. The thermal expansion coefficient mismatch of the copper and BeO layers makes the substrate bend toward to the front surface, or the side with BCP channels. This deformation of the substrate makes the solder layer at the center of the diode bar thicker. Also the average solder layer thickness is thicker than the ideal case. Solder layers always cause an extra thermal resistance because the thermal conductivity of the solder is lower than copper and BeO. Consequently, the temperature difference between the central emitters and the edge emitters on the same bar is higher than the simulation suggests. Moreover, because of the bending, the overall solder thickness for each single bar is also thicker in order to compensate the bending. The failed bar and emitters might also be related to improper soldering.

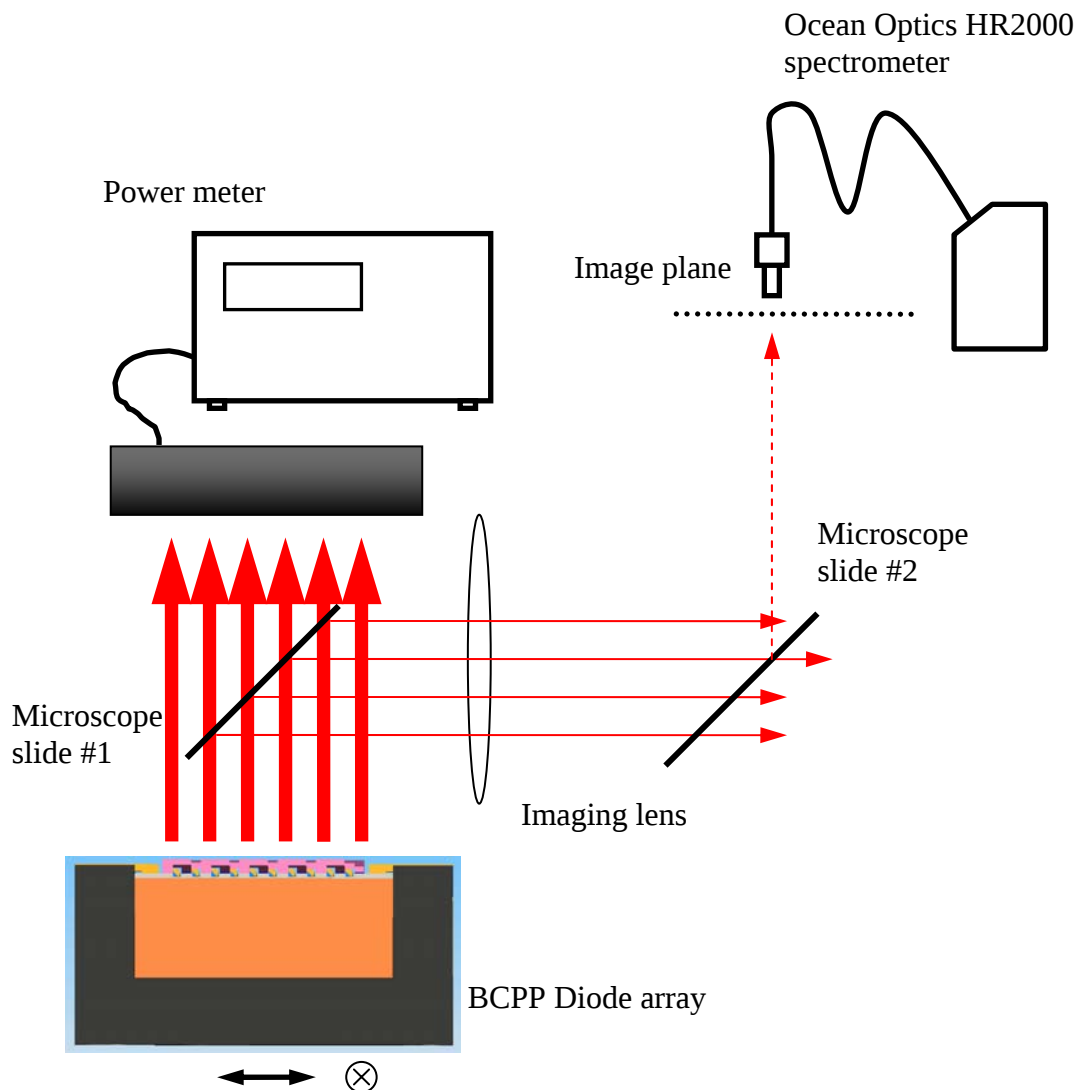


Figure 5.35: Experimental setup for measuring the output spectrum of each emitter

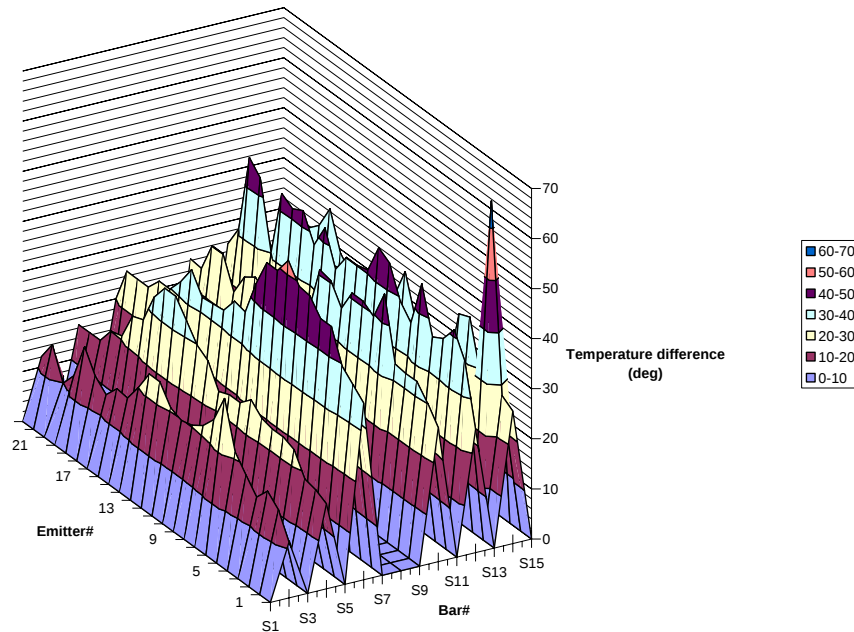


Figure 5.36: Measured temperature distribution of a BCPP diode laser array

5.3 CONCLUSION

We demonstrated the capability of low pressure water spray cooling to handle a typical diode laser array with heat flux of about 400 W/cm^2 . The individual emitter spectra we measured enabled us to estimate the emitter temperatures by using the known temperature coefficient of the diode laser wavelength. The result shows the emitter temperatures are quite high for the array tested. There are three main ways to reduce the diode temperature. One is to improve the cooling technique by using another liquid, for example Ammonia or R134a, for spray cooling. The second is to improve the nozzle design and optimize the droplet velocity and size to achieve an even higher heat transfer coefficient and the third is to develop a new packaging design which shortens the cooling path and/or uses materials with higher thermal conductivity.

A better package design is essential for high power diode laser arrays. Optical, mechanical, thermal and electrical issues of the diode laser array packaging have been discussed thoroughly. Several designs of the beam control prism package have been proposed to achieve better cooling. Based on the BCPP design, a 10 bar diode array was built and tested. The experimental results show trend similar to the FEA results. However, because the thermal expansion coefficient mismatch induces deformation of the substrate, the overall solder thickness increases and the thermal resistance is higher. Also, the thickness of the solder layer might not be uniform which makes the thermal resistance uneven and consequently the temperature distribution is not uniform.

6. CONCLUDING REMARKS

In this final report, we have concluded the following tasks. First, the characteristics of power MOSFETs have been demonstrated with simulation and fabrication. The computer-generated model of the device was simulated using the semiconductor device modeling and simulation package from Integrated System Engineering (ISE) TCAD. Simulations of the device were performed for a range of parameters. Also, several wafers were fabricated with different parameters. The aim was for all the devices fabricated to show the characteristics of a MOSFET. However, it is very difficult to obtain a high yield. As a result, the real devices showed huge variations. Many of the devices showed diode characteristics and some behaved as resistors. Two devices behaving as a power MOSFET were obtained.

Second, the MSL system that was designed and built at UCF addresses these needs and is the first step in making this a valuable method for MEMS. With more work and optimization, the MSL method could be ready for commercial use in a relatively short period of time. As more research in this area is completed, it is clear that the cost of this technology will decrease. This novel approach has never been attempted and would entail very strict timing requirements. However, if successful it would lead to the possibility of making very large parts on the macro scale with the resolution and accuracy of an MSL system. Currently we have the ability to begin trying research in this area with our current hardware. This may be an important method to examine the feasibility of a future for MSL.

Third, spray cooling of a large area was tested and reported. Since a nozzle usually can cool only 1 cm^2 , a nozzle array must be designed for cooling a large area in practice. Spray nozzles from Spraying System were modified so that they can fit tightly together with a center-to-center spacing of 1 cm. The array was positioned such that the nozzles are about 1.5 cm away from the surface. When multiple spray cones impinge onto a flat surface, the sprays interact with each other. The liquid was forced to escape off the surface along lines that are the intersection lines of the spray patterns. As more nozzles were added to the array, there was less area for the run-off liquid to flow through as it tries to exit the impingement surface. This reducing ratio of flow area to liquid volume flow rate results in an increase in the thickness of the liquid layer. To properly manage the excess liquid impinging on the surface from the multiple nozzles, grooves were machined on the heater surface. Suction tubes located at the corners can be used for excess liquid removal. Suction effectiveness

was helped greatly by adding extra siphons outside the spray area. Additionally, suction effectiveness was also increased by adding small slits to the sides of the siphons. Thus, this design presents a solution to the problem of surface flooding in a multiple pressure atomized spray nozzle array.

In the end, we demonstrated spray cooling in a high power laser application. Spray cooling has been used in industry, but typically the applications have involved very high surface temperatures. Due to the advancing requirements of the electronics industry, attention has been focused on high heat flux removal with lower surface temperatures. The potential of using spray cooled laser diode arrays is feasible. A short thermal conduction path is essential to reducing the thermal resistance between the emitter and the coolant. Evaporative spray cooling can provide temperature uniformity for all emitters in the diode array. Our novel BCPP design fulfills the requirements needed for diode array – low thermal resistance, temperature uniformity among emitters, low coolant flow rate and simplicity. Modifying the substrate configuration can compensate for the effect of edge cooling on the end emitters.

APPENDICES

A. Calculations for Ion Implantation

$$T := 1373 \text{ K}$$

$$x_j := 2 \cdot 10^{-4} \cdot \text{cm}$$

$$N_0 := 5 \cdot 10^{16} \cdot \text{cm}^{-3}$$

$$N_b := 2 \cdot 10^{14} \cdot \text{cm}^{-3}$$

$$D := 2.9 \cdot 10^{-13} \cdot \frac{\text{cm}^2}{\text{s}}$$

$$t := \frac{(x_j)^2}{D \cdot 4 \cdot \ln\left(\frac{N_0}{N_b}\right)}$$

$$t = 6.245 \cdot 10^3 \text{ s}$$

$$Q_1 := N_0 \cdot \sqrt{\pi \cdot D \cdot t}$$

$$Q_1 = 3.772 \cdot 10^{12} \cdot \text{cm}^{-2}$$

$$Q := 2 \cdot Q_1$$

$$Q = 7.543 \cdot 10^{12} \cdot \text{cm}^{-2}$$

Total dose for $x_j=2 \text{ um}$ is $7.543 \cdot 10^{12}$

$$E := 20$$

$$R_p := 0.0662 \cdot 10^{-4}$$

$$\Delta R_p := 0.0283 \cdot 10^{-4}$$

$$N_p := \frac{Q}{\sqrt{2 \cdot \pi} \cdot \Delta R_p}$$

$$N_p = 1.063 \cdot 10^{18} \cdot \text{cm}^{-2}$$

$$E := 25$$

$$R_p := 0.08 \cdot 10^{-4}$$

$$\Delta R_p := 0.03 \cdot 10^{-4}$$

$$N_p := \frac{Q}{\sqrt{2 \cdot \pi} \cdot \Delta R_p}$$

$$N_p = 1.003 \cdot 10^{18} \text{ cm}^{-2}$$

$$x_j := 1.5 \cdot 10^{-4} \cdot \text{cm}$$

$$t := \frac{(x_j)^2}{D \cdot 4 \cdot \ln\left(\frac{N_0}{N_b}\right)}$$

$$t = 3.513 \cdot 10^3 \text{ s}$$

$$Q_1 := N_0 \cdot \sqrt{\pi} \cdot D \cdot t$$

$$E := 30$$

$$R_p := 0.0987 \cdot 10^{-4}$$

$$\Delta R_p := 0.0371 \cdot 10^{-4}$$

$$N_p := \frac{Q}{\sqrt{2 \cdot \pi} \cdot \Delta R_p}$$

$$N_p = 8.111 \cdot 10^{17} \text{ cm}^{-2}$$

$$Q1 = 2.829 \cdot 10^{12} \text{ cm}^{-2}$$

$$Q := 2 \cdot Q1$$

$$Q = 5.657 \cdot 10^{12} \text{ cm}^{-2}$$

Total dose for xj=1.5 um is 5.657.10e12

$$E := 20$$

$$Rp := 0.0662 \cdot 10^{-4}$$

$$\Delta Rp := 0.0283 \cdot 10^{-4}$$

$$Np := \frac{Q}{\sqrt{2 \cdot \pi \cdot \Delta Rp}}$$

$$Np = 7.975 \cdot 10^{17} \text{ cm}^{-2}$$

$$E := 25$$

$$Rp := 0.08 \cdot 10^{-4}$$

$$\Delta Rp := 0.03 \cdot 10^{-4}$$

$$Np := \frac{Q}{\sqrt{2 \cdot \pi \cdot \Delta Rp}}$$

$$Np = 7.523 \cdot 10^{17} \text{ cm}^{-2}$$

$$E := 30$$

$$Rp := 0.0987 \cdot 10^{-4}$$

$$\Delta Rp := 0.0371 \cdot 10^{-4}$$

$$Np := \frac{Q}{\sqrt{2 \cdot \pi \cdot \Delta Rp}}$$

$$Np = 6.083 \cdot 10^{17} \text{ cm}^{-2}$$

B. Fabrication Process Procedure

Obtain a wafer from the lab and record the following information:

- a. thickness of the epilayer :3.5 μm
- b. Majority carrier type of substrate: n type
- c. Resistivity: 17-23 $\Omega\text{-cm}$
- d. Orientation of wafer: $\langle 100 \rangle$

Clean the wafer with trichloroethylene, acetone, methanol and deionized water.

Initial Oxidation

An initial oxide layer is grown on the wafer. The wafer is subjected to following conditions:

Wet Oxidation

Time: 27 minutes

Oxide thickness: 4000 $\text{\AA}$

Mask 1 Process

1. Place the wafer on the spinner, ensuring it is centered and apply a liberal amount of Futurrex NR8-1500 negative photoresist on wafer.
2. Spin the wafer at 3000 rpm for 30 seconds.
3. Place the wafer in a small bake plate and bake the silicon wafer with the photoresist material in the hardbake oven for 3 minutes.
4. Allow the silicon wafer to cool for 1 minute.
5. Place the wafer in the Mask aligner (Karl Suss) and expose for 10 sec using Mask # 1.
6. Pour Futurrex Developer in a Petri dish.
7. Place the wafer in the petri dish to develop the wafer for 6 minutes and 30 seconds. Then rinse the wafer with DI water and blow dry with Nitrogen. Inspect the wafer under a microscope to test the completions of developing. If the wafer still has photo-resist in the areas that were not exposed, repeat the procedure until the wafer is completely developed.
8. Place the silicon wafer in a small bake plate and bake the wafer in the hardbake oven for 10 minutes 20 seconds.
9. Allow the silicon wafer to cool for 1 minute. Meanwhile prepare the 9:1 BOE solution that will be used to etch the oxide and pour the solution in a plastic petri dish.

10. Place the wafer in the 9:1 BOE solution etch for 8 minutes and inspect the wafer. Assuming a 600 Angstrom etch rate the wafer should have etched in about 7 ½ minutes. When the oxide is etched completely the water will bead up and roll away very quickly on the backside of the wafer. The window should appear white.
11. Once the etching is complete, strip the photo-resist from the wafer using acetone, then strip the acetone with Methanol, then strip the Methanol with DI water. Inspect under a microscope for completeness.

Dry Oxidation For Implantation Process

For the Peak concentration of the implanted Boron to be at the silicon-silicon dioxide interface, an oxide of 1000°A is required as from the calculations made. Dry oxidation is carried out.

1. Ensure the Dry Oxidation furnace is set as follows:
 - a. Furnace Temperature: 1100° C
 - b. Flow rate: 4.0 on flowmeter.
 - c. Oxygen Pressure 5 psi
2. Load the wafer in the boat and put it in the furnace at a push-in rate of about 3 minutes.
3. Once the boat is in the furnace, close the door and allow the samples to drive-in for 51 minutes.
4. After 51 minutes, remove the samples from the furnace and let the wafer cool for some time.

Boron Implantation

The wafer was given for boron implantation to the company, Core Systems. The following data was supplied for implantation.

Implantation species: Boron

1. Implanter energy , $E = 30 \text{ KeV}$
2. Implanter dose , $Q = 5.657 \times 10^{12} \text{ atoms/cm}^2$
3. Projected Range, $R_p = 0.0987 \text{ um}$
4. Normal straggle, $\Delta R_p = 0.0371 \text{ um}$
5. Peak surface concentration (after implantation) = $6.083 \times 10^{17} / \text{cm}^3$
6. Area of each wafer:

Wafer # 5 = 6.8468 sq.cm

Wafer # 6 = 8.2296 sq.cm

Wafer # 7 = 8.1126 sq.cm

Wafer # 8 = 6.8472 sq.cm

Boron Drive-in

1. Ensure the Dry Oxidation furnace is set as follows:
 - a. Furnace Temperature: 1100° C
 - b. Flow rate: 4.0 on flowmeter.
 - c. Oxygen Pressure 5 psi
2. Place wafer in the boat and load boat in the furnace at a push-in rate of about 3 minutes.
3. Once the boat is in the furnace, close the door and allow the samples to drive-in for 32 minutes.
4. After 32 minutes, remove the samples from the furnace and let the wafer cool for some time.

Mask 2 Process

1. Place the wafer on the spinner, ensuring it is centered and apply a liberal amount of Futurrex NR8 – 1500 negative photoresist on wafer.
2. Spin the silicon wafer at 3000 rpm for 30 seconds.
3. Place the silicon wafer in a small bake plate and bake the silicon wafer with the photo-resist material in the hardbake oven for 3 minutes 20 seconds.
4. Allow the silicon wafer to cool for 1 minute.
5. Place the mask plate into the mask aligner and expose the silicon wafer for 10 seconds using Mask # 2.
6. Pour Futurrex Developer In a Petri dish.
7. Place the wafer in the petri dish to develop the wafer for 5 minutes, and then rinse the wafer DI water and blow dry the wafer with Nitrogen. Inspect the wafer under a microscope to test the completion of developing. If the wafer has still photresist in the areas that were not exposed, repeat the procedure until the wafer is completely developed.
8. Place the silicon wafer in a small bake plate and hard bake the wafer in the hardbake for 10 minutes 20 seconds.

9. Allow the silicon wafer to cool for 1 minute. Meanwhile prepare the 9:1 BOE solution that will be used to etch the silicon wafer and pour the solution in a plastic petri dish.
10. Place the wafer in the 9:1 BOE solution etch for 6 minutes and inspect the wafer.
11. Once the etching is complete, strip the photo-resist from the wafer using acetone, then strip the acetone with Methanol, then strip the Methanol with DI water. Inspect under a microscope for completion.

Phosphorus Predep

1. Remove the boat from the furnace and place the samples into the slots close to the phosphorous diffusion sources. Ensure the furnace for the pre-dep is as follows:
 - a. Furnace Temperature: 950° C
 - b. Nitrogen Pressure 5 psig
 - c. Flow rate: 4.0 on flowmeter.
2. Load the boat back into the furnace at a push-in rate of about 1 minute.
3. Once the boat is in the furnace, close the door and allow the samples to predep for 10 minutes.
4. After the 10 minutes mark, remove the boat at a pull out rate of 1 minute, and allow the sample to cool.

Gate Oxidation

1. Etch off all the oxide with 9:1 BOE solution for 11 minutes.
2. Set the Dry Oxidation furnace to following conditions.
 - a. Furnace Temperature: 1100° C
 - b. Flow rate: 4.0 on flowmeter.
 - c. Oxygen Pressure 5 psi
3. Place the wafer in the boat and load boat in the furnace at a push-in rate of about 3 minutes.
4. Once the boat is in the furnace, close the door and allow the samples to drive-in for 32 minutes.
5. After 32 minutes, remove the samples from the furnace and let the wafer cool for some time.
6. Remove the samples at a full-rate of 3 minutes.
7. The oxide thickness from the drive in process is 300°A.

Mask 3 Process

1. Place the wafer on the spinner, ensuring it is centered and apply a liberal amount of Futurrex NR8 – 1500 negative photoresist on wafer working from the inside out. Spin the silicon wafer at 3000 rpm for 30 seconds.
2. Place the silicon wafer in a small bake plate and bake the silicon wafer with the photoresist material in the hardbake oven for 3 minutes 20 seconds.
3. Allow the silicon wafer to cool for 1 minute and make preparations to the mask aligner.
4. Place the Mask-3 mask plate into the Karl Suss aligner and expose the silicon wafer for 10 seconds.
5. Pour Futurrex Developer in a petri dish.
6. Place the wafer in the petri dish to develop the wafer for seven minutes and thirty seconds. Then rinse the wafer with DI water and blow dry the wafer with Nitrogen. Inspect the wafer under a microscope to test the completion of developing. If the wafer still has photoresist in the areas that were not exposed, repeat the procedure until the wafer is completely developed.
7. Place the silicon wafer in a small bake plate and bake the wafer in the hardbake oven for 10 minutes 20 seconds.
8. Allow the silicon wafer to cool for 1 minute. Meanwhile prepare the 9:1 BOE solution that will be used to etch the silicon wafer and pour the solution in a plastic petri dish.
9. Place the wafer in the 9:1 BOE solution etch for 3 minutes and inspect the wafer. When the oxide is etched completely the water will bead up and roll away very quickly on the back side of the wafer.
10. Once the etch is complete, strip the photo-resist from the wafer using acetone, then strip the acetone with Methanol, then strip the Methanol with DI water. Inspect under a microscope for completeness.

Aluminum Deposition

For metallization, aluminum was deposited on the wafer in a vacuum system.

Mask 4 Process

1. Place the wafer on the spinner, ensuring it is centered and apply a liberal amount of Shipley 1400-27 positive photoresist on wafer.
2. Spin the silicon wafer at 3000 rpm for 30 seconds.
3. Place the silicon wafer in a small bake plate and soft bake the silicon wafer with the photoresist material in the hardbake oven for 3 minutes.

4. Allow the silicon wafer to cool for 1 minutes and make preparations to the mask aligner.
5. Place the Mask-4 into the Cobilt Mask aligner and expose the silicon wafer for 10 seconds.
6. In Pertri dish pour 40mL of Shipley developer and develop for about 4 minutes and 30 seconds then rinse the wafer with DI Water.
7. Dry the wafer with Nitrogen gun and inspect the wafer under a microscope for developing completeness. If the wafer still has photo-resist in the areas that were exposed, repeat the procedure until the wafer is completely developed.
8. Place the silicon wafer in a small bake plate and hard bake the wafer in the oven for 3 minutes.
9. Allow the silicon wafer to cool for 1 minute. Meanwhile prepare the Aluminum etch solution.
10. Place the wafer in the etch solution on a burner pad for 3 minutes and rinse the wafer with DI water. Inspect the wafer under a microscope to see if the etch is completed.
11. Once you are sure the aluminum is fully etched, remove the photoresist using acetone. Then rinse the wafer with methanol, DI water, and blow-dry with Nitrogen.

Deposit Aluminium on the back side of wafer for the drain contact.

C. Command File of DIOS

```
Title ("dmos")

grid (x(0.0,12.0), y(-4.0,0.0),nx=27)

comment('n-substrate')
substrate(orientation=100,element=P,conc=2E14,ysubs=0.0)

Replace(control(ngraphic=10))

graph(triangle=on,plot)


diffusion(temperature=1100degc,atmosphere=H2O,thickness=400nm,TH2O=98degc)

mask(material=resist, thickness=0.8um,XLeft=0.0,XRight=4.25)
mask(material=resist,thickness=0.5um,XLeft=11.0,XRight=12.0)

etch(material=0X,stop=sigas,rate(Isotropic=60))
etch(material=resist)

diffusion(temperature=1100degc,atmosphere=02,thickness=100nm)

comment('p-region')
implant(element=B,dose=5.66E12,energy=30Kev, tilt=7)

diffusion(element=B, temperature=1100degc,time=58min,atmosphere=02)

comment('n-region')
mask(material=resist,thickness=0.8um,XLeft=0.0,XRight=5.9)
mask(material=resist,thickness=0.5um,XLeft=9.5,XRight=12.0)

etch(material=0X,stop=sigas,rate(Isotropic=60),over=20%)
etch(material=resist)

diffusion(element=P,time=10min,temperature=950degc)

etch(material=0X,stop=sigas,rate(Isotropic=60),over=20%)

comment('gate oxidation')

diff:(dth=2nm)
diffusion(temperature=1100degc,atmosphere=02,thickness=35nm)

comment('contact windows')
mask(material=resist,thickness=0.8um,XLeft=0.0,XRight=6.5)
mask(material=resist,thickness=0.5um,XLeft=11.2,XRight=12.0)

etch(material=0X,stop=sigas,rate(Isotropic=60),over=20%)
etch(material=resist)
```

```

comment('full device structure')
reflect(reflect=0.0)

comment('metal contacts')
mask(material=al, thick=0.03, x(-11.5, -6.2, -5.0, 5.0, 6.2, 11.5))

comment('save final DIOS simulation file')

save(file=n@node@)

comment('save final structure for device simulation')

save(file='n@node@', type=MDRAW, synonyms(al=metal),
contacts(
contact1(name='source1', -9.0, -7.0)
contact2(name='gate', -2.0, 2.0)
contact3(name='source2', 7.0, 9.0)
contact4(name='drain', location=bottom)
),
species(NetActive, BTotal, PTotal), control(y0=0.0)
)

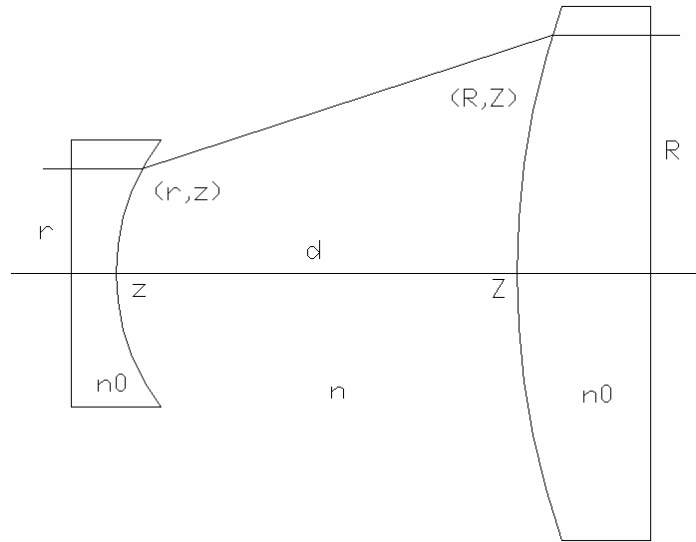
end

```

D. Beam Shaping Design

For the microstereolithography machine we need a uniform profile. To do this there are a few different methods. The first method of creating a uniform beam is to just absorb more of the higher intensity beam and let the lower intensity edges pass through. This uses a plano-convex gray lens cemented to a plano-concave clear lens (Figure 15). Of course this design has a major disadvantage, only 30% of the beam is transmitted through the system. This fact makes this method to shaping a beam from gaussian to uniform a very poor approach.

Another approach is the geometrical methods technique. Using geometrical methods to shape a laser beam profile involves application of geometrical optics to solve the optical design problem. Specifically, the laws of reflection and refraction are used along with ray tracing, conservation of energy within a bundle of rays, and constant optical path length condition to design laser beam profile shaping optical system. Interference or diffraction effects are not considered as part of the design process. Only lenses and mirrors are used for the optical components of the laser beam profile shaping system. This is the approach that we want to take with the laser system. These are the calculations to arrive at a system of two lenses.



Beam expander design

We are given a beam with an intensity profile of

$$\sigma(r) := e^{-2 \cdot \left(\frac{r}{r_0} \right)^2} \quad \text{Equation 1}$$

We want to achieve an energy density of

$$\Sigma = \text{constant} \quad \text{Equation 2}$$

R can be evaluated by the following equation:

$$R(r) = \left[\frac{r_0^2}{2 \cdot \Sigma} \left[1 - e^{\left(\frac{-2 \cdot r^2}{r_0^2} \right)} \right] \right]^{\frac{1}{2}} \quad \text{Equation 3}$$

where Σ is given by

$$\Sigma = \frac{r_0^2}{2 \cdot R_{\max}^2} \cdot \left(1 - e^{\left(\frac{-2 \cdot r_{\max}^2}{r_0^2} \right)} \right) \quad \text{Equation 4}$$

I used values of 2 mm for $r_{\max}$ and 25 mm for $R_{\max}$.

Using Snell's law we can arrive at the following differential equation for the refracted light beams.

$$\left[(1 - \gamma^2)(Z - z)^2 - \gamma^2(R - r)^2 \right] \cdot \left(\frac{d}{dr} z \right)^2 + 2(R - r)(Z - z) \cdot \left(\frac{d}{dr} z \right) + (r - r)^2 = 0 \quad \text{Equation 5}$$

Where

$$\gamma = \frac{n}{n_0} \quad \text{Equation 6}$$

Also using the constant optical path condition to get a plane of light from the output we get the equation:

$$(Z - z) = \frac{n(n - n_0) \cdot d + \left[n_0^2 \cdot (n - 1)^2 \cdot d^2 + (n^2 - n_0^2) \cdot (R - r)^2 \right]^{\frac{1}{2}}}{n^2 - n_0^2} \quad \text{Equation 7}$$

For this calculation I chose a distance $d=250$ mm. Combining equations 8 and 10 we can solve for z with a simple integration.

$$z(r) = \int f(r) dr + C \quad \text{Equation 8}$$

also it follows from equations 6 and 10

$$Z(r) = z(r) + g(r) \quad \text{Equation 9}$$

where $g(r)$ is given by:

$$g(r) := \frac{n \cdot (n - n_0) \cdot d + \left[n_0^2 \cdot (n - 1)^2 \cdot d^2 + \left(n^2 - n_0^2 \right) \cdot \left[\sqrt{\frac{r_0^2}{2 \cdot \Sigma} \cdot \left[1 - e^{\left(\frac{-2 \cdot r^2}{r_0^2} \right)} \right]} - r \right]^2 \right]^{\frac{1}{2}}}{n^2 - n_0^2} \quad \text{Equation 10}$$

Similarly we can rewrite equation 8 as:

$$a(r) \cdot \left(\frac{d}{dr} z \right)^2 + b(r) \cdot \left(\frac{d}{dr} z \right) + c(r)^2 = 0 \quad \text{Equation 11}$$

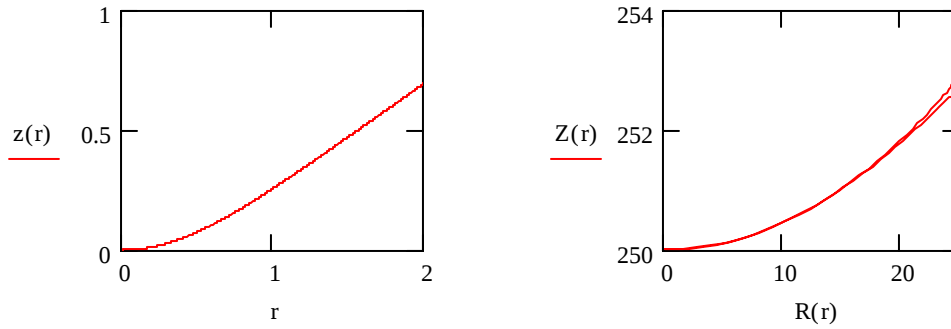
Where:

$$c(r) := \left[\sqrt{\frac{r_0^2}{2 \cdot \Sigma} \cdot \left[1 - e^{\left(\frac{-2 \cdot r^2}{r_0^2} \right)} \right]} - r \right]$$

$$b(r) := 2 \cdot c(r) \cdot g(r)$$

$$a(r) := \left(1 - \gamma^2 \right) \cdot g(r)^2 - \gamma^2 \cdot c(r)^2 \quad \text{Equation 12}$$

Using the quadratic equation we can solve for z prime and then insert it into equation 12. Integrating from $-r$ to r lead to the following plots of (r, z) and (R, Z) . These are the profiles of the shapes of the aspheric lenses.



Lens profiles for beam shaping

These lenses will shape the beam from a gaussian profile to a uniform intensity while increasing the diameter from 2mm to 40 mm.

REFERENCES

- An, Linan, Wenge Zhang, Victor M. Bright, Martin L. Dunn, and Rishi Raj. "Development of injectable polymer-derived ceramics for high temperature MEMS." IEEE: Micro Electro Mechanical Systems (MEMS) (2000).
- Baliga, B. J., Power Semiconductor Devices, p.2, PWS Publishing Company, 1995a
- Baliga, B. J., Power Semiconductor Devices, p.335, PWS Publishing Company, 1995b
- Baliga, B. J., Power Semiconductor Devices, p.336, PWS Publishing Company, 1995c
- Baliga, B. J., "Semiconductor for high-voltage, vertical channel FETs," Journal of Applied Physics, vol. 53, p-1759, 1982
- Baliga, B.J., IEEE Electron Device Letters, vol. 10, p.455, 1989
- Bauer, W., and R. Knitter, "Development of a rapid prototyping process chain for the production of ceramic microcomponents." Journal of Materials Science Volume 37 (2002): 3127-3140
- Baysinger, K. M., K. L. Yerkes and T. E. Michalak, "Design of a Microgravity Spray Cooling Experiment", 42nd AIAA Aerospace Conference, Reno, NV, Paper No. AIAA-2004-0966
- Beluze, Laurence, Arnaud Bertsch, and Philippe Renaud. "Microstereolithography: a new process to build complex 3D objects." Design, Test, and Microfabrication of MEMS and MOEMS SPIE Volume 3680 (1999): 808-817.
- Bertsch, A., H. Lorenz, and P. Renaud. "3D microfabrication by combining microstereolithography and thick resist UV lithography." Sensors and Actuators Volume 73 (1999): 14-23.
- Bertsch, A., J.Y. Jezequel, and J.C. Andre. "Study of the spatial resolution of a new 3D microfabrication process: the microstereophotolithography using a dynamic mask-generator technique." Journal of Photochemistry and Photobiology Volume 107 (1997a): 275-281.
- Bertsch, Arnuaud, Paul Bernhard, Christian Vogt, and Philippe Renaud. "Rapid Prototyping of small size objects." Rapid Prototyping Journal Volume 6 Number 4 (2000): 259-266.
- Bertsch, Arnuaud, Paul Bernhard, and Philippe Renaud. "Microstereolithography: Concepts and applications." IEEE: Emerging Technologies and Factory Automation (ETFA) (2001): 289-298.

Bertsch, A., S. Zissi, J.Y. Jezequel, S. Corbel, and J.C. Andre. "Microstereolithography using a liquid crystal display as dynamic mask-generator." *Microsystem Technologies* (1997b):42-47.

Canali, C., G. Majni, R. Minder, and G. Ottaviani, "Electron and Hole Drift Velocity Measurements in Silicon and their Relation to Electric Field and Temperature," *IEEE Transactions on Electron Device*, vol. ED-22, pp. 1045-1047, 1975

Chatwin, Chris, Maria Farsari, Shiping Huang, Malcolm Heywood, Philip Birch, Rupert Young, and John Richardson. "UV microstereolithography system that uses spatial light modulator technology." *Applied Optics* Volume 37 Issue 32 (1998): 7514-7522.

Chen, R. H., L. C. Chow and J. E. Navedo, "Effects of Spray Cooling Characteristics on Critical Heat Flux in Subcooled Water Spray Cooling", *Int. J Heat Mass Transfer*, vol. 45, 2002, pp.4033-4043

Chen, X. Q., L. C. Chow and M. S. Sehmbe, "Thickness of Film Produced by a Pressure Atomizing Nozzle", 30th AIAA Thermophysics and Heat Transfer Conference, San Diego, CA, Paper No. AIAA-95-2103

Chinnock, Chris. "Home Theater Projectors; The Next Big Thing?" *Proceedings of SPIE:Projection Displays VII* Volume 4657 (2002): 31-37.1-10.

Chizhov, A. V. and K. Takayoma "The Impact of Compressible Liquid Droplets on hot rigid surface," *Int. J. Heat Mass Transfer*, vol. 47, 2004, pp.1391-1401

Chow, L. C., M. S. Sehmbe and M. R. Pais, "High Heat Flux Spray Cooling", *Annual Review of Heat Transfer*, vol. 8, 1997, pp. 291-318

David P. DeWitt Frank P. Incropera, *Introduction to Heat Transfer*, 3 ed. (John Wiley & Sons, 1996)

Farsari, M., F. Claret-Tournier, S. Huang, C.R. Chatwin, D.M. Budgett, P.M. Birch, R.C.D. Young, and J.D. Richardson. "A novel high-accuracy microstereolithography method employing an adaptive electro-optic mask." *Journal of Materials Processing Technology* Volume 107 (2000): 167-172.

Frank P. Incropera and David P. DeWitt, (Wiley, New York, 1996), pp. xxiii.

Grant, D. A. and J. Gower, *Power MOSFETS: Theory and Applications*, p. 5, John Wiley & Sons, New York, 1989a

Grant, D. A. and J. Gower, *Power MOSFETS: Theory and Applications*, John Wiley & Sons, New York, 1989b

Huang, S., M.I. Heywood, R.C.D. Young, M. Farsari, and C.R. Chatwin. "Systems control for a micro-stereolithography prototype." *Microprocessors and Microsystems Volume 22* (1998): 67-77.

Huddle, JJ, A. Marcos, T. Chung, S. J. Lindauer II, S. Lei, D. P. Rini, L. C. Chow, M. Bass, and P. J. Delfyett, SAE Power System Conference, San Diego, CA, Oct 2000.

Ikuta, Koji, Ken Hirowatari, and Tsukasa Ogata. "Three dimensional micro integrated fluid systems (MIFS) fabricated by stereo lithography." *IEEE: Micro Electro Mechanical Systems (MEMS)* (1994): 1-6.

Ikuta, Koji, Shoji Maruo, and Syunsuke Kojima. "New micro stereolithography for freely movable 3D micro structure." *IEEE: Micro Electro Mechanical Systems (MEMS)* (1998): 290-295.

ISE TCAD Manual, p. 12-237, Release 6, Volume 3

Jacobs, Paul F. *Rapid Prototyping & Manufacturing: Fundamentals of Stereolithography*. Dearborn, MI: Society of Manufacturing Engineers, 1992.

Jaeger, Richard C., *Introduction to Microelectronic Fabrications*, 1993

Knitter, Regina, Werner Bauer, Dieter Gohring, and Peter Risthaus. "RP process chains for ceramic microcomponents." *Rapid Prototyping Journal Volume 8 Number2* (2002): 76-82.

Kunzman, A., J. O'Connor, D. Segler, T. Migl, and K. Bell. "Advancing the DMD™ Device for High-Definition Home Entertainment." *Proceedings of SPIE: Digital Cinema and Microdisplays Volume 4207* (2000):

Liew, Li-Anne, Ruiling Luo, Yiping Liu, Wenge Zhang, Linan An, Victor M. Bright, Martin L. Dunn, John W. Daily and Rishi Raj. "Fabrication of multi-layered SiCN ceramic MEMS using photo-polymerization of precursor." *IEEE: Micro Electro Mechanical Systems Conference (MEMS)* (2001a): 86-89.

Liew, Li-Anne, Wenge Zhang, Victor M. Bright, Linan An, Martin L. Dunn, and Rishi Raj. "Fabrication of SiCN ceramic MEMS using injectable polymer-precursor technique." *Sensors and Actuators Volume A 89* (2001b): 64-70.

Lin, L., R. Ponnappan, K Yerkes and B. Hager, "Large Area Spray Cooling" 42nd AIAA Thermophysics and Heat Transfer Conference, Reno, NV, Paper No. AIAA-2004-1340

Maruo, Shoji and Koji Ikuta. "New microstereolithography (Super-IH Process) to create 3d freely movable micromechanism without sacrificial layer technique." *IEEE: Micromechantronics and human science* (1998): 115-120.

Maruo, Shoji, Koji Ikuta, and Hayato Korogi. "Remote light-driven micromachines fabricated by 200 nm microstereolithography." IEEE: Nanotechnology (2001a): 507-512.

Maruo, Shoji, Koji Ikuta, and Korogi Hayato. "Light-driven MEMS made by high-speed two-photon microstereolithography." IEEE: Micro Electro Mechanical Systems Conference (MEMS) (2001b): 594-597.

Maruo, Shoji, Koji Ikuta, and Toshihide Ninagawa. "Multi-polymer microstereolithography for hybrid opto-MEMS." IEEE: Micro Electro Mechanical Systems Conference (MEMS) (2001c): 151-154.

Mauriello, Robert "Simulation and Fabrication of the Power MOSFET under Cryogenic Conditions", Thesis, 1998

Mazzoni, O. S., Low Temperature Operation Power Semiconductors. May 1993

Monneret, Serge, Christophe Provin, and Herve Le Gall. "Micro-scale rapid prototyping by stereolithography." IEEE: Emerging Technologies and Factory Automation (ETFA) (2001): 299-304.

Monneret, Serge, Virginie Loubere, and Serge Corbel. "Microstereolithography using a dynamic mask generator and a non-coherent visible light source." Design, Test and Microfabrication of MEMS and MOEMS SPIE Volume 3680 (1999): 553-561.

Nakamoto, Takeshi, Katsumi Yamaguchi, Petros A Abraha, and Kunihiro Mishima. "Manufacturing of three-dimensional micro-parts by UV laser induced polymerization." Journal of Micromechanics and Microengineering Volume 6 (1996): 240-253.

Nelson, T.J., and J.R. Wullert II. Electronic Information Display Technologies. Singapore: World Scientific Publishing, 1997.

Oxner, E. S., Power FET's and their Applications, P.88, Prentice Hall, Englewood Cliffs, CA, 1982

Pelly, B. R., "Power MOSFET's- A Status Review," Power Electronics Conference, P. 19-32, 1983

Sehmbey M. S. Louis C. Chow, and Martin R. Pais, in *Annual Review of Heat Transfer* (1997), Vol. 8, pp. 291.

Shams, Shaikh Ferdouse" Simulation of Silicon Carbide Power Metal Oxide Semiconductor Field Effect Transistors at High Temperature", Thesis, 1998

Streetman, B. G., Solid State Electronic Devices, Fourth Edition, Prentice Hall Series, 1995

Sun, Cheng and Xiang Zhang. "The influence of the material properties on ceramic micro-stereolithography." *Sensors and Actuators Volume A* 101 (2002): 364-370.

Sun, S. C. and J. D. Plummer, "Modeling of the On-Resistance of LDMOS, VDMOS, and VMOS Power Transistors," *IEEE Transactions on Electron Devices*, vol. Ed-27, P. 356-367, 1980

Taylor, B. E., *Power MOSFET Design*, P. 38, New York, Wiley-Interscience, 1993

Unger, Marc A., Hou-Pu Chou, Todd Thorsen, Axel Scherer, and Stephen R. Quake. "Monolithic microfabricated malves and pumps by multilayer soft lithography." *Science Magazine Volume* 288 (2000): 113-116.

Varadan, V.K., and V.V. Varadan. "Micro stereo lithography for fabrication of 3D polymeric and ceramic MEMS." *MEMS Design, Fabrication, Characterization and Packaging SPIE Volume* 4407 (2001): 147-157.

Ventura, S., S. Narang, P. Guerit, S. Liu, D. Twait, P. Khandelwal, E. Cohen, and R. Fish. "Freeform Fabrication of functional silicon nitride components by direct photo shaping." *Materials Research Society Symposium Volume* 625 (2000): 81-89.

Vogt, C., A. Bertsch, P. Renaud, and P. Bernhard. "Methods and algorithms for the slicing process in microstereolithography." *Rapid Prototyping Journal Volume* 8 Number3 (2002): 190-199.

Yipping, Liu, Li-Anne Liew, Ruiling Luo, Linan An, Victor M. Bright, Martin L. Dunn, John W. Daily, and Rishi Raj. "Fabrication of SiCN MEMS structures using microforged molds." *IEEE: Micro Electro Mechanical Systems Conference (MEMS)* (2001): 118-121.

Young, D F., Munson, Okiihi, *A Brief Introduction to Fluid Mechanics*, John Wiley & Sons, Inc. 1997, pg. 406

Zhang, X., X.N. Jiang, and C. Sun. "Micro-stereolithography of polymeric and ceramic microstructures." *Sensors and Actuators Volume* 77 (1997): 149-156.

Zhang, X., X.N. Jiang, and C. Sun. "Micro-stereolithography for MEMS." *Micro-Electro-Mechanical Systems (MEMS) ASME DSC Volume* 66 (1998): 3-9.

Zissi, S., A. Bertsch, J.Y. Jezequel, S. Corbel, D.J. Lougnot, and J.C. Andre. "Stereolithography and microtechniques." *Microsystem Technologies* 2 (1996): 97-102.