

AD-A257 586



RL-TR-92-28
Final Technical Report
March 1992



EHF FIBER OPTIC TRANSMITTER MODULLE

Hughes Research Laboratory

R. R. Hayes

DTIC
ELECTE
NOV 24 1992
S E D

APPROVED FOR PUBLIC RELEASE; DISTRIBUTION UNLIMITED.

92-30081

4500

Rome Laboratory
Air Force Materiel Command
Griffiss Air Force Base, New York

This report has been reviewed by the Rome Laboratory Public Affairs Office (PA) and is releasable to the National Technical Information Service (NTIS). At NTIS it will be releasable to the general public, including foreign nations.

RL-TR-92-28 has been reviewed and is approved for publication.

APPROVED:



PAUL SIERAK
Project Engineer

FOR THE COMMANDER:



JOHN A. GRANIERO
Chief Scientist
Command, Control & Communications Directorate

If your address has changed or if you wish to be removed from the Rome Laboratory mailing list, or if the addressee is no longer employed by your organization, please notify RL (C3DB), Griffiss AFB NY 13441-5700. This will assist us in maintaining a current mailing list.

Do not return copies of this report unless contractual obligations or notices on a specific document require that it be returned.

REPORT DOCUMENTATION PAGE

Form Approved
OMB No. 0704-0188

Public reporting burden for this collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington, VA 22202-4302, and to the Office of Management and Budget, Paperwork Reduction Project (0704-0188), Washington, DC 20503.

1. AGENCY USE ONLY (Leave Blank)		2. REPORT DATE March 1992		3. REPORT TYPE AND DATES COVERED Final	
4. TITLE AND SUBTITLE EHF FIBER OPTIC TRANSMITTER MODULE				5. FUNDING NUMBERS C - F39602-88-C-0039 PE - 63726F PR - 2863 TA - 92 WU - 22	
6. AUTHOR(S) R. R. Hayes					
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Hughes Research Laboratory 3011 Malibu Canyon Road Malibu CA 90265				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES) Rome Laboratory (C3DB) Griffiss AFB NY 13441-5700				10. SPONSORING/MONITORING AGENCY REPORT NUMBER RI-TR-92-28	
11. SUPPLEMENTARY NOTES Rome Laboratory Project Engineer: Paul Sierak/C3DB(315)330-4092					
12a. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution unlimited.				12b. DISTRIBUTION CODE	
13. ABSTRACT (Maximum 200 words) Some aspects of the theory of free-carrier-induced refractive index changes in semiconductors are covered as well as some device structures aimed at exploiting this effect to provide for high frequency analog optical modulation are covered in this report. High optical loss and the requirement of large current changes at high frequencies prevented exploiting this effect initially at 10 GHz and then at greater than 30 GHz in the directional coupler architecture. A thermal switch was incidentally explored which had a risetime of 0.3 microseconds and required 0.8 watts to operate.					
14. SUBJECT TERMS Optical Modulation, electroptic modulators, freecarrier induced index of refraction changes.				15. NUMBER OF PAGES 120	
				16. PRICE CODE	
17. SECURITY CLASSIFICATION OF REPORT UNCLASSIFIED	18. SECURITY CLASSIFICATION OF THIS PAGE UNCLASSIFIED	19. SECURITY CLASSIFICATION OF ABSTRACT UNCLASSIFIED	20. LIMITATION OF ABSTRACT U/1.		

TABLE OF CONTENTS

SECTION		PAGE
1	INTRODUCTION.....	1
2	FREE-CARRIER EFFECTS IN SEMICONDUCTORS.....	4
3	DEPLETION-EDGE MODULATORS.....	10
	3.1 Depletion Edge Translation in an Optical Waveguide.....	10
	3.2 Optical Modulators.....	21
4	FREE CARRIER INJECTION MODULATORS.....	58
	4.1 Injection in PIN Junctions.....	58
	4.2 Time and Spatial Response of a PIN Device.....	61
5	EXPERIMENTAL RESULTS.....	66
	5.1 Optical Waveguide Fabrication and Measurement.....	66
	5.2 Free-Carrier Injection Modulators.....	73
	5.3 PIN Guiding-Antiguinding Modulator.....	79
	5.4 PIN Directional Coupler Modulator.....	83
	5.5 Electrooptic and Thermal Modulators.....	93
	SUMMARY.....	103
	REFERENCES.....	105

DTIC QUALITY INSPECTED 4

Accession For	
NTIS CRA&I	☒
DTIC TAB	☐
Unannounced	☐
Justification	
By	
Distribution /	
Availability Codes	
Dist	Avail and/or Special
A-1	

LIST OF ILLUSTRATIONS

FIGURE		PAGE
3.1a	Cross-sectional view of a semiconductor waveguide structure.....	11
3.1b	The same structure with p and n-type dopants....	12
3.2	Waveguide structure showing the change in the depletion width, dD , that results from a change in the bias voltage, and $\psi(y)$	17
3.3	Profile of the electric field within the depletion region.....	19
3.4	Top view of a waveguide directional coupler.....	22
3.5	Optical field intensities in the coupler of fig. 3.4 as a function of position along its length.....	22
3.6	The traveling wave RIMFET.....	25
3.7	Cross-sectional view of the RIMFET with the corresponding index profiles mapped below.....	26
3.8	RIMSCHOTT traveling wave modulator.....	27
3.9	Cross-sectional view of the interguide region of the RIMSCHOTT device, showing how the depletion region (the dotted area) is located near the minimum of the optical field strength, $\psi(y)$	29
3.10	A top-view of a Mach-Zehnder interferometer, showing the optical phase shifts ϕ that must be induced in each arm to obtain modulation.....	31
3.11	The same Mach-Zehnder interferometer with traveling-wave electrodes.....	31
3.12	Waveguide dimensions: t is the thickness of the guiding layer, D is the depletion width, and W is the waveguide width.....	35
3.13	Drive power required for 50% depth of modulation for a 1.0 cm long device using either the free-carrier (FC) or the electrooptic (EO) effect, but not both.....	36

LIST OF ILLUSTRATIONS (Continued)

FIGURE		PAGE
3.14	Same as fig. 3.13, but including the EO effect in the FC results (as would happen in a real device).....	37
3.15	Required drive power for the device of example 2, table II.....	38
3.16	Required drive power for the device of example 3.....	39
3.17	A Mach-Zehnder modulator using coplanar strip (CPS) electrodes.....	41
3.18	Cross-sectional view of the CPS device of fig. 3.17.....	42
3.19	Equivalent circuit of the CPS structure of fig. 3.18.....	42
3.20	Cross-sectional view of a PN stack (a) and the equivalent electrical circuit (b).....	48
3.21	Simplified cross-sectional view of the electrode region of the CPS interferometer showing the location and direction of the B-field in the n^+ strip (the dotted region).....	52
4.1	Cross-sectional view of a PIN diode with current flowing.....	59
4.2	Equivalent circuit of a PIN diode with a narrow top electrode that shows how the finite resistivity of the top layer reduces the current flow through outlying sections of the structure.....	62
4.3	Free-carrier density as a function of distance away from the center of the injecting electrode.....	64
5.1	Optical waveguides formed by wet-chemical etching.....	68
5.2	Experimental setup used to measure waveguide loss, mode profiles, and modulator response times.....	69

LIST OF ILLUSTRATIONS (Continued)

FIGURE		PAGE
5.3	Variation of waveguide transmission as the guide is slowly heated.....	71
5.4	Optical waveguides using a single heterostructure.....	72
5.5	Heterostructure design for free-carrier injection modulator.....	74
5.6	I-V trace for diodes fabricated from the material of fig. 5.5.....	76
5.7	Logarithmic plot of the forward-biased characteristics of the same diode.....	77
5.8	Plot of the emitted light intensity (as measured by a Silicon photodiode) as a function of the injected current.....	78
5.9a	PIN "anti-guiding" modulator.....	80
5.9b	Change in transmission (bottom trace) versus applied voltage (top trace) for the anti-guiding modulator.....	82
5.10	Cross-sectional view of PIN directional coupler modulator.....	84
5.11	Top view of PIN directional coupler modulator showing the wide electrode used to insure uniform current injection.....	85
5.12	Photographs of the PIN coupler modulator before up-plating.....	87
5.13	Profiles of the optical intensity at the output facets of the PIN coupler modulator for two different injection currents.....	88
5.14a	Change in transmission (bottom trace) of the electrooptic switch of section 5.5 driven by a drive pulse (top trace).....	90
5.14b	Change in transmission of one channel of the PIN directional coupler modulator (bottom trace) for a 100 ma injected current pulse (top trace).....	91

LIST OF ILLUSTRATIONS (Continued)

FIGURE		PAGE
5.15	Computed density profiles (top graph) for the PIN coupler modulator described in the text (1 ns between traces).....	92
5.16	Top (a) and cross-sectional (b) view of directional coupler switch with an interguide Schottky electrode.....	94
5.17	(a) SEM photograph of cleaved facet of the structure of fig. 5.16., showing the two guides and the interguide electrode. (b) The I-V response of the Schottky diode.....	96
5.18	Intensity profiles at output facet of the device of fig. 5.16 as a function of reverse bias voltage.....	97
5.19	Intensity profiles for three different values of forward bias current.....	99
5.20	Change in device transmission (bottom trace) for a 180 ma forward-bias current pulse (top trace).....	100
5.21a	Lateral temperature profile for a 0.8 watt heat step confined to the inter-guide (dotted) region.....	101
5.21b	Temperature difference between center of dotted region and middle of each guide (the small crosses at the top of fig. a).....	101

SECTION 1

INTRODUCTION

Optical modulators vary the amplitude of a lightwave by changing either the real or the imaginary part of the index of refraction. Devices that change the imaginary part usually function by operating near a tuneable resonance of the material, so that the optical transmittance can be varied by shifting this resonant frequency. Devices that utilize changes in the real part of the index of refraction either employ coherent interferometric techniques, or vary the angle of total internal reflection. Regardless of the approach taken, one must be able to change the index of refraction by means of some external stimulus.

This stimulus is quite often an electric field. There exists a class of non-centrosymmetric materials for which the real part of the index of refraction can be changed by the application of a static or alternating electric field. The effect, called the electrooptic effect, is normally quite small, and leads to usable devices only because integrated-optic devices allow one to concentrate the field in small volumes. Even then, the index change for a typical material, such as GaAs, is only 10^{-3} for achievable field strengths. Such a small change makes the practical realization of compact, efficient devices difficult.

Better modulators with lower operating power and shorter lengths could be made if other physical phenomena could be found that would allow larger changes in the index of refraction. One candidate is the free-carrier effect in semiconductors, in which the free-carriers, which behave like particles in a neutral plasma, generate a negative polarizability at optical frequencies. By increasing or decreasing the free-carrier

concentration, one thus decreases or increases the index of refraction.

There are several features of the free-carrier effect that make it particularly attractive for integrated optic applications. One is the magnitude of the effect which, for reasonable carrier densities, is almost an order of magnitude larger than that of the electrooptic effect in semiconductors. Another is the increased design flexibility, i.e., the ability to design modulators in materials that have no intrinsic electrooptic effect, such as silicon. The third and perhaps most alluring feature, however, is the possibility of integrally combining this relatively large effect with the excellent high-frequency capabilities of certain semiconductor devices (such as MESFETS) in such a way that one has the best of both worlds: efficient modulation at very high microwave frequencies.

Historically, the scope of this contract was to do just that. In particular, the goal was to fabricate an optical modulator that used the depletion region of a microwave field effect transistor to modulate the free carrier density in the coupling region of an optical directional coupler. However, it was realized part way through the program that there were several insurmountable difficulties, both theoretical and fabrication, that prevented the successful development of such a device. In light of these findings, all work on this contract was stopped, and this final report documenting the work to date was prepared.

Although the contractual goal was not reached, various results were obtained that are of interest on their own. A theoretical model for an improved high-frequency semiconductor modulator that combined free-carrier and electrooptic effects was formulated, two types of current injection modulators (a variable directional coupler and a variable confinement waveguide) were fabricated and

tested, and a novel idea for a thermal switch was accidentally discovered and reduced to practice.

This final report, which shall discuss all of this work, is meant to be more than a mere compilation of monthly reports. The theory of free-carrier effects is briefly reviewed, the design for a Mach-Zehnder interferometer systematically developed, the benefits and shortcomings of free-carrier devices discussed, and the experimental findings presented, all in sufficient detail so that future experimenters in this field should be able to benefit from this acquired knowledge. We hope you find this report readable, entertaining, and instructive.

SECTION 2

FREE-CARRIER EFFECTS IN SEMICONDUCTORS

The electrons and holes in a semiconductor will move under the influence of an external electric field, with the positive charges (holes) moving away from the negative charges (electrons). For a static or low frequency field, the polarization produced by this motion will lag the applied field by 90° , leading to the resistive current flow normally associated with Ohm's law. As the frequency of the applied field is increased, however, a point is eventually reached, somewhere in the infrared, where the period of oscillation will be shorter than the time between scattering collisions. The motion of the electrons and holes will then be 180° out-of-phase with the applied field, yielding a negative polarization.

The magnitude of this effect can be determined by analyzing the response of classical unbound charged particles to an electric field. In the absence of collisions, the one-dimensional equation of motion for a free electron in a sinusoidally varying electric field is given by

$$m_e d^2x/dt^2 = -eE_0 \sin \omega t$$

where m_e is the electron mass, E_0 the field amplitude, and ω the frequency of oscillation. Integrating this directly, one has

$$x = eE_0 \sin \omega t / (m_e \omega^2)$$

Assuming that each electron is paired with a (noninteracting) stationary particle of opposite sign, so that charge neutrality is maintained, one can compute a dipole moment per unit volume

$$P = -nex = -ne^2E_0\sin\omega t/(m_e\omega^2)$$

where n is the electron density. If the stationary particles are allowed to move, or if mobile positive particles paired with stationary negative charges are added to this mixture, one will have an additional term, yielding

$$P_{fc} = P_n + P_p = -e^2E_0\sin\omega t\{n/m_e + p/m_p\}/\omega^2$$

where p is the positive particle density, and P_{fc} the total free-carrier (positive and negative particle) polarization. Now, from Maxwell's equation,

$$D = \epsilon_0 E + P = \epsilon_0 \epsilon E$$

where ϵ_0 is the free-space permittivity, ϵ the relative dielectric constant for the system, and P the sum of the material (i.e., the semiconductor crystal) and free-carrier polarizations, i.e.,

$$P = P_m + P_{fc}$$

Solving for ϵ , one has

$$\begin{aligned}\epsilon &= 1 + P/(\epsilon_0 E) = 1 + P_m/(\epsilon_0 E) + P_{fc}/(\epsilon_0 E) \\ &= \epsilon_m - e^2\{n/m_e + p/m_p\}/(\epsilon_0 \omega^2)\end{aligned}$$

where ϵ_m is the relative dielectric constant of the solid material in which the free carriers move. The index of refraction N , is found by taking the square root of the relative dielectric constant. Because the free-carrier polarization is

much smaller than that of the host material, one can expand the square root binomially, yielding

$$N = \sqrt{\epsilon} \simeq N_s - 0.5e^2\{n/m_e + p/m_p\}/(N_s\epsilon_0\omega^2)$$

This classical model can be applied to semiconductors by replacing the actual masses with the effective masses of the electrons and holes, respectively; the equation thus modified correctly describes the free-carrier index change due to intraband transitions. One sees that both the electrons and holes contribute to this change, although the hole contribution will usually be smaller due to the hole's larger effective mass. Both contributions are directly proportional to the carrier concentration.

There are two other quantum effects, however, that can significantly alter this result. The first is related to the interband transition between the light and heavy hole bands in GaAs. The second, and more significant of the two, is band-filling, often referred to as the Burstein-Moss effect, which occurs in all semiconductors. In this effect the presence of carriers in a semiconductor fills the low-lying quantum states in the conduction band, so that the "effective" band gap, i.e., that energy needed to bring an electron from the valence band into the conduction band, is increased. This increase in bandgap shifts the dielectric response due to the band transition upward in frequency, which has the effect of lowering the index of refraction at lower frequencies. Bandfilling, like the free-carrier plasma effect, thus decreases the index of refraction when the carrier concentration is increased.

An analysis which includes interband effects has been performed by Stern ⁽¹⁾, and one that includes band-filling effects has been

discussed by Mendoza-Alvarez et. al ⁽²⁾. The magnitude of these additive effects in GaAs at 1.3 μm wavelength are given in Table I. The net result of these corrections is to increase the magnitude of the free-carrier effect by roughly a factor of two. Note that, once again, all of these effects are proportional to carrier concentration.

The assumption of zero lattice scattering used above to derive the classical equation for the free carrier effect is not entirely correct. There will always be some loss associated with the motion of both electrons and holes, and this loss will produce an undesirable absorption of the optical signal as it passes through the material. Spitzer and Whelan⁽³⁾ measured the absorption coefficient for electrons in bulk GaAs, and found that

$$\alpha_n = 5 \text{ cm}^{-1} \quad (\lambda = 1.3 \text{ } \mu\text{m}, n = 10^{18}/\text{cc})$$

The absorption, which was anomalously constant between 1 and 5 μm wavelength, was thought to be due to excitation from the conduction band to higher lying minima. For longer wavelengths the absorption had the expected λ^3 dependence associated with impurity-scattering. Henry et al⁽⁴⁾ have measured the same parameter for holes, finding that

$$\alpha_p = 13 \text{ cm}^{-1} \quad (\lambda = 1.3 \text{ } \mu\text{m}, p = 10^{18}/\text{cc})$$

The absorption, which they attributed to intervalence band transitions, changed with wavelength, increasing to 25 cm^{-1} at 1.6 μm wavelength. As was the case for the real part, both α 's scale with carrier concentration.

The free-carrier effects discussed above are significantly larger than those associated with the electrooptic effect in semiconductors. They will only be useful, however, if there is a

Table I. Change in refractive index due to various effects. The values shown are for a free-carrier decrease of $10^{18}/\text{cc}$, and for a wavelength of $1.3 \mu\text{m}$. All terms have the same sign, so that the effects are additive.

9127-03-011

	n	p
INTRABAND	0.00350	0.00058
INTERBAND	~ 0	0.00070
BAND FILLING	0.00290	0.00120
TOTAL	0.00640	0.00250

way to rapidly change the free carrier density in the semiconductor. There are two techniques that can be used to do this. One is to "inject" free carriers into a confined volume by driving current through a PIN junction, as is done in a semiconductor laser. The other is to dope the material with either donors or acceptors, to form a Schottky or PN junction at the surface or at the PN interface, and to then reverse-bias this junction, so that a depletion region is formed which is devoid of mobile carriers. It is this second technique, charge depletion, that we shall examine first, returning to charge injection in a subsequent section.

SECTION 3

DEPLETION-EDGE MODULATORS

3.1 Depletion Edge Translation in an Optical Waveguide

An optical waveguide is formed when the cladding material surrounding the central guiding region has an index of refraction lower than that of the material in the guiding region. A single-mode glass fiber, for example, typically has a 5-10 μm circular core with an index that is roughly .002-.005 higher than the sheath material. The light within the core undergoes total internal reflection as it bounces off the core-cladding interface, ricocheting obliquely from one side of the core to the other as it propagates along the guide length.

In semiconductors, one is restricted to planar geometries, so that the waveguide core is typically rectangular or trapezoidal, not circular. One lowers the index above and below this core by using different materials, such as air above and AlGaAs below, and lowers the index to the left and right by thinning the top cladding and/or guiding layer, as shown in fig. 3.1a. The effect of thinning the stack of layers is to reduce the average index of refraction outside the guiding region.

If, in addition to using different materials, one also switches the dopant type from p to n as one passes downward through the center of the guiding region, a PN junction diode will be formed with a depletion region located symmetrically about the center of the guide (fig. 3.1b). The width of this depletion region can be changed by varying the voltage applied to the diode. This voltage variation "translates" the edges of the depletion region within the guiding region, so that the index of refraction

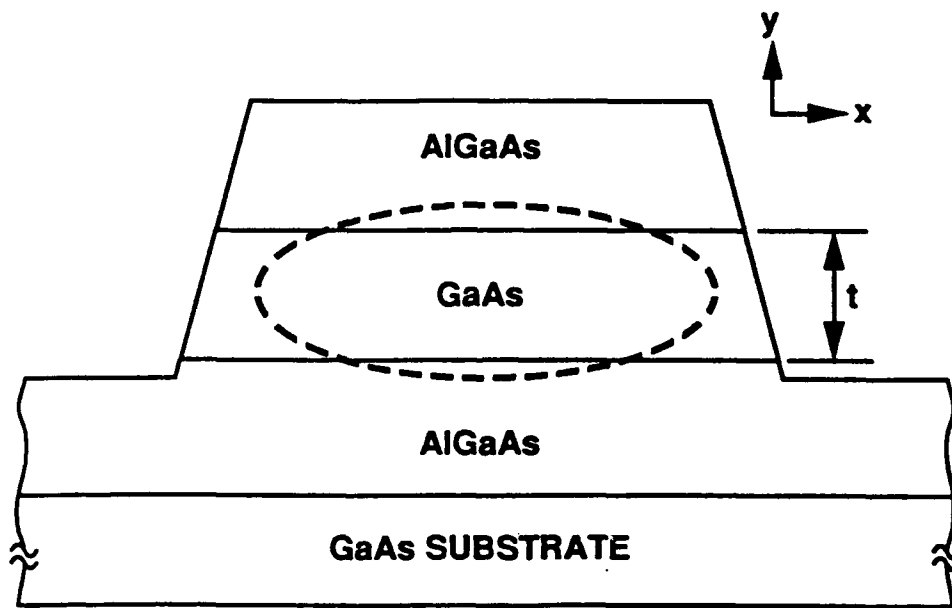


Fig. 3.1a Cross-sectional view of a semiconductor waveguide structure. The dotted ellipse surrounds the region of maximum optical field strength.

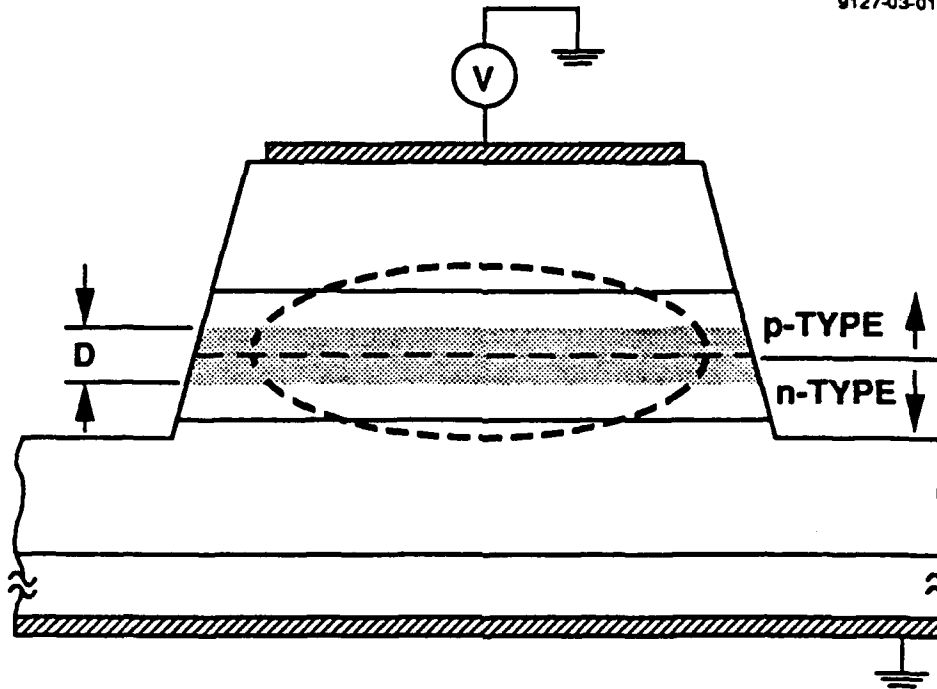


Fig. 3.1b The same structure with p and n-type dopants. The diode depletion region (the dotted area) is purposely located in the region of maximum optical intensity.

changes locally from that of a semiconductor with free carriers to that of a semiconductor without free carriers.

In order to change the width of the depletion region, one must either add or remove free carriers, which means that there will be a flow of current accompanying any change in this width. The parameter that determines just how much charge will be moved for a given change in junction voltage is the depletion layer capacitance, defined as

$$C = dQ/dV$$

For a symmetric PN junction in which the donor density of the N material is equal to the acceptor density in the P material, the depletion width is given by ⁽⁵⁾

$$D = (4\epsilon\epsilon_0(V_{bi}+V)/en)^{1/2}$$

where V_{bi} is the built-in potential, V the applied voltage, and n the donor (and acceptor) density. Using the chain rule, one has⁽⁵⁾

$$dQ/dV = (dQ/dD)(dD/dV) = (0.5nWe)(2\epsilon\epsilon_0/(enD)) = \epsilon\epsilon_0 W/D$$

where W is the width of the junction and D is the depletion width. Upon inspection, one sees that the incremental capacitance per unit length of a PN junction is identical to that of a parallel plate capacitor having identical dimensions and a plate separation equal to the depletion width, D , an interesting and somewhat unexpected result.

In order to calculate the efficiency of any modulator using this sandwich construction, one must be able to calculate the change in the effective index of refraction of the guided optical wave

for a given change in junction voltage. The parameter of interest is actually β , the propagation constant, defined by

$$\beta = (2\pi/\lambda)N$$

where λ is the free-space wavelength. For small changes in N , one can use perturbation theory to find the change in β .

Extending the results of reference 6 to two dimensions, one has

$$\delta\beta = \frac{(2\pi/\lambda) \int \Psi(x,y)^2 \delta N(x,y) dx dy}{\int \Psi(x,y)^2 dx dy}$$

where $\Psi(x,y)$ is the unidirectional (and typically transverse) optical electric field in the medium, and δN is the local change in index.

This equation has a very simple physical interpretation: the change in β is equal to the intensity-weighted average of δN over the mode multiplied by the free-space propagation constant. From this one sees that much greater changes in β are realized when the region of changing N is positioned at the field maximum. In fact, as pointed out in ref. 6, one achieves the maximum possible change by localizing δN in a small region at the center of the guide, rather than by spreading the change over the entire guide. As we shall see, the use of a PN junction allows us to do just that.

To proceed further, one must know the functional form for the fields in the guiding region. If the lateral index difference is small compared to the vertical, and the guiding layer thickness small compared to the guide width, a good approximation is one in which the field is separable, i.e.,

$$\Psi(x,y) = \Phi(x)\psi(y)$$

and in which Φ and ψ are the one-dimensional solutions for slab modes. This is the basis of the effective index approach, which is known to give good agreement with more exact techniques for guides with moderate confinement. With this simplification, the integrals separate into a product of integrals, one for each dimension.

Integration over x gives Γ , the fraction of optical power within the guide width. This term can easily be calculated using the effective index technique and the equations for the slab modes given below, yielding value close to unity (0.85 or larger) for most of the guide geometries that we shall consider. In order to simplify the equations that follow, we shall (somewhat arbitrarily) set Γ equal to one.

To integrate over y , one must use the functional forms for ψ , given by⁽⁶⁾

$$\psi = \cos(2uy/t) \quad |y| < t/2$$

$$\psi = \cos(u) \exp[w(1-2|y|/t)] \quad |y| > t/2$$

where

$$u = (t/2) (k^2 N_g^2 - \beta^2)^{1/2}$$

$$w = (t/2) (\beta^2 - k^2 N_c^2)^{1/2}$$

$$\tan(u) = w/u$$

t is the thickness of the guiding layer, and N_g and N_c are the indices of refraction in the guiding and cladding regions. The

last equation is the dispersion equation for the slab, which must be solved numerically to find β .

A further simplification results if the change in the depletion width for a given change in voltage is small. The region in which $\delta N(y)$ is non-zero will then be sufficiently narrow (fig. 3.2) that one can remove $\psi(y)^2$ from the integrand. One then has

$$\delta\beta = \frac{(2\pi/\lambda)\psi(D/2)^2 B dD}{\int \psi(y)^2 dy}$$

where B is the change in the index of refraction for a given change in carrier density ($B = dN/dn$), and dD is the change in the depletion width.

$$dD = (dD/dV)dV = (2e\epsilon_0/enD)dV$$

Evaluating this and also the integral in the denominator⁽⁶⁾, one has

$$\delta\beta_{fc} = (2\pi/\lambda)B(2e\epsilon_0/(enD))dVf_{fc}$$

where f_{fc} is a dimensionless integration factor given by

$$f_{fc} = \frac{\cos^2(uD/t)}{(t/2)(1+1/w)} \quad \delta < t$$

$$f_{fc} = \frac{\cos^2(u) \exp[2w(1-D/t)]}{(t/2)(1+1/w)} \quad \delta > t$$

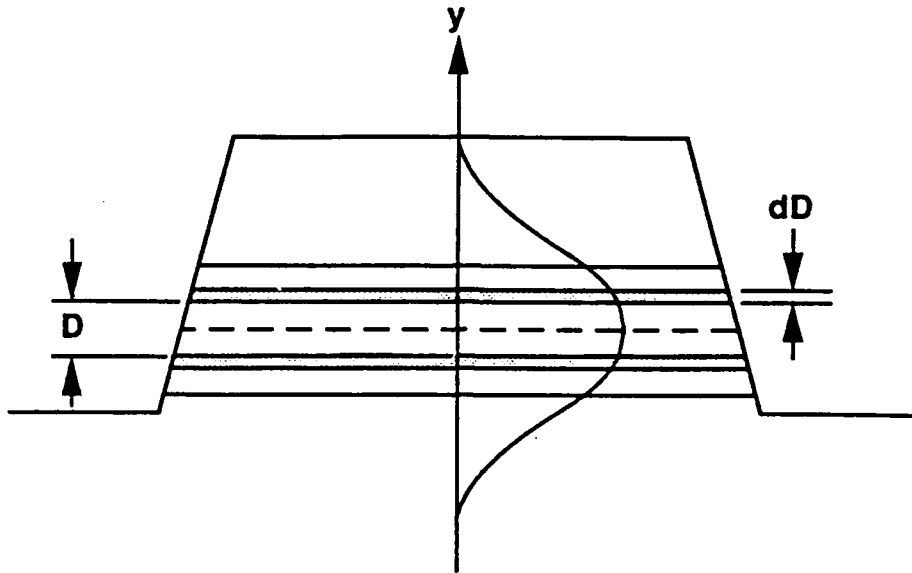


Fig. 3.2 Waveguide structure showing the change in the depletion width, dD , that results from a change in the bias voltage, and $\psi(y)$. Note that $\psi(y)$ is relatively constant across the region dD .

There will also be a contribution from the electrooptic effect for those materials having a finite electrooptic coefficient, such as GaAs. The electrooptically induced change in the index of refraction for crystals having the zinc blend structure (such as GaAs), for a field applied in the (100) direction, and for light propagating in either the (011) or (011) direction is given by

$$\delta N(y) = (1/2)rN^3E(y)$$

where r is the appropriate electrooptic coefficient, (in this case, r_{41}), and N the index of refraction of the host material.

The electric field in the depletion region is triangular, and is given by⁽⁵⁾

$$E(x) = -en(a-y)\epsilon\epsilon_0$$

$$a = (\epsilon\epsilon_0V/en)^{1/2} = D/2$$

However, as shown in fig. 3.3, the change in the field with applied voltage is constant across the region, a useful result which greatly simplifies the mathematics. Referring to fig. 3.3 and using the above equations, one has

$$dE/dV = (dE/da)(da/dV) = -1/D$$

so that the change in β is given by

$$\delta\beta = \frac{(2\pi/\lambda)((1/2)rN^3dV/D)\int\psi(y)^2dy}{(t/2)(1+1/w)}$$

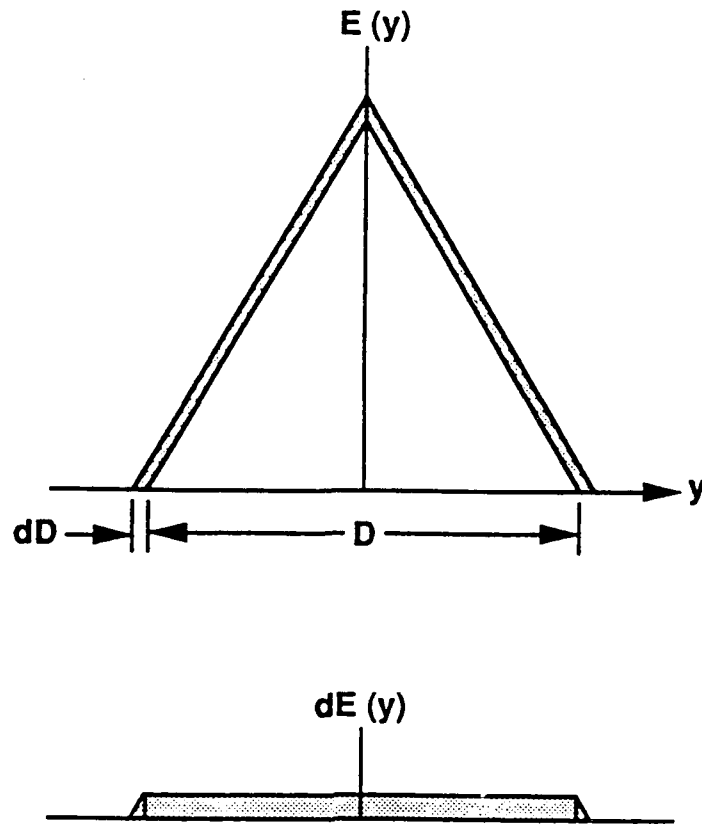


Fig. 3.3 Profile of the electric field within the depletion region. One sees that although the field $E(y)$ is triangular, the change in the field, $dE(y)$, for a given change in bias voltage is uniform.

Performing the required integration over the depletion region only, one has

$$\delta\beta_{\bullet\bullet} = (2\pi/\lambda)(0.5rN^3dV/D)f$$

where f is a unitless integration factor given by⁽⁶⁾

$$f_{\bullet\bullet} = \frac{D + (t/2u)\sin(2uD/t)}{t(1+1/w)} \quad \delta < t$$

$$f_{\bullet\bullet} = 1 - \frac{u^2 \exp[2w(1-D/t)]}{(w+1)(u^2+w^2)} \quad \delta > t$$

The change in β due to the free-carrier and electrooptic effects will be given by

$$\delta\beta = \delta\beta_{fc} + \delta\beta_{\bullet\bullet}$$

These effects can either add or subtract, depending upon the sign of the electrooptic coefficient and the direction of light propagation in the semiconductor crystal. In the discussions that follow, it will be assumed that this direction has been chosen so that the two effects add, thereby maximizing modulator efficiency.

3.2 OPTICAL MODULATORS

3.2.1 The Directional Coupler Switch

A simple directional coupler, shown in fig. 3.4, functions by means of the weak coupling between two closely-located optical waveguides. This small amount of coupling allows light from one guide to be transferred into the other. Because of the coherent nature of the process, and because of the subtractive effect the light in the second guide has on the first, all of the light will be drawn into the second guide after it has traveled a certain length known as the transfer length. The closer the guides, the greater the coupling, and the shorter the transfer length. Typical transfer lengths for integrated optic couplers can vary anywhere from 0.5 to 10 mm.

One could make a switch out of such a coupler by changing the coupling strength between the waveguides so that the actual length of the coupler corresponds first to one transfer length, then to two. The field patterns that would result for the "on" and "off" states for a device having 6 μm wide waveguides and a 4 μm (edge-to-edge) separation are shown in fig. 3.5; here the on-state corresponds to one transfer length, the off-state to two. Note that in the on-state, the power is transferred into the bottom guide, while in the off-state it remains in the top.

The degree of coupling is determined by the closeness of the waveguides and by the index of refraction between the guides. Although dynamically changing the physical location of the guides is obviously not practical, changing the index of refraction between the guides is, and one could do this using either the free-carrier or electrooptic effects discussed earlier. The basic drawback to this approach is that one is changing the index of refraction in the region between the guides, where the optical

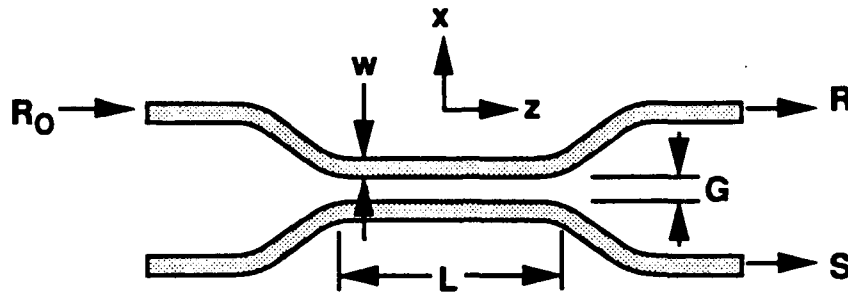


Fig. 3.4 Top view of a waveguide directional coupler.

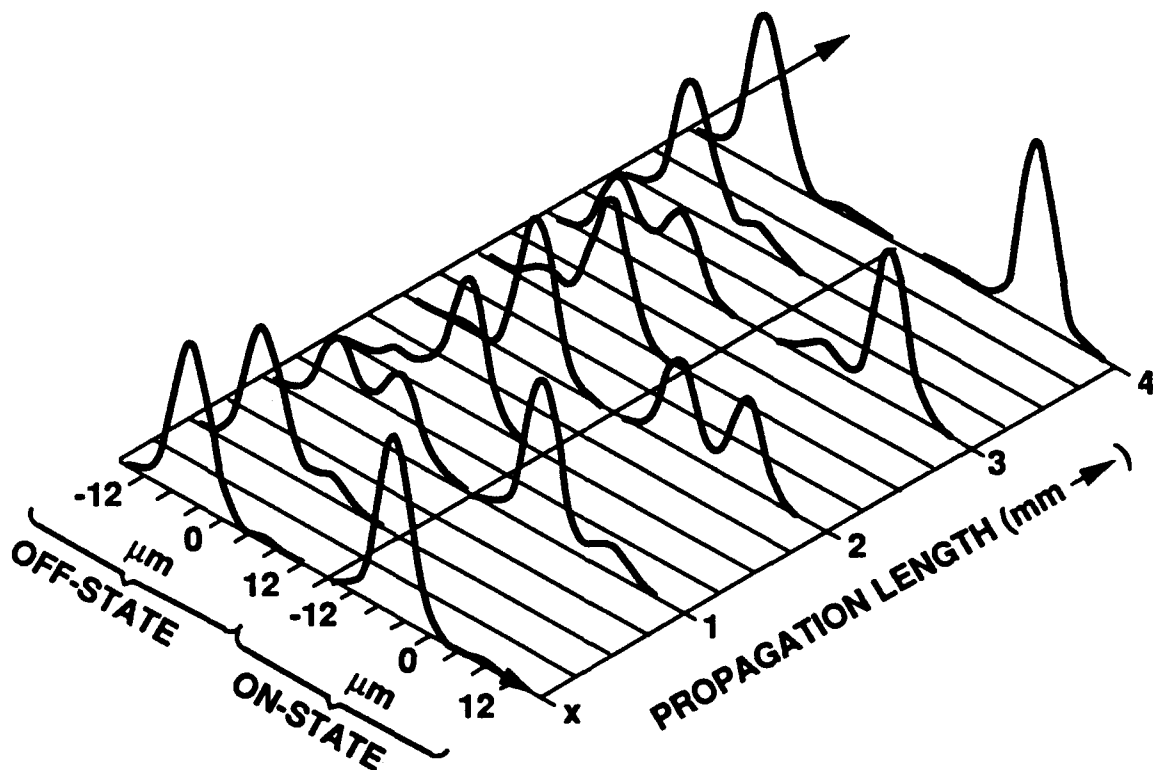


Fig. 3.5 Optical field intensities in the coupler of fig. 3.4 as a function of position along its length. Note that in the "on" state, light entering the top port of fig. 3.5 emerges from the bottom, while in the "off" state, just the opposite occurs.

intensity is relatively low (fig. 3.5). This is contrary to the results of section 3.1, which show that the greatest effect will result when one changes the index of refraction at an intensity maximum. One should thus expect that a directional coupler switch using variable coupling would be less efficient than other types of devices, and in fact such switches do indeed require lengths that are from 5 to 7 times longer than Mach-Zehnder modulators for equivalent depths of modulation. This type of device is, however, somewhat easier to fabricate than other types of modulators, relatively insensitive to the lossy nature of any components installed in the interguide region, and amenable to the traveling-wave geometries needed for high-frequency performance, all of which make it attractive for the two devices of the next section.

3.2.2 The RIMFET and RIMSCHOTT Concepts

The original aim of this contract was to develop an optical modulator that would use the depletion region under the gate of a microwave field effect transistor to modulate the coupling between two waveguides in a directional coupler switch via the free-carrier effect. Figures 3.6 and 3.7 show top and cross-sectional views of such a device, denoted RIMFET for Refractive Index Modulating FET, while fig. 3.8 shows a related device that uses a Schottky diode (RIMSCHOT) in place of a FET gate.

The intent was to marry the high-frequency properties of MESFETs with the large index-changing property of free-carriers in such a way as to produce a highly efficient Ka-band modulator. As contractual work commenced, however, it soon became obvious that this marriage was suffering from RH incompatibility. The very properties that allow MESFETs to operate at very high microwave frequencies - $0.1 \mu\text{m}$ gate widths and narrow pinch-off regions - are in opposition to those needed for control of optical signals, i.e., micron by multi-micron cross-sectional areas over which the index of refraction can be changed. Furthermore, because the FET is a majority carrier device, the drain-source current flow has no effect on the free-carrier concentration. Any change that occurs, occurs because of changes in the depletion volume beneath the gate. Because the depletion region beneath a FET gate has no particular advantage over the depletion region beneath a Schottky diode, the use of a FET structure is essentially superfluous. The only conceivable advantage to a FET structure would be that one might be able, with additional circuitry, to produce an active traveling wave structure that would have gain and would hence reduce the microwave signal level required to drive the device.

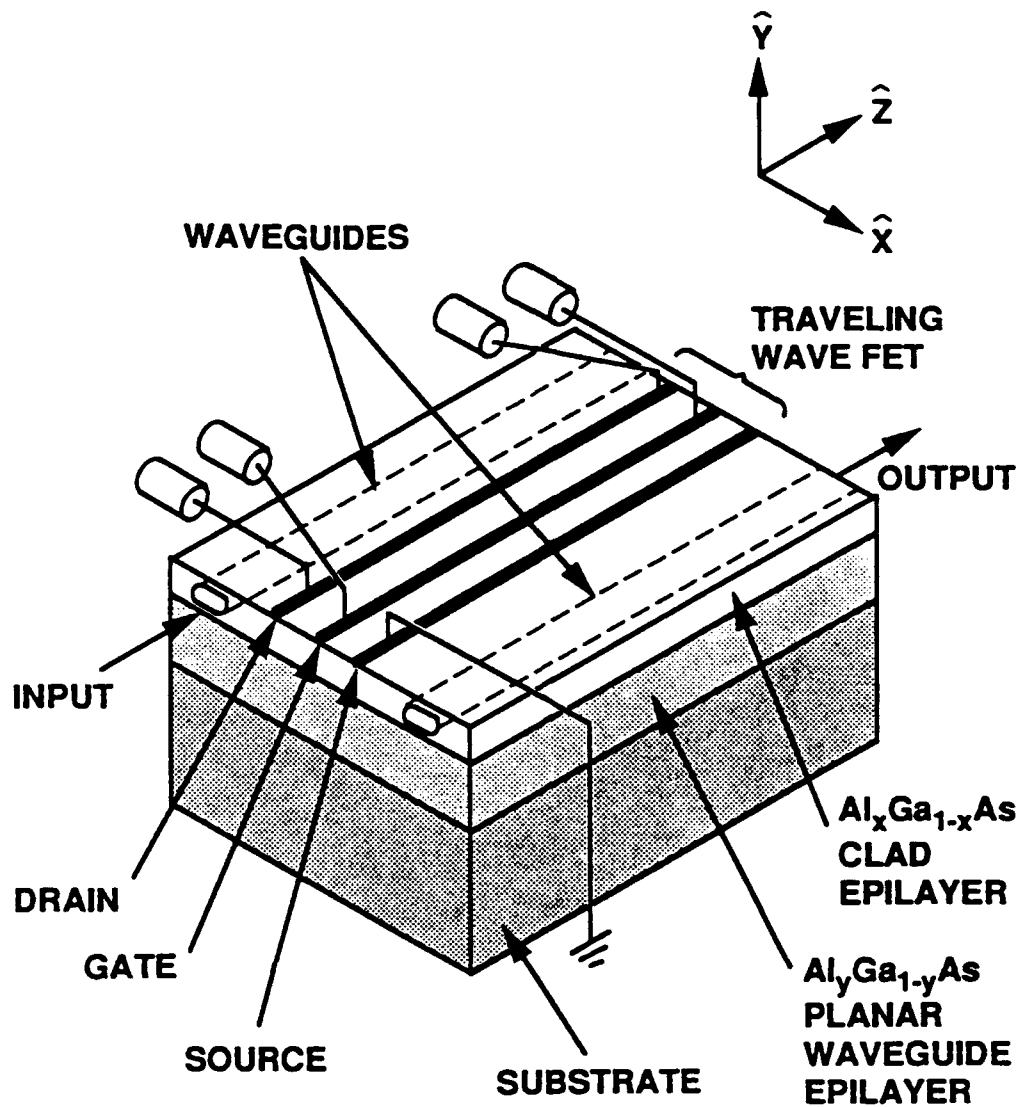


Fig. 3.6 The traveling wave RIMFET.

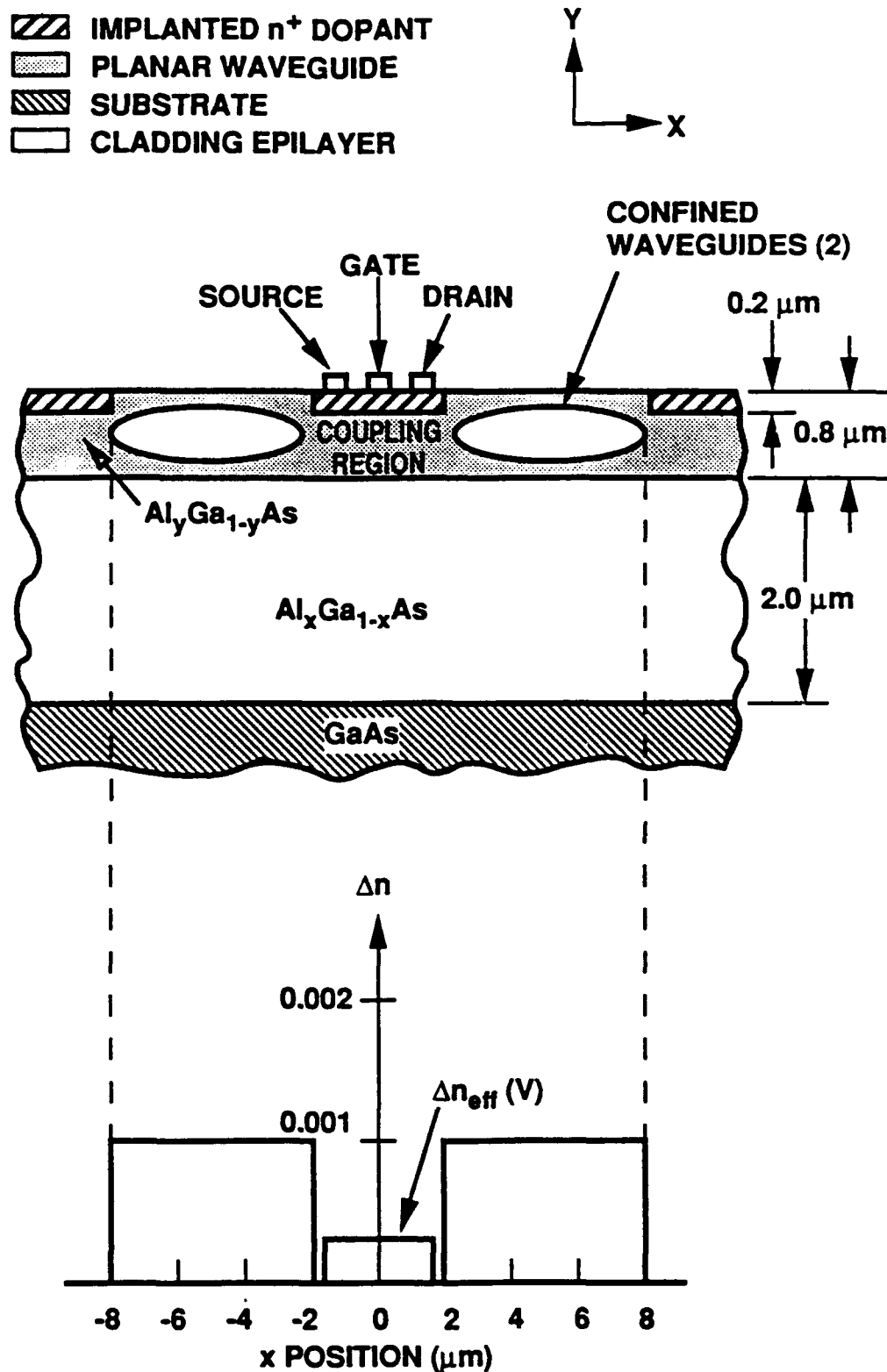


Fig. 3.7

Cross-sectional view of the RIMFET with the corresponding index profiles mapped below. Note that the effective index under the gate is a function of the applied voltage.

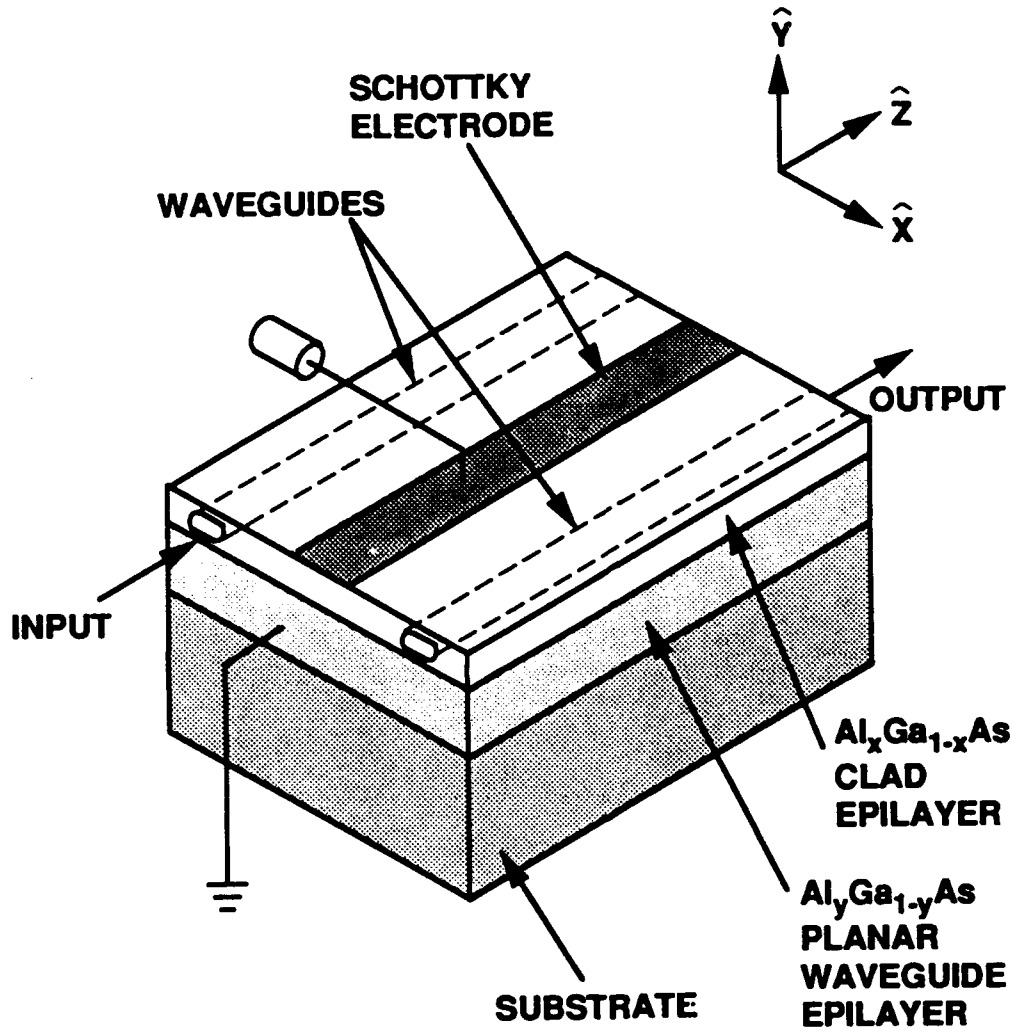


Fig. 3.8 RIMSCHOTT traveling wave modulator.

The RIMSCHOT concept, on the other hand, is not without merit, although it does have the following shortcomings:

1. The depletion region under the electrode of a heavily-doped Schottky diode is located at the top of the guiding region, so that changes in the carrier concentration occur where the optical field is at a minimum in the vertical direction (fig. 3.9), thereby yielding a minimal change in β .
2. Changing the index of refraction in the interguide region, is, as mentioned in section 3.2.1, an inefficient way to produce modulation, so that very long coupler lengths would be required for reasonable efficiency.

These problems, however, are not insurmountable. One might solve the first by changing the doping profile under the electrode so that the dopant concentration is very low adjacent to the electrode, and higher at the center of the guiding layer. This would move the depletion edge into the center of the guiding region, where further translation would have maximum effect. However, the incremental change in the depletion width for a given change in applied voltage would be small, so that overall modulator efficiency would suffer. A better approach would be to change from a Schottky to a PN diode structure, as in fig. 3.1b. This also moves the depletion region down into the volume of maximum optical intensity, so that one produces the largest $\delta\beta$ for a given change in depletion width. Furthermore, the change in the depletion width per unit applied voltage goes up as $1/D$, so that, in principle, high efficiencies are possible. We shall adopt this second approach in the analyses that follow.

The second problem can be solved by choosing the modulator configuration that requires the minimum $\delta\beta$ necessary for a given

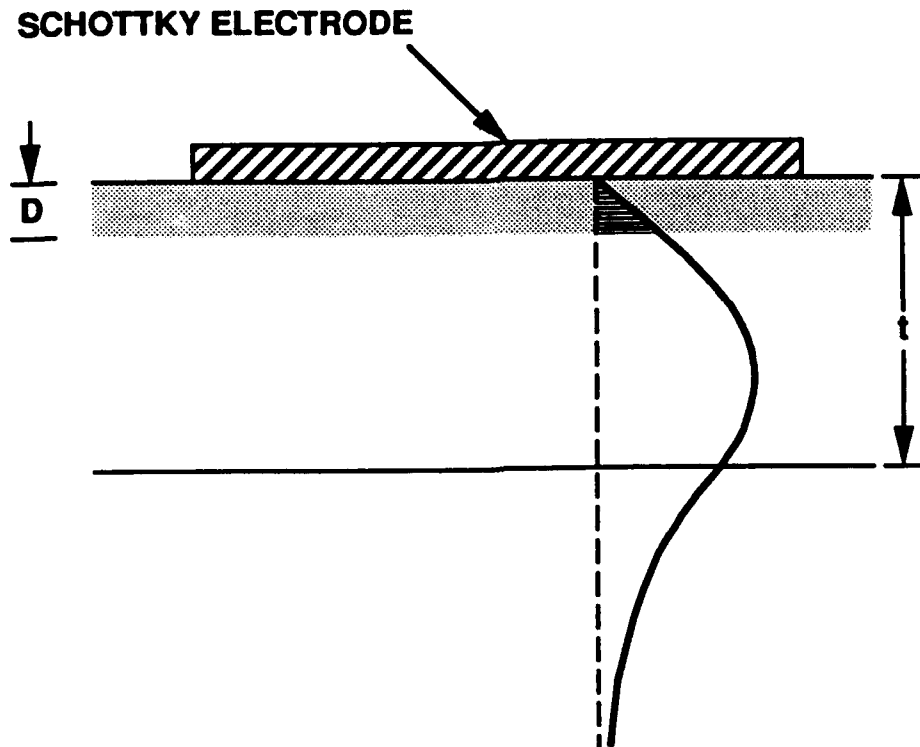


Fig. 3.9 Cross-sectional view of the interguide region of the RIMSCHOTT device, showing how the depletion region (the dotted area) is located near the minimum of the optical field strength, $\psi(y)$.

depth of modulation, and that will be compatible with the PN sandwich-type construction. That modulator is the Mach-Zehnder interferometer.

3.2.3 The Mach-Zehnder Traveling Wave Modulator

The Mach-Zehnder interferometer, shown schematically in fig. 3.10, is the integrated-optic equivalent of the Michelson interferometer. A guided optical signal is split at the first Y-junction, is oppositely shifted in phase in each arm, and is then summed at the second junction, where the two signals interfere constructively or destructively, depending on the relative phase shift in each arm. This interferometer is typically biased at $\pm 45^\circ$ for intensity modulation, and gives 50% depth of modulation for 15° excursions about this point.

At low modulation frequencies one generates the necessary phase shifts by changing the index of refraction uniformly in each arm. At higher frequencies, however, where the modulator length can be an appreciable fraction of the guided wavelength, one must use a traveling wave structure in which the microwave signal is launched at one end of the interferometer, and travels forward with the lightwave signal (fig. 3.11). If the two signals have identical propagation velocities, the phase shift in each arm will increase monotonically with length, so that longer modulators will produce greater modulation for a given power level. If the velocities are not equal, however, the signals become unsynchronized, and the latter part of the traveling wave structure cancels the modulation produced by the forward part, leading to a high-frequency rolloff. The 3 dB bandwidth for a velocity mismatched device is given by

$$\omega_{3 \text{ dB}} = 1.4 / (L(1/v_m - 1/v_l))$$

where v_m and v_l are the microwave and light propagation velocities, and L is the length of the phase shifting arms.

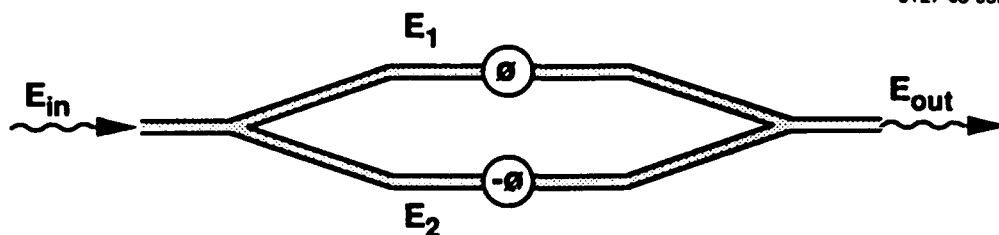


Fig. 3.10 A top-view of a Mach-Zehnder interferometer, showing the optical phase shifts ϕ that must be induced in each arm to obtain modulation.

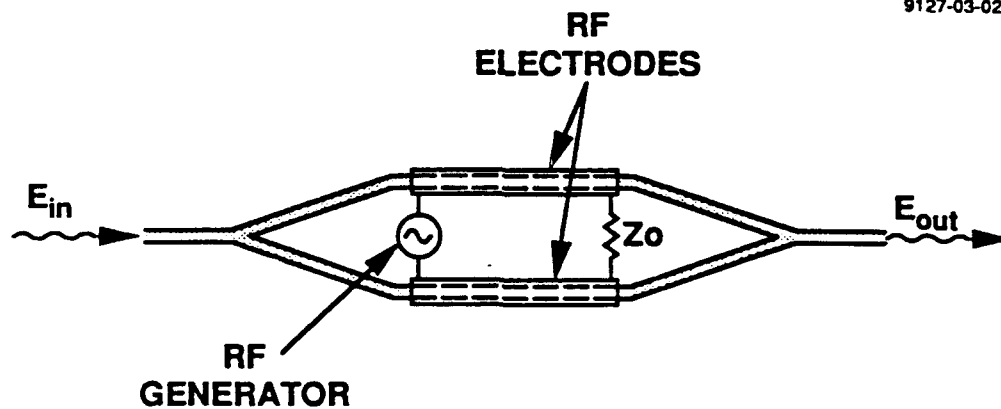


Fig. 3.11 The same Mach-Zehnder interferometer with traveling-wave electrodes. In this figure, both the optical and microwave signals propagate from left to right.

Obviously, then, one would like to match the velocities of the two signals whenever possible. The optical velocity is set by the index of the material, so that significant changes in this velocity are not possible. However, one can change the microwave velocity, and one does this either by moving the fields partially out of the guiding material, so that the microwave index of refraction is some combination of that of the material and that of air, or by adding lumped reactive elements.

A lossless transmission line is characterized by two reactances: the inductance per unit length, L , and the capacitance per unit length, C . From these one can find the characteristic impedance and the microwave propagation velocity, given by

$$Z_0 = (L/C)^{1/2}$$

$$v_m = (LC)^{-1/2}$$

By properly choosing L and C , one can adjust v so that it is equal to the speed of light in the medium, i.e.,

$$v_m = (LC)^{-1/2} = c/N$$

where c is the speed of light, and N is the optical index of refraction. This is the condition for velocity matching. We shall assure that this condition can be met, and shall defer till later a discussion of how this can be accomplished in practice.

In a practical implementation, the traveling wave structure of the Mach-Zehnder interferometer will be terminated with its characteristic impedance (fig. 3.11). If this structure is driven by a generator having the same characteristic impedance, the power dissipated will be given by

$$P = V^2/Z_0$$

where V is the rms drive voltage, and Z_0 is the characteristic impedance of the line. Substituting the equation for Z_0 , invoking the condition of velocity match, and using the value for the incremental capacitance of a PN junction derived earlier, one has

$$P = V^2/Z_0 = V^2/(L/C)^{1/2} = V^2 c C/N = V^2 c \epsilon \epsilon_0 W/(ND)$$

Now, the efficiency of a modulator will be determined by how much modulation (or phase shift) one gets for a given drive power. To calculate this, let us first replace the differential voltage of the earlier analyses with a sinusoidally varying signal, i.e.,

$$dV = \sqrt{2}V \exp(j\omega t)$$

where V is, again, the rms drive voltage. The rms phase shift in one arm of an interferometer of length l is then

$$\begin{aligned} \phi &= \delta\beta l = (2\pi/\lambda) \{ B(2\epsilon\epsilon_0/(enD))f_{tc} + (1/2)(rN^3/D)f_{\bullet\bullet} \} V l \\ &= (2\pi/\lambda) G V l \end{aligned}$$

Solving for V and substituting in the equation for P , one has

$$P = (c/N) (\epsilon\epsilon_0/D) (W/l^2) (\phi\lambda/(2\pi G))^2$$

One sees immediately that the required drive power goes up quadratically with required phase shift and optical wavelength, up linearly with guide width (W), and down quadratically with length. The other terms are a bit more complicated, however, and must be evaluated numerically. Note, interestingly enough, that

the free-carrier density, n , seems to have disappeared from these results. One must remember, however, that the depletion width, D , is a function of both V_{bias} and n , so that the dependence on n has simply been folded into D .

Figures 3.13 through 3.16 show the microwave power required for 50% depth of modulation ($\phi_{\text{rms}} = 10.6^\circ$ in each arm) as a function of depletion width for the three waveguide geometries listed in table II, for a 1 cm modulator length, for a lateral confinement factor (Γ) of unity, and for the waveguide dimensions shown in fig. 3.12. In each case the index of refraction of the cladding layers and the lateral width of the guide have been adjusted to give the largest possible mode size consistent with single-mode operation. The curves labeled EO show the drive power that would be required if the free carrier effect were absent (i.e., $B = 0$). This is also the power that would be required by a conventional EO modulator having a uniform field between two hypothetical plates separated by a fixed distance δ .

Several features are immediately obvious. The first is the overall lower drive power needed for guides with smaller dimensions, a direct consequence of the higher fields that result when voltage is applied over a smaller vertical dimension, and also of the lower capacitance associated with smaller guide widths (W). The second is the diminishing contribution of free-carriers for large values of D . This is also to be expected; for large D the depletion edge will have been moved well out of the high-field region, so that the changes in carrier density, which occur at this edge, will now be occurring where the fields are weak, and the effect on β will be minimal.

The third, and perhaps most interesting, however, is the prediction of zero drive power for vanishingly small D . This is due to the fact that the change in depletion width with voltage

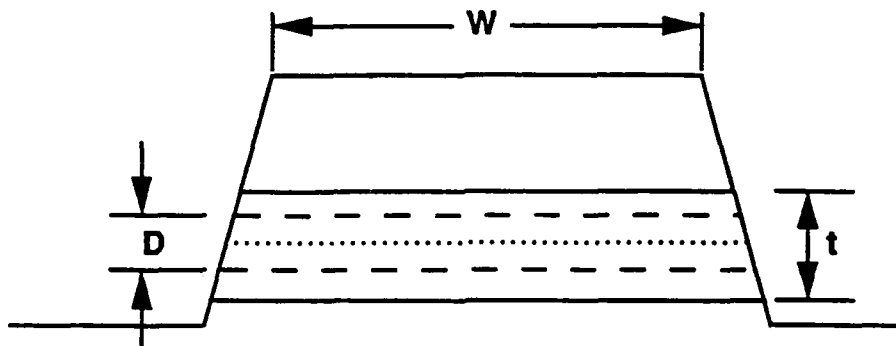


Fig. 3.12 Waveguide dimensions: t is the thickness of the guiding layer, D is the depletion width, and W is the waveguide width.

Table II. Waveguide parameters for three different modulators. N_g and N_c are the refractive indices of the guiding and cladding regions, respectively.

MODULATOR	W	t	N_g	N_c
1	1.1	0.5	3.45	3.20
2	2.5	1.2	3.45	3.40
3	4.5	2.2	3.45	3.435

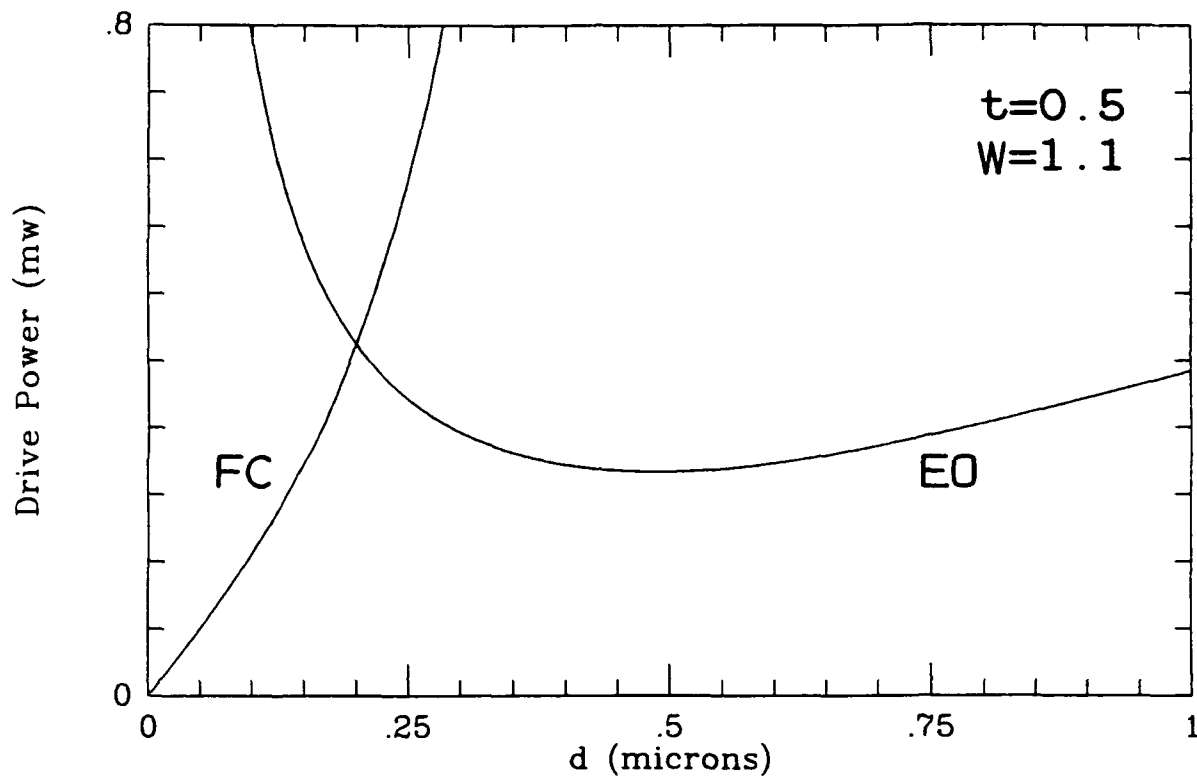


Fig. 3.13 Drive power required for 50% depth of modulation for a 1.0 cm long device using either the free-carrier (FC) or the electro-optic (EO) effect, but not both. The waveguide dimensions are those of example 1 in table II.

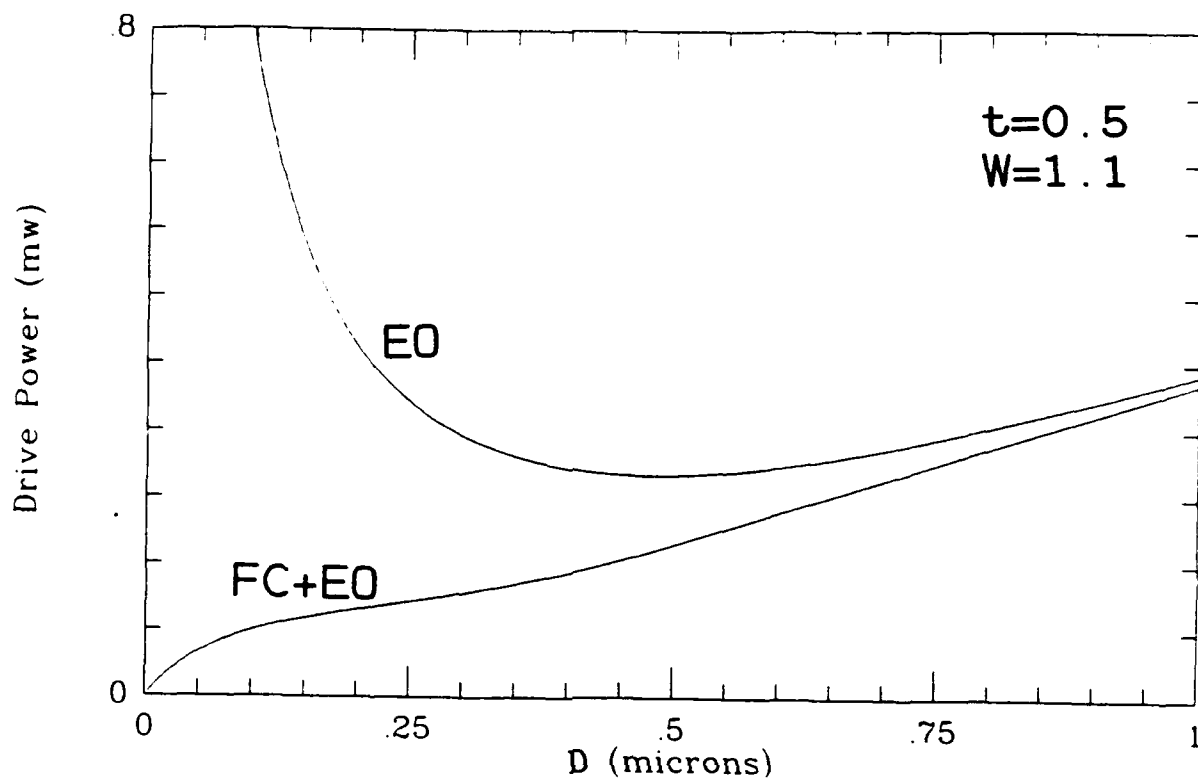


Fig. 3.14 Same as fig. 3.14, but including the EO effect in the FC results (as would happen in a real device).

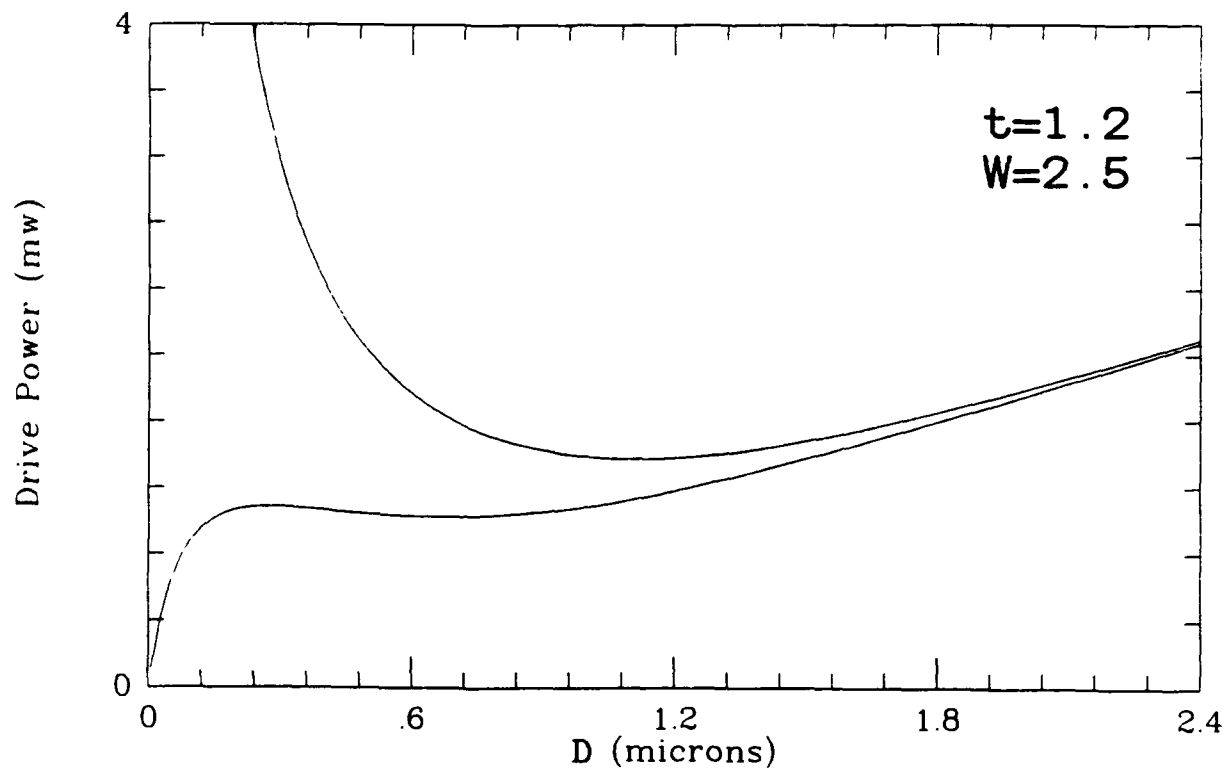


Fig. 3.15 Required drive power for the device of example 2, table II.

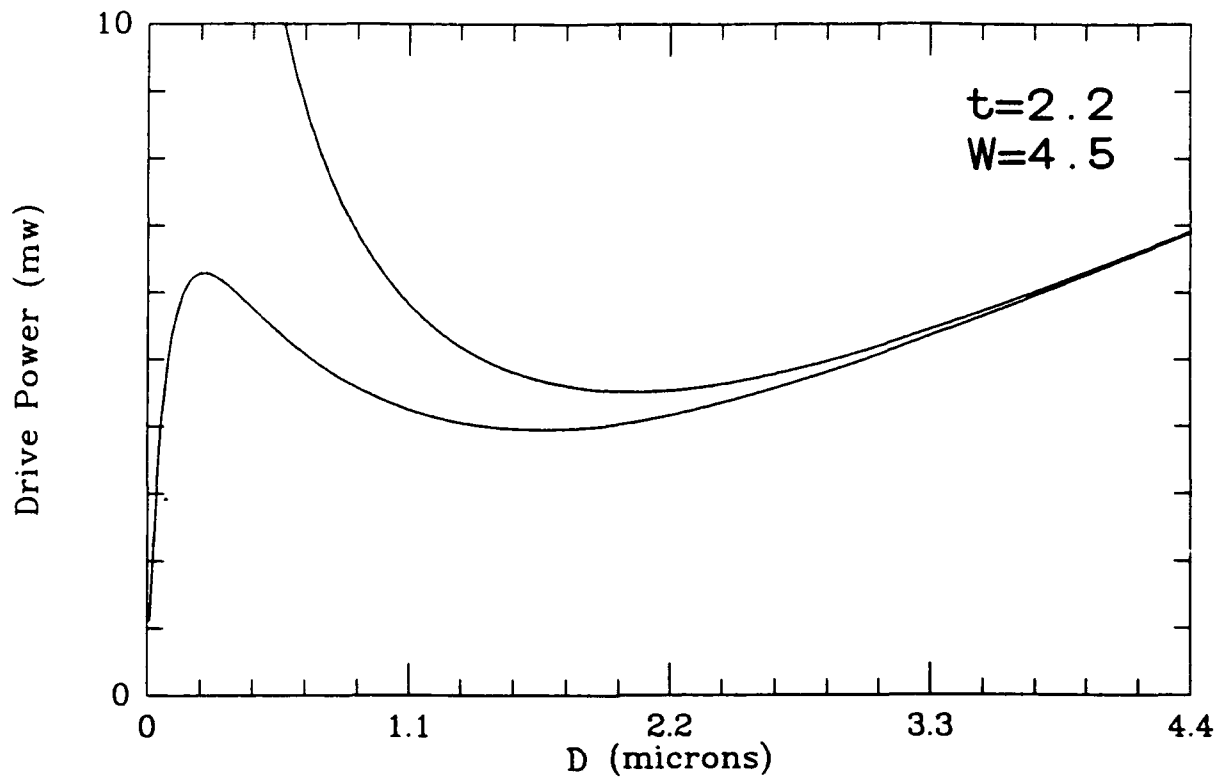


Fig. 3.16 Required drive power for the device of example 3.

is proportional to $1/D$, so that the smaller D , the lower the voltage needed to effect a given phase change. Unfortunately, this operational point cannot be reached in practice for a number of reasons, the primary one being the difficulty of designing a suitable low-loss traveling wave structure for D/W values of less than 0.25. This limitation will be discussed in the next section.

3.2.3 Transmission Line Design

A proposed traveling wave geometry that uses a Coplanar Strip Line (CPS) in a Mach-Zehnder configuration is shown in figs. 3.17 and 3.18. The two arms of the interferometer consist of the PN junctions discussed earlier, with ohmic contacts to the tops, and a heavily doped ($n \approx 10^{18}$) strip connecting the bottoms. This connecting strip, used by Walker et al⁽⁷⁾ to fabricate high-frequency electrooptic modulators in GaAs, allows one to connect the capacitances of each arm in the push-pull configuration of fig. 3.19, so that the RF voltage will bias one junction in the forward direction, and the other in the reverse direction, thus giving the opposite phase changes necessary for modulation. DC biases would be applied using the through-the-substrate technique of Walker⁽⁸⁾, although it might be possible to use the built-in voltage of each junction to eliminate the need for such a bias. In this latter case, external phase trimmers would then have to be used to provide the static 90° phase offset between arms necessary for intensity modulation. The entire structure would be located on top of a slab of semi-insulating GaAs.

Referring to the equivalent circuit of fig. 3.19, one sees that the total capacitance per unit length is given by

$$C_{total} = C_{cps} + C_{pn}/2$$

where C_{pn} is the capacitance of the waveguide's PN junction and C_{cps} is the capacitance of a CPS transmission line having conductors of width w and a distance between conductors of s . C_{cps} is given by⁽⁹⁾

$$C_{cps} = 0.5(\epsilon_r + 1)\epsilon_0 K'(k)/K(k)$$

$$k = s/(s+2w)$$

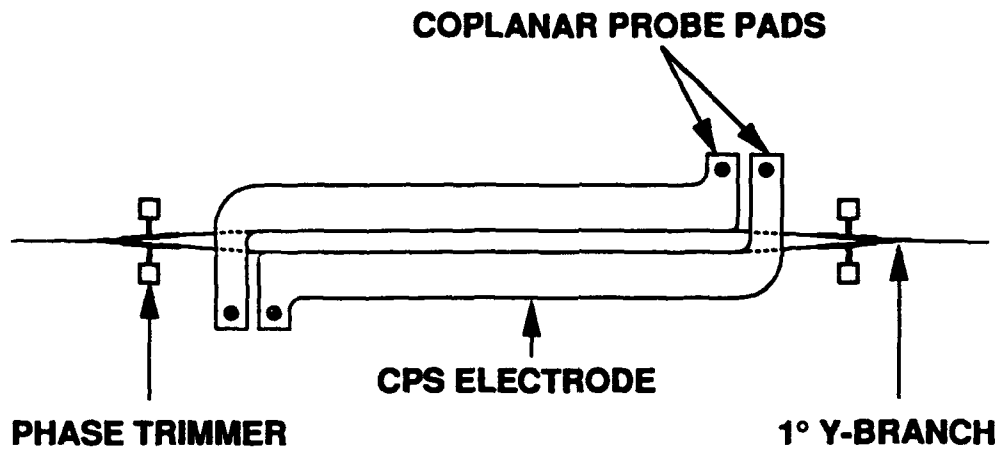


Fig. 3.17 A Mach-Zehnder modulator using coplanar strip (CPS) electrodes. The small phase trimmers would be needed only if the device were operated in the self-bias mode (see text).

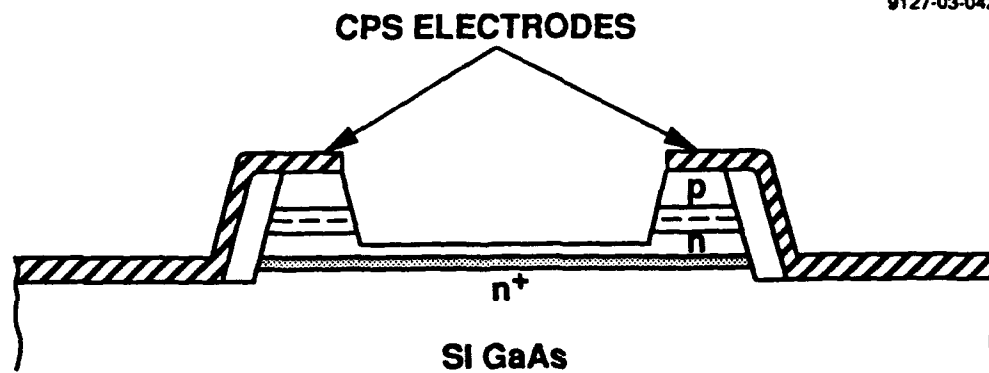


Fig. 3.18 Cross-sectional view of the CPS device of fig. 3.17. Electrical contact between the two arms is provided by the n^+ strip connecting the two waveguides.

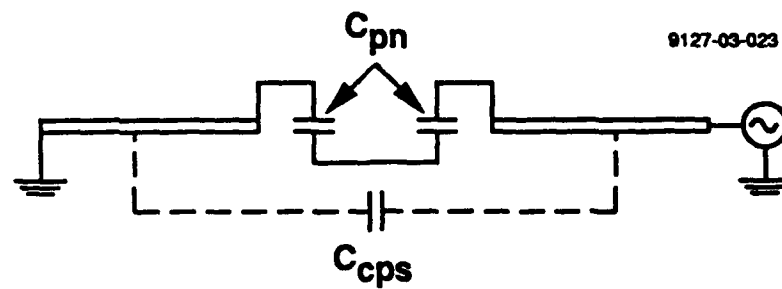


Fig. 3.19 Equivalent circuit of the CPS structure of fig. 3.18

$K(k)$ is the complete elliptic integral of the first kind, and

$$K'(k) = K(k'), \quad k' = (1-k^2)^{1/2}$$

The ratio of elliptic functions can be approximated by⁽⁹⁾

$$\begin{aligned} K(k)/K'(k) &= (1/\pi) \ln\{2(1+\sqrt{k})/(1-\sqrt{k})\} && \text{for } .707 < k < 1 \\ &= \pi / (\ln\{2(1+\sqrt{k}')/(1-\sqrt{k}')\}) && \text{for } 0 < k < .707 \end{aligned}$$

The inductance of the CPS line is determined by invoking the fact that the velocity must equal the speed of light when the relative permittivity, ϵ_r , goes to one. This yields

$$L = K(k)/(K'(k)\epsilon_0 c^2) = \chi/(\epsilon_0 c^2)$$

where χ denotes the ratio of the elliptic functions. Assuming that the inductance of the junction is negligible, so that the total inductance is just that of the CPS line, substituting the value for C_{pn} , and requiring that the microwave and optical velocities match, one has

$$(N/c)^2 = LC_{total} = (\chi\epsilon_0/\epsilon_0 c^2) \{(\epsilon_r+1)/(2\chi) + \epsilon_r W/(2D)\}$$

For the case of GaAs, where $\epsilon_r = 13.1$ and $N = 3.45$, this reduces to

$$\chi = K(k)/K'(k) = .741 D/W$$

One can show that, for small values of D/W

$$k \simeq 4\exp(-2.12W/D)$$

$$w/s \approx .125 \exp(2.12W/D)$$

This last result shows why, as stated earlier, it is difficult to design CPS transmission lines with D/W less than 0.25.

Substituting 0.25 in the equation, one has

$$w/s = 602$$

For a typical guide separation distance of about $20 \mu\text{m}$, one would have electrodes with widths of 1.2 cm, a dimension which is quite large for integrated-optic applications.

The reason for the lower limit on D/W can be understood as follows. A smaller D means a larger capacitance, which in turn must be offset by a decrease in inductance in order to maintain a velocity match. Lowering the inductance is achieved by making the conductor broader. However, because inductance goes as the inverse logarithm of w , a small decrease in L requires an exponential increase in width, which soon leads to physically unrealizable dimensions.

A small depletion width could have other undesirable consequences as well, such as a low line impedance and consequent high conductor losses. However, the CPS geometry, with its wide electrodes, minimizes the effect of loss, and the choice of suitably small waveguide dimensions will give reasonable impedances. The geometries of examples 1 and 2 in table II, for example, have characteristic impedances of 45 and 20Ω for a D of $0.25 \mu\text{m}$, values easily matched to an external 50Ω line. The losses are also quite reasonable. For a general transmission line, this loss is given by

$$\alpha(\text{dB/cm}) = 8.7 R/2Z_0$$

where R is the resistance per unit length of the conductors. For millimeter widths, copper electrodes, and a skin depth of $0.3 \mu\text{m}$, this value will be less than 0.2 and 0.5 dB/cm for said examples.

The limitation on D/W derived above is for this particular geometry. One could reduce this limit somewhat by using some other material, such as glass, for the substrate, so that C_{sp} would be significantly smaller. Doing this, however, would only reduce the allowable D/W ratio to 0.19, a relatively insignificant improvement.

Another approach might be to stack several PN diodes vertically, thereby reducing the capacitance and spreading the effect out over a larger volume. Unfortunately, a simple PNP...PN arrangement does not produce any change in free-carrier concentration beyond that produced by one PN junction. If the first junction (PN) is reversed biased, the second (NP) will be forward biased, so while one junction width is being increased, the other is being decreased, and no net change in free-carrier density would occur for the pair. Even if one could somehow "destroy" the matching NP junctions, the problem would still exist. Basically, holes and electrons would flow and reside near the former location of the NP junction under changing reverse bias. These free-carriers, removed from one region, would thus exist in an adjacent region, so that no net change in free-carrier concentration over the total volume would occur.

3.2.4 Free Carrier Losses

The above results have assumed that the free carriers do not introduce any type of loss, which is not the case. In section 2 we saw that both electrons and holes introduced optical loss, with values approaching 20 and 54 dB/cm, respectively, at concentrations of 10^{18} /cc. In fact, the loss for holes is sufficiently high that one might want to consider making an absorption modulator using free carriers! (This is, by the way, a distinct possibility. At $1.6 \mu\text{m}$, the absorption increases to 108 dB/cm, so that 10^{17} /cc concentrations over a 1 cm length could vary the amplitude by 10 dB.) Aside from optical losses, however, the finite conductivity introduced by free carriers at microwave frequencies can lead to conductive losses in the semiconductor that will attenuate the high-frequency traveling wave, thereby reducing the effective length of the device and raising the required drive power for a given depth of modulation.

These conductive losses can be most easily understood by considering the PN sandwich of fig. 3.20a. For very high dopant concentrations, the doped regions between the electrodes and the depletion region will have very high conductivity, and will thus appear as zero-reactance shunts connecting the electrodes to the depletion capacitance. As the concentration is lowered, the shunting reactance becomes primarily resistive, introducing loss. For very low concentrations, these regions become lossy capacitors in series with the depletion capacitance.

One can determine the attenuation factor for the modulator by first calculating the complex susceptance of the PN sandwich, and then inserting this into the transmission line equation for the propagation constant. The resistance and capacitances per unit length of the various elements in fig. 3.20b are given by

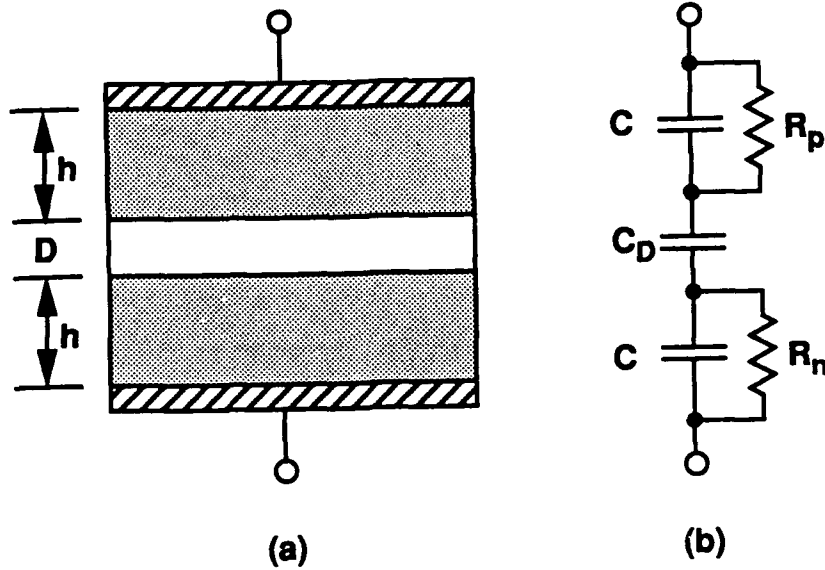


Fig. 3.20 Cross-sectional view of a PN stack (a) and the equivalent electrical circuit (b).

$$R_i = h/(W\sigma_i)$$

$$C = \epsilon\epsilon_0 W/h$$

$$C_d = \epsilon\epsilon_0 W/D$$

where σ_i is the conductivity of either region (i.e., $i = n$ or p). Using simple circuit theory, one sees that the total admittance of the stack is given by

$$\begin{aligned} Y_s &= (j\omega\epsilon\epsilon_0 W/D) \{1 + (h/D) (1 + \sigma_p/(j\omega\epsilon\epsilon_0))^{-1} + (h/D) ((1 + \sigma_n/(j\omega\epsilon\epsilon_0))^{-1})\}^{-1} \\ &= (j\omega\epsilon\epsilon_0 W/(D(1+\eta))) \end{aligned}$$

which is seen to be the admittance of the PN junction capacitance modified by a unitless complex factor. The total admittance of the transmission line will then be

$$\begin{aligned} Y_t &= j\omega C_{pn}/(2(1+\eta)) + j\omega C_{cps} \\ &= (j\omega\epsilon\epsilon_0 W/(2D(1+\eta))) + j\omega\epsilon_0(\epsilon_r+1)/(2\chi) \end{aligned}$$

The propagation factor is given by

$$\gamma = \alpha + j\beta = (ZY)^{1/2}$$

where Z is the impedance per unit length, $j\omega L$. In the limit of infinite conductivity, this will reduce to the velocity-matched value of

$$\gamma = j\beta = j(2\pi/\lambda)N$$

where N is again the optical index of refraction. Expanding in powers of η , and assuming that $\sigma \gg \omega\epsilon\epsilon_0$, one finds that, to first order, the power attenuation coefficient is given by

$$\alpha_p = 2\alpha \simeq N\omega^2(\epsilon\epsilon_0/c)(h/D)(\rho_p + \rho_n)/[1 + D(\epsilon+1)/(\epsilon W\chi)]$$

where the ρ 's are the resistivities ($\rho=1/\sigma$) and c is the speed of light. The term in the denominator accounts for the fact that only part of the electromagnetic energy is stored in the junction capacitance, so that the losses are reduced accordingly.

The attenuation at 30 GHz for the GaAs transmission line of fig. 3.18, with p and n dopant concentrations of $10^{17}/\text{cc}$, a χ of .25, and an h/D ratio of 3, is

$$\alpha_p = 7.6 \text{ dB/cm}$$

This loss is due almost entirely to the holes, which have a conductivity that is roughly 1/20 that of the electrons.

These results are for the transverse component of the electric field only, and do not include the losses due to the longitudinal component that will exist for inhomogeneous structures. The origin of the longitudinal field has to do with the fact that regions with moderate carrier densities (10^{15} to $10^{18}/\text{cc}$) have a conductivity that is large enough to electrostatically exclude the electric field, but not large enough to support the currents necessary to exclude the magnetic field. The magnetic field thus threads its way through these conductive regions of the semiconductor just as it would in the absence of carriers, satisfying Maxwell's equations by producing a small longitudinal electric field.

The losses associated with this longitudinal field increase with increasing conductivity, in direct contrast to the transverse losses, (which decrease with increased conductivity), and are proportional to the square of the magnetic field. For reasons having to do with the flow pattern of the magnetic field lines and the moderate conductivities used in the PN junctions, the waveguide PN stacks will not contribute significantly to this

type of loss. However, the n^+ strip that connects the two stacks has a relatively high conductivity, and is located where the magnetic field is maximum. An appreciable longitudinal current will thus flow in this strip, causing measurable loss.

The magnitude of this loss can be calculated by making a few simplifying assumptions. The first is that the magnetic field pattern in the gap between capacitors is essentially that of a CPS line without the capacitors and the n^+ strip. The second is that the transverse electric field in the gap is constant, and is given by

$$E_x = V/d$$

where V is the line drive voltage, and d is the gap spacing (see fig. 3.21). (We shall see later that this approximation of a uniform field, rather than underestimating, actually overestimates the loss by about a factor of two.) The magnetic field in the gap will be orthogonal to the electric field, and will, for this particular geometry, have a magnitude given by

$$B_y = (1/v)E_x = (N/c)E_x$$

where v is the propagation velocity which, because of velocity matching, is equal to c/N . Now, in the n^+ strip, E_x , the transverse electric field will be essentially zero. The presence of a time-varying magnetic field in a region with zero electric field, however, is a violation of Maxwell's 3rd equation, which states that

$$-dB/dt = \text{curl } E$$

For the geometry of fig.3.21, where B is in the y -direction only, this reduces to

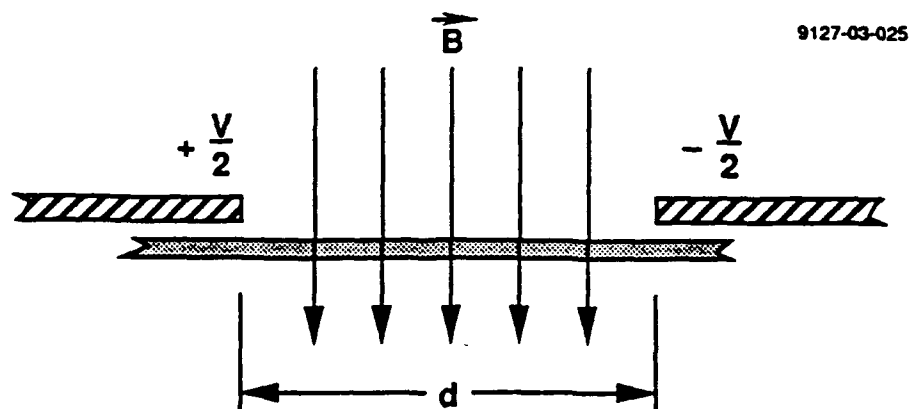


Fig. 3.21 Simplified cross-sectional view of the electrode region of the CPS interferometer showing the location and direction of the B-field in the n^+ strip (the dotted region).

$$-j\omega B_y = \partial E_x / \partial z - \partial E_z / \partial x$$

For a normal TEM mode in carrier-free material, this equation would be satisfied by the first spatial derivative, i.e., by a large transverse field (E_y) that varies slowly in the z direction. However, because the conductivity forces this term to zero, one must now have a finite second term, i.e., a small longitudinal field (E_z) that varies rapidly in the x direction. One thus has

$$-j\omega B_y = 0 - \partial E_z / \partial x = - \partial E_z / \partial x$$

Substituting and integrating, one has

$$E_z = (N\omega/c)Vx/d$$

The power dissipated per unit length is then

$$P = (1/2) \int J \cdot E d\tau = (1/2) \sigma \int E_z^2 d\tau$$

Performing the integration over the strip volume, and using the fact that the power attenuation coefficient is just the dissipation per unit length divided by the power flow, one has

$$\alpha_p = P_{dis}/P_{flow} = P_{dis}/(V^2/2Z_0) = \sigma(N\omega/c)^2 t d Z_0 / 12$$

where t is the n^+ strip thickness, σ the conductivity of the strip, and Z_0 the impedance of the CPS transmission line (with capacitors). A more exact approach gives an almost identical result, but with the factor $1/12$ (.083) replaced with .047. Using this more exact result, together with an n^+ concentration of $2 \times 10^{18}/cc$, a strip thickness of $1 \mu m$, a gap spacing of $20 \mu m$ and a Z_0 of 50Ω , one has

$$\alpha_p = 1.92 \text{ dB/cm}$$

3.2.5 Predicted Modulator Bandwidth

The frequency of a semiconductor Mach-Zehnder modulator is limited by the degree of velocity mismatch, by the transit time of the electrons and holes across the depletion region, and by the frequency-dependent microwave losses in the traveling wave structure. The first of these has been eliminated by matching the microwave velocity to that of the lightwave, as discussed in section 3.2.3. The second will be determined by the saturation velocity of electrons and holes in the semiconductor material; for GaAs, this value is 6×10^6 cm/s, so that for the roughly $0.1 \mu\text{m}$ widths of the P and N sections of the depletion region, one has a transit time of about 2 picoseconds, which corresponds to an 80 GHz bandwidth. The third mechanism, microwave loss, limits the bandwidth by reducing the effective length of the modulator at higher frequencies, thereby reducing the degree of modulation.

In an earlier section we saw that the total optical phase shift in one arm of the modulator was given by $\phi = \delta\beta l$, where $\delta\beta$ was the incremental change in the propagation constant, and l was the modulator length. When microwave loss is included, the total phase shift will be given by

$$\phi = \delta\beta \int \exp\{-\alpha_p(\omega)z/2\}dz$$

where integration is over the length of the modulator. At low frequencies, where α is zero, ϕ will again be $\delta\beta l$. For higher frequencies, however, ϕ will be reduced. Performing the integration, setting ϕ equal to 70% of its DC value, and solving the resulting transcendental equation for αl , one finds that the 3 dB rolloff point for a modulator of length l will occur when

$$\alpha_t(\omega)l = 1.48$$

where $\alpha_t(\omega)$ is the sum of the various loss processes. Using this relation and the results of section 3.2.4, one finds that a 5 mm long modulator having the waveguide dimensions of example 1, table II, and having the various carrier concentrations and dimensions given for the examples of transverse and longitudinal loss, would have a 35 GHz bandwidth and a drive power requirement of 3.2 milliwatts (+ 5 dBm).

Higher bandwidths could be obtained by shortening the length of the modulator. Doing this, however, raises the required drive power quadratically with inverse length. Aside from requiring more powerful drive amplifiers, this might lead to other undesirable effects that have not been discussed here, such as the generation of 2nd harmonics and 3rd order intermodulation products due to nonlinearities of the PN junction. One could also increase the bandwidth by increasing the PN carrier concentrations to reduce the transverse losses, but only at the expense of decreased optical throughput due to the increased optical losses associated with free-carriers. The $10^{17}/\text{cc}$ concentration used in the transverse loss example is already causing a 3.5 dB/cm optical loss which, although frequency independent, is nonetheless objectionable.

3.2.6 Summary of Findings for Depletion Edge Translation Modulators

The predicted performance of a free-carrier modulator using a variable-width depletion region of a PN junction shows that these devices, when properly designed, are capable of reaching Ka-band frequencies with reasonably low (3 mw) drive power. There are, however, several disadvantages to using the free-carrier depletion effect.

The first is that the very free carriers which provide the index-changing effect also provide the loss mechanism which ultimately limits the high-frequency performance. This limit, caused by the poor conductivity of the holes in the p-layer, could be raised threefold by converting to a device that used n-type dopants only. However, we have not yet found a geometry that would allow this without introducing unacceptable optical losses or excessively high drive power requirements.

The second disadvantage is that fundamental transmission line limitations prevent operation at the smaller depletion widths that would provide greater operating efficiency. The large capacitance associated with smaller depletion widths slows the microwave signal, eventually preventing velocity matching, so that the free-carrier contribution to the phase shift in a velocity-matched modulator can at best equal to, and is often less than, that due to the electrooptic effect. The fact that the larger depletion widths still have a relatively low required drive power is due primarily to the PN junction's ability to generate large electric fields, and thereby a large electrooptic effect, in the center of optical waveguides.

The third and perhaps most objectionable disadvantage, however, and one that has not been mentioned earlier, is associated with

the small optical waveguide dimensions needed to obtain an appreciable free-carrier effect. In addition to making fabrication more difficult, these small dimensions greatly complicate the task of efficiently coupling the light to and from optical fibers. One would have to design waveguides tapered in both vertical and horizontal dimensions in order to convert from the 7-9 μm diameter dimension of a single-mode fiber to the 1 by 0.5 μm dimensions of the waveguide. However, if the modulator were used in an integrated-optic circuit, where all guides had these dimensions, and where coupling to an external fiber was not required, this would not be a problem.

Several of these shortcomings would be solved or at least mitigated by the use of the multiply-stacked BRAQWETS discussed earlier. Another approach, based upon a more mature technology, is to use free-carrier injection. Although carrier injection has serious limitations that greatly limit its usefulness at higher frequencies, it does not require a small mode size for efficient operation, and may thus be preferable for certain applications. This approach and its low-frequency implementation will be discussed in some detail in the next section.

SECTION 4

FREE CARRIER INJECTION MODULATORS

4.1 Injection in PIN Junctions

The depletion edge devices discussed in section 3 used free-carriers that were supplied by immobile donor or acceptor atoms introduced into the semiconductor at the time of growth. Another technique for producing free-carriers is to generate them by simultaneously injecting electrons and holes into an intrinsic (undoped) material. These electrons and holes will eventually recombine, so that their numbers must be continuously replenished, but while they exist they produce the same reduction in the index of refraction produced by donors and acceptors.

The most efficient way of injecting free electrons and holes into a semiconductor is to flow a current through the double heterostructure sandwich shown in fig. 4.1. This is very similar to the PN junction used in section 3, but has an additional layer of Intrinsic material sandwiched between the P and N-type materials (Hence, PIN). The top and bottom doped layers have a relatively large aluminum concentration which, in addition to lowering the index of refraction, also increases the band gap in these layers. The net effect is to provide a dual rectifying action, so that both electrons and holes can flow over the interfacial barriers when the junction is forward biased, but neither can escape once in the intrinsic material. The high degree of carrier confinement, which occurs for an aluminum concentration of 25% or greater, is responsible for the high efficiency of AlGaAs/GaAs lasers using these double heterostructures.

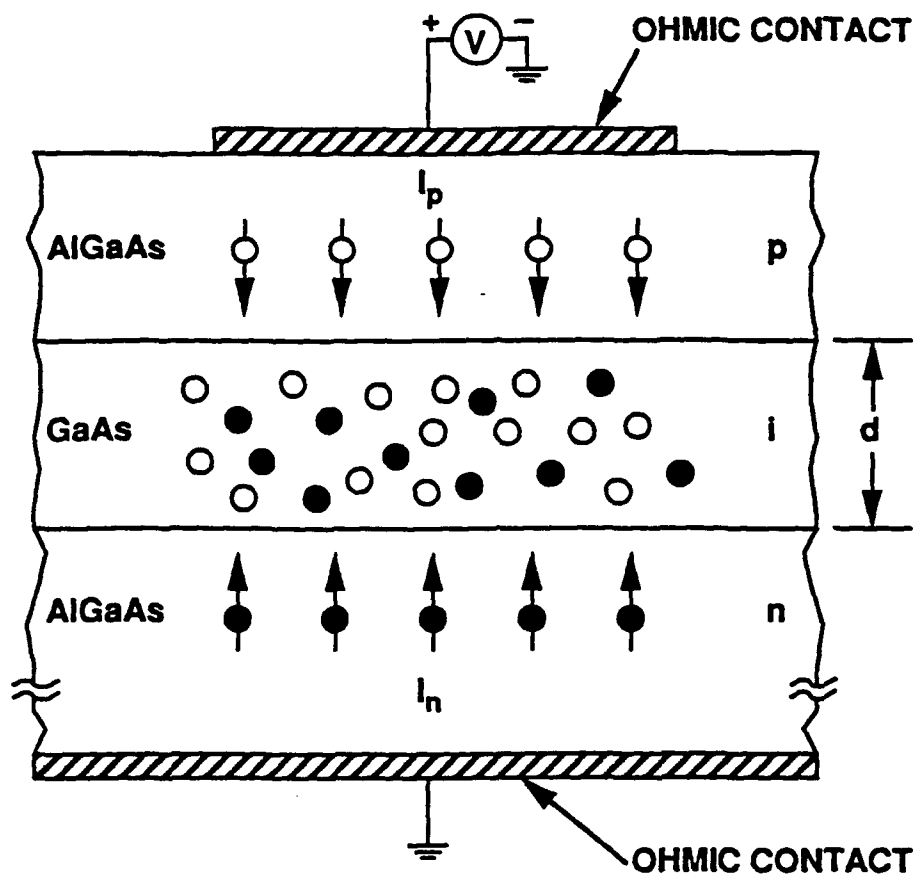


Fig.4.1 Cross-sectional view of a PIN diode with current flowing. The injected holes (circles) and electrons (dots) intermingle in the GaAs guiding layer, forming a plasma that changes the dielectric constant of this layer.

There are several things that can happen to the holes once they are in the intrinsic region. They can recombine at the junction interfaces, they can recombine in the bulk of the material, and they can diffuse laterally. The first process, interfacial recombination, is minimal in AlGaAs structures, and will be henceforth ignored. Bulk recombination, on the other hand, is significant, and occurs at a rate that is determined by the product of the concentrations and by the low (but finite) dopant background, giving recombination times that can vary anywhere from about 1 microsecond for very low concentrations ($10^{15}/\text{cc}$) and pure material, to about 1 nanosecond for very high concentrations ($\sim 10^{19}/\text{cc}$). Lateral diffusion moves the carriers to the left and right, reducing the densities and thus slowing the recombination rate somewhat. The primary effect of diffusion, however, is to spread the charge out over a larger volume.

The recombination time determines the upper frequency limit on a device using a PIN construction. The time it takes for the free-carriers to reach an equilibrium value once the current is turned on and the time it takes for these carriers to disappear once the current is turned off is determined by the recombination time. This time can be reduced somewhat by artificially increasing the number of recombination centers, but this has the undesirable side-effect of increasing the rate at which carriers must be replaced, which means a higher operating current. A better way to increase the upper frequency limit is to use a modification of this configuration, such as a transistor geometry⁽¹¹⁾, in which the charges are pulled out of the intrinsic region and allowed to recombine elsewhere. The PIN geometry is, however, adequate for modulators or switches requiring bandwidths of less than 100 MHz. In light of this fact, and because the physical phenomena associated with the PIN diode are basic to the transistor modulator, we decided to investigate the PIN geometry in some detail before moving on to other configurations.

4.2 Time and Spatial Response of a PIN Device

To better appreciate some of the issues involved, we shall consider first the case of current injection from the finite-width electrode of fig. 4.1. If the thickness of the intrinsic region in the underlying PIN stack is much less than the diffusion length, as it will be for all the examples in this and following sections, the carrier density will be approximately uniform in the vertical direction. The lack of any means for confinement in the lateral direction, however, means that the carriers will diffuse horizontally out of the injection region, so that the concentration will decrease as one moves away from the center. The one-dimensional equation that governs the time-dependent concentration is

$$D_a \partial^2 n / \partial x^2 - Bn^2 - n/\tau_0 + J(x)/(ed) = \partial n / \partial t$$

where n is the free-carrier concentration ($n=p$), B is the bimolecular recombination rate, τ_0 the background recombination rate, $J(x)$ the injected current, e the electron charge, d the intrinsic layer thickness, and D_a the ambipolar diffusion constant given by

$$D_a = 2D_n D_p / (D_n + D_p)$$

Because of the finite conductivity of the top p-layer, current will flow not only downwards, but also horizontally, so that $J(x)$ will not be a constant. This horizontal flow is best understood in terms of the equivalent circuit of fig. 4.2. One sees that if the resistances connecting the individual diodes is small, all of the diodes will have approximately the same current. For larger sheet resistivities, however, the center diode will carry the lion's share, with the current through the other diodes falling off rapidly with x . By assuming that the underlying n-layer has

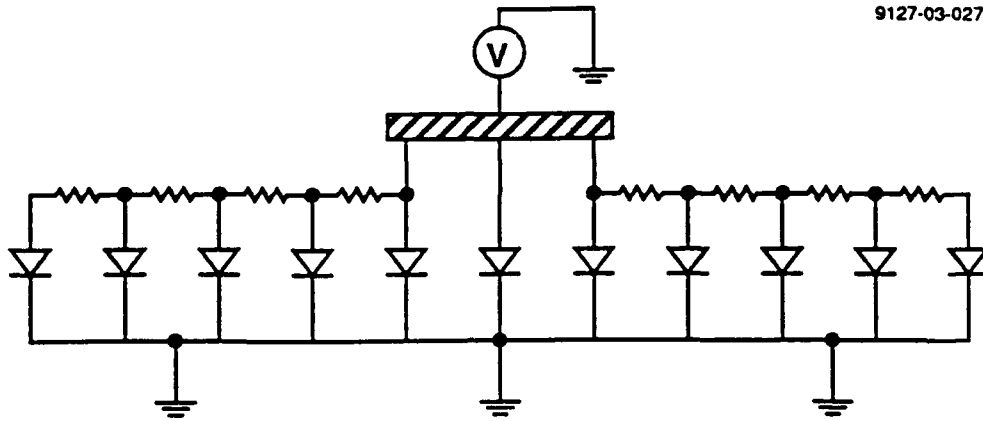


Fig. 4.2 Equivalent circuit of a PIN diode with a narrow top electrode that shows how the finite resistivity of the top layer reduces the current flow through outlying sections of the structure.

infinite conductivity, and by using an incremental model for the voltage across a PN diode, Yonezu et al.⁽¹²⁾ have developed an approximate expression for $J(x)$ which is as follows:

$$J(x) = J_0 \quad x < w/2$$

$$J(x) = J_0 / (1 + ((|x| - S/2) / l_0)^2) \quad x > w/2$$

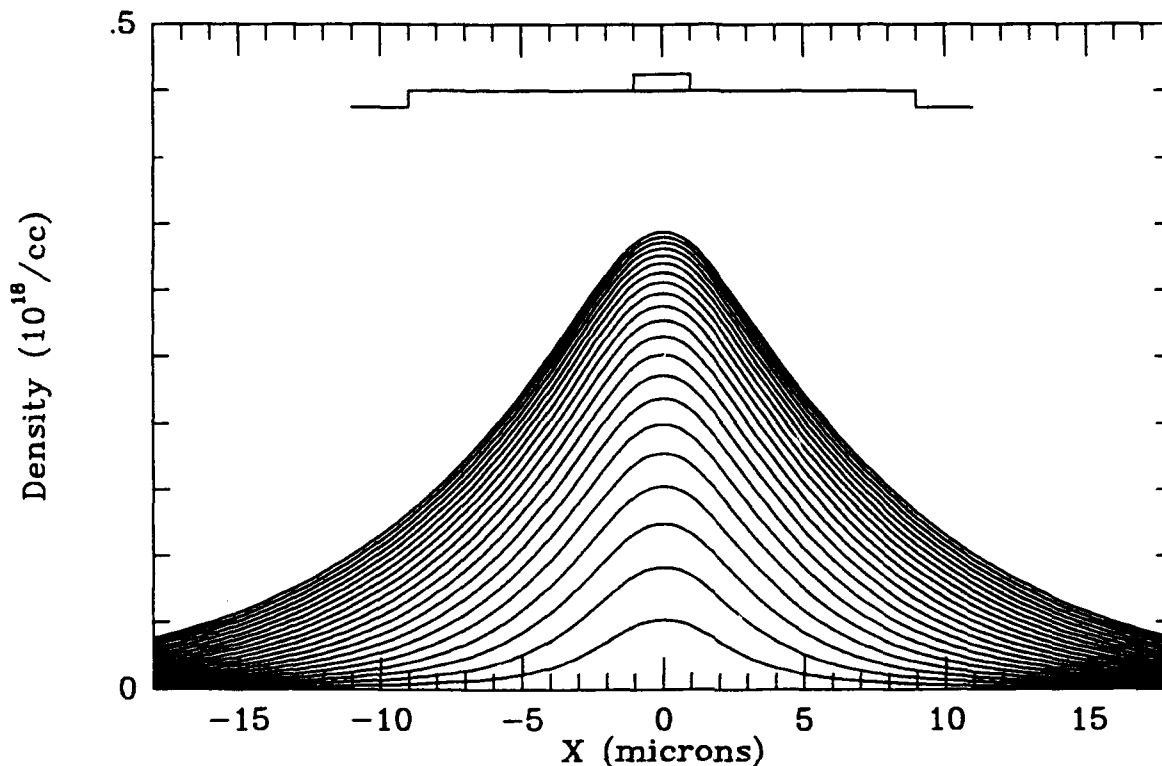
$$l_0 = S / ((1 + S\beta R I_t / 2)^{1/2} - 1)$$

$$J_0 = 2 / (\beta R l_0^2)$$

where S is the electrode width and R the sheet resistivity of the top p-layer. β is the junction diode parameter given by $\beta = e/\eta kT$; η is approximately 2 for AlGaAs/GaAs junctions.

The differential equation for $n(t)$ is nonlinear in n , and therefore does not yield a simple closed-form analytic solution. It is a trivial matter, however, to solve this one-dimensional equation numerically. The program which results consists of about 20 lines of code, runs on a PC, and gives complete solutions in a matter of minutes. A typical solution is shown in fig. 4.3 for a pn sandwich having the dimensions and resistivities of an actual device that we have fabricated (and which will be discussed in more detail in a subsequent section). The plots are "snapshots" of the carrier density, plotted at one nanosecond intervals following the abrupt turn-on of the electrode current. The density at the center increases rapidly at first, then slows down, ultimately reaching an equilibrium value.

One sees that the effect of lateral diffusion is significant, yielding profile widths that are several times that of the 2 micron wide injecting electrode. This spreading can be undesirable if the net effect is to change the index of



$Be = 2.0(e10)$ $I = 100$ ma $d1 = 0.8$ $d2 = 0.7$
 $G = 2.0$ $W = 8.0$, $L = 2.0$ mm $TAUzero = 1000$ ns
 $Da = 18.77$ Sheet Resistivity = $1800/d1$ $Beta = 18.4$
 1.0 nanosecond between traces
 RHDIFF Program

Fig. 4.3 Free-carrier density as a function of distance away from the center of the injecting electrode. The time interval between traces is one nanosecond. The density is seen to increase rapidly from zero, slowing down as it approaches its steady-state value.

refraction uniformly throughout the particular device employing this effect. This will not always be the case, however, and, as we shall see in the next section, devices can be fabricated that perform reasonably well in spite of this spreading.

SECTION 5

EXPERIMENTAL RESULTS

5.1 Optical Waveguide Fabrication and Measurement

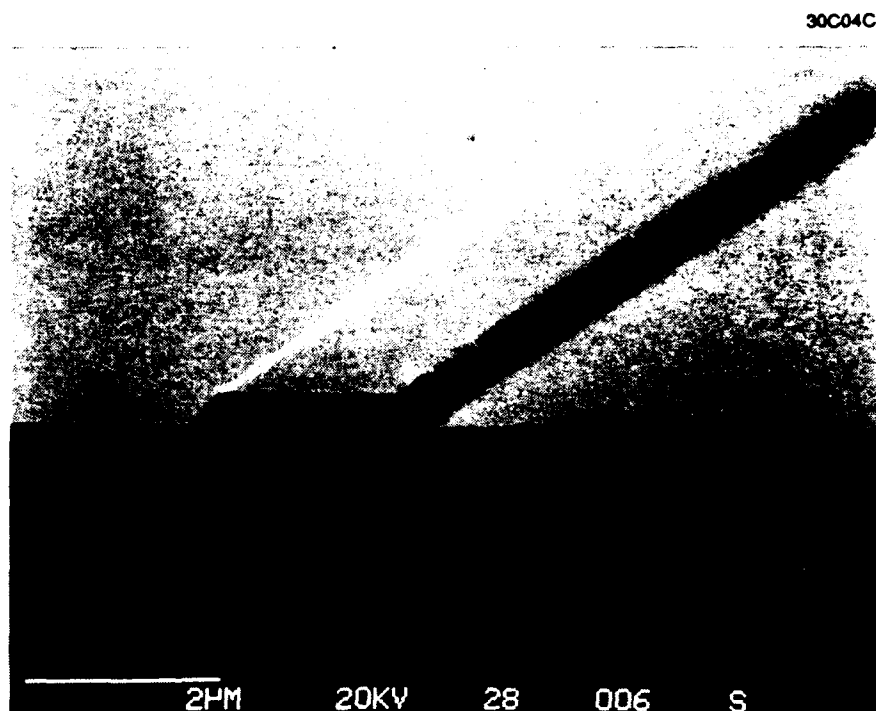
Semiconductor optical waveguides consist of ridges or ribs that have been "machined" from the surface of a multilayer epitaxial material. This micro-machining can be done with wet chemical etching, reactive ion etching, or ion-milling techniques, all of which have certain advantages and disadvantages. Wet etching gives the smoothest sidewalls, but cannot be used for closely located guides because of severe undercutting due to the omnidirectionality of the etch. Reactive ion etching, on the other hand, which has no undercutting, and ion beam milling, which has some, can give 1 micron spacings between guides, but also produce greater sidewall roughness. All of the devices produced on this contract used wet chemical etching; however, we have more recently developed RIE and ion beam milling techniques for devices requiring closer guide spacings.

If the ridge heights are not too large, one can use the effective index technique⁽¹³⁾ to determine the profiles and propagation constants for all of the modes that will be supported by such a guide. We have developed a program that uses a modified version of this technique to calculate both the real and imaginary part of the propagation constant for waveguide structures, and which includes the effects of lossy materials, metal electrodes and finite bottom cladding layers (i.e., radiation into the substrate). We have used this program to determine the guiding layer thickness, index differentials, and lateral rib dimensions necessary for single-mode operation, the cladding thicknesses needed to adequately reduce the attenuation due to the electrodes

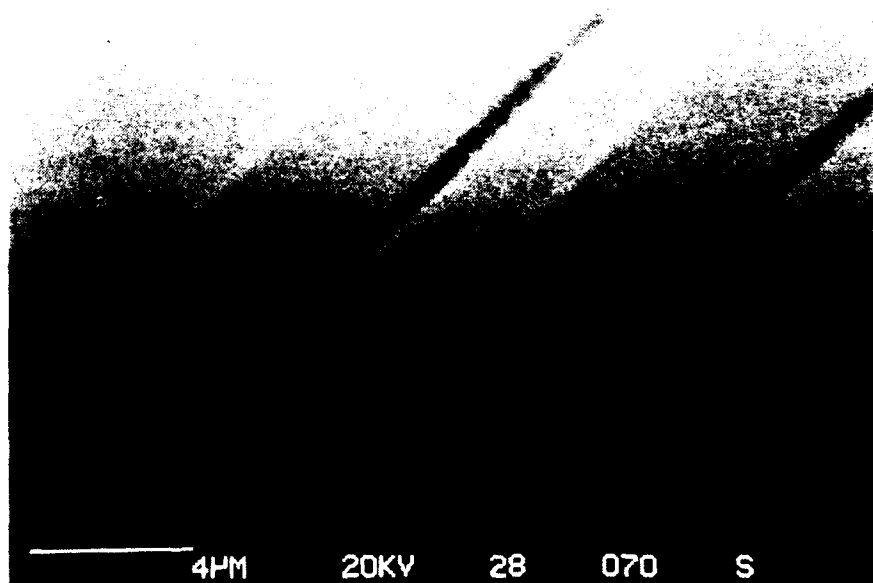
and substrate radiation, and the transfer lengths for directional couplers.

Once the appropriate parameters were determined, wafers were grown with the necessary epitaxial layers, and then etched to form the ribs. A photomask set was developed for this program that allowed us to fabricate single and paired waveguides with or without top electrodes. The ribs were formed by first coating the top of the wafer with photosensitive resist, transferring an image of the photomask to this resist by contact exposure, and then developing the resist. The resulting pattern consisted of 4 to 10 μm -wide by one inch long lines of resist that protected the underlying material, so that the subsequent etching step would remove only the material adjacent to the guide. After etching, the wafer was cleaved perpendicular to the guide at two points along the guide, thereby forming the input and output facets. Figure 5.1 shows the cleaved surface and top structure of single and paired waveguides that were formed by wet chemical etching.

Light can be coupled into the waveguide by butt-coupling to an optical fiber. The emerging light is then either coupled to another fiber, or collected with a lens. For waveguide characterization, this second technique is preferred, as it allows one to measure the mode profile as well as the transmitted power. Fig. 5.2 shows the experimental setup that was used for all of the waveguide measurements performed on this program. The waveguide light at the output facet was imaged onto either a detector, or onto a frosted glass screen so that it could be recorded with an IR vidicon. Monitoring with a vidicon allows one to measure the spatial profile of the waveguide mode, thereby determining the mode order as well as the degree of confinement (examples of this profiling will be shown in later sections), whereas imaging onto a detector allows one to measure the absolute power, as well as the high-frequency response of



(a)



(b)

Fig. 5.1 Optical waveguides formed by wet-chemical etching.

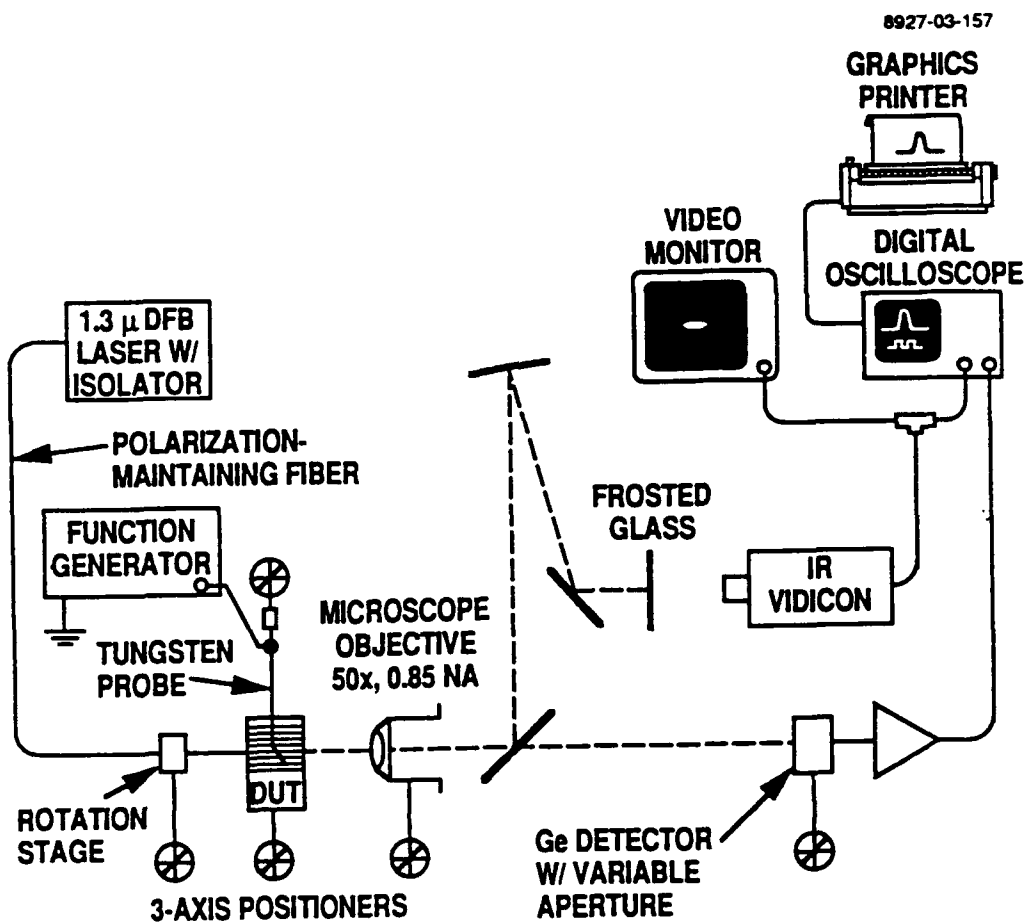


Fig. 5.2 Experimental setup used to measure waveguide loss, mode profiles, and modulator response times.

modulators and switches. An overhead microscope and a variety of 3-axis positioners and probes were available to facilitate alignment of and connection to the devices under test.

Waveguide loss was determined using the Fabry-Perot technique described in reference 14. A GaAs waveguide of finite length having cleaved facets at each end forms an etalon with 30% reflectivity. If one heats the waveguide a few degrees, the index change that results will change the optical length of the etalon, so that the output power will go through a series of maxima and minima. By knowing the length of the guide and the facet reflection coefficient, one can determine the internal guide loss. Fig. 5.3 shows the power throughput of a guide that had been heated by flowing current through a small resistor positioned just above the guide. The loss for this particular guide was 0.39 dB/cm.

A variety of waveguiding structures were grown and measured. The first batch used the single heterostructure of fig. 5.4, and were essentially undoped, i.e., the dopant concentrations were the background levels associated with the MOVPE growth process, typically $10^{15}/\text{cc}$. Guides which were fabricated from this material had losses varying from a low of 0.39 dB/cm to a high of several dB/cm, with typical values of about 1 dB/cm. These loss values were adequately low for the applications contemplated. Other devices required a more complex heterostructure with higher dopant levels, and consequently had higher guide losses. Each of these structures shall be discussed in the sections that follow.

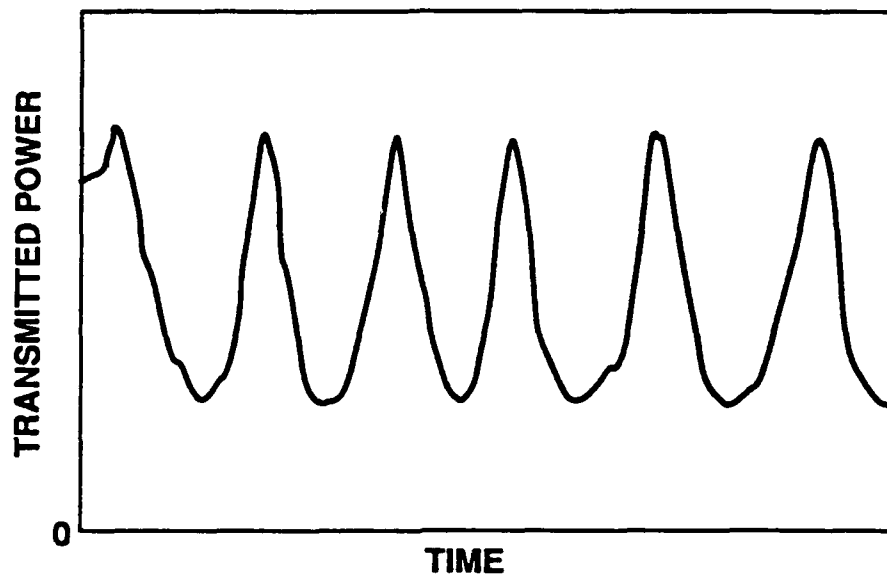


Fig. 5.3 Variation of waveguide transmission as the guide is slowly heated. The periodic response is that of a Fabry-Perot etalon with varying plate separation.

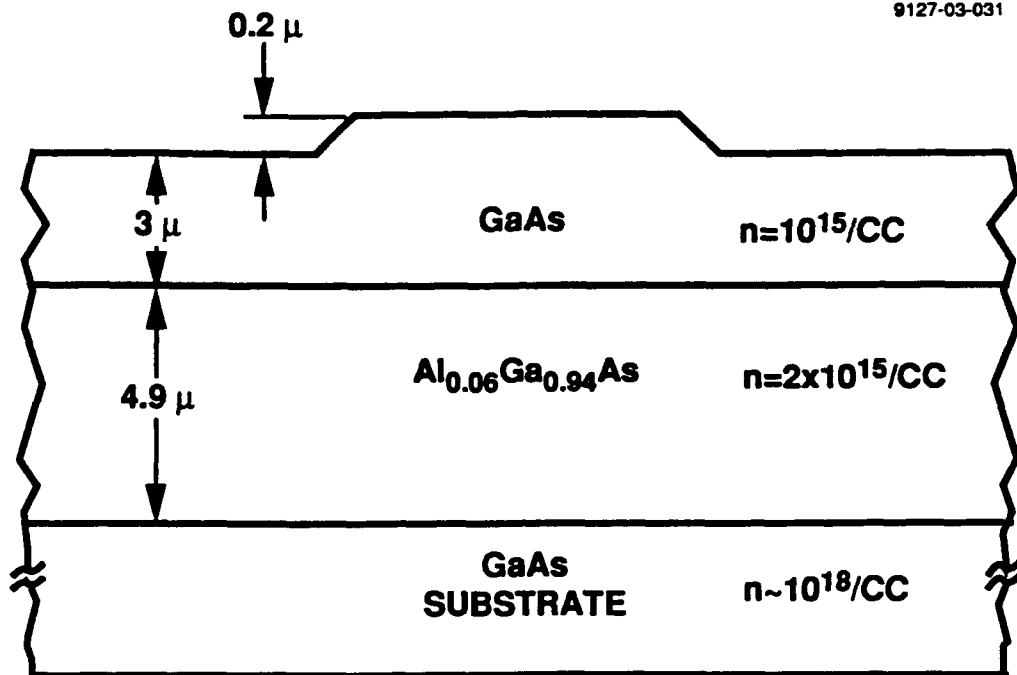


Fig. 5.4 Optical waveguides using a single heterostructure.
The top cladding layer in this case is air.

5.2 Free-Carrier Injection Modulators

The heterostructure sandwich needed for a free-carrier injection modulator must satisfy several criteria. First, the bandgap of the cladding layers must be large enough to confine the injected carriers in the vertical direction. Casey and Panish⁽¹⁵⁾ have shown that x-values of 0.25 will give excellent confinement at GaAs/Al_xGa_{1-x}As interfaces, and we have chosen this value for our first (and, as it turned out, our only) free-carrier injection modulator. The thicknesses of the cladding layers having this particular value of x must be then be adjusted so that the losses due to the top electrode and due to light leakage into the substrate are held to acceptable values. The guiding layer thickness was chosen to give the largest value consistent with single-mode operation in the vertical direction. A value for the p and n dopant levels was selected to minimize losses during injection while holding the optical losses to an acceptable value (~ 2 dB/cm). The dopant concentration in the intrinsic guiding layer was made as low as possible, in this case approximately 10¹⁵/cc for the growth process used.

The structure that evolved from this design procedure is shown in fig. 5.5, which shows the design goals and the actual values (in parenthesis) measured for one of the completed wafers. Four wafers were grown to this specification using our in-house MOVPE reactor. The layer thicknesses were determined by etching a cleaved facet of the wafer with a solution that dissolves AlGaAs faster than GaAs, thereby producing a physical demarcation at the material interfaces; the distance between these interfaces was then measured with a scanning electron microscope. Aluminum concentrations were measured using a double-crystal X-ray technique that measures the lattice parameters for the two types of material. The n and p carrier concentrations were measured and adjusted during calibration runs, and the concentration of

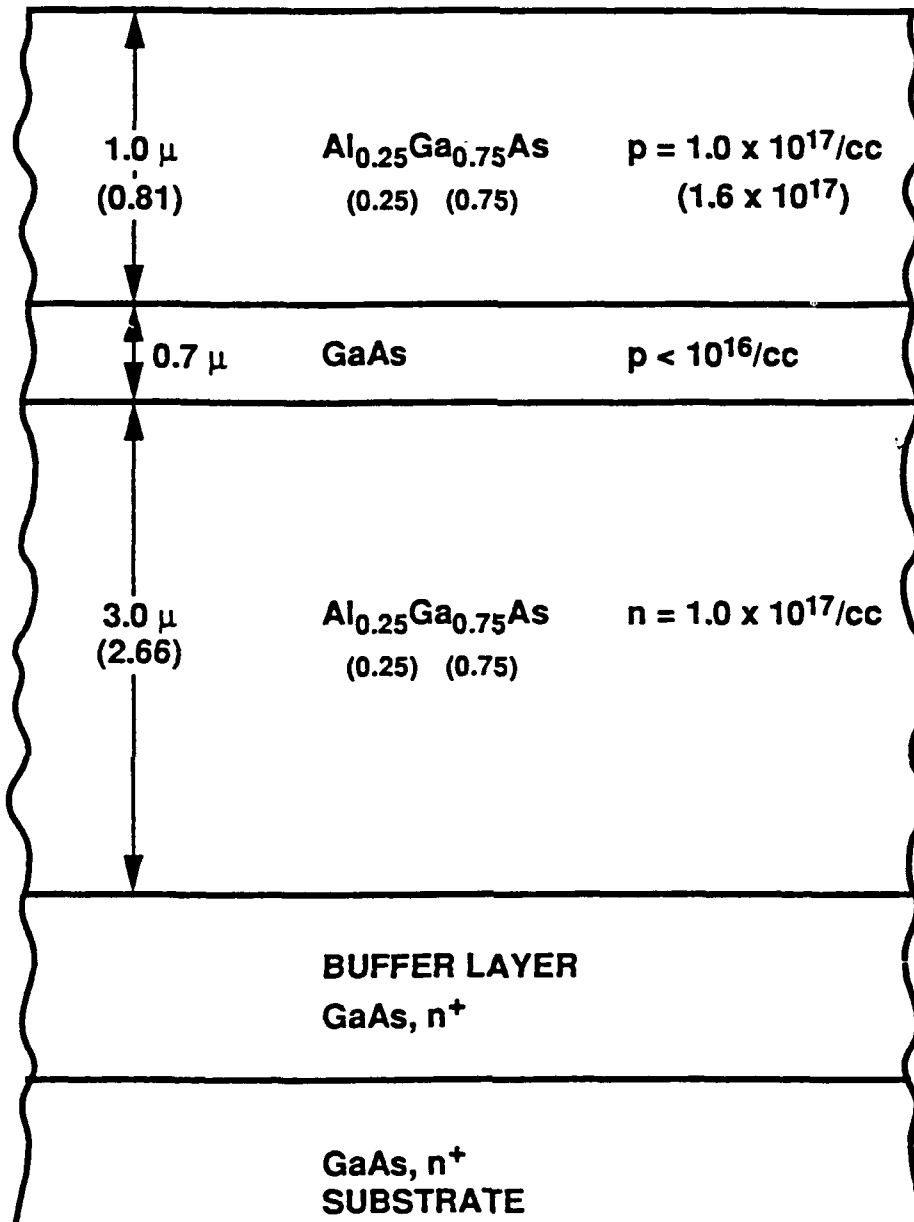


Fig. 5.5 Heterostructure design for free-carrier injection modulator. The numbers in parenthesis are the measured values for the grown wafer.

the top p-layer confirmed with Hall measurements once the wafer was completed.

Diodes were made from this material by forming ohmic contacts to the top and bottom surface. The top contact was alternately sputtered Au/Zn/Au/Zn/Au, while the bottom was e-beam evaporated AuGe/Ni/Au. After deposition, the contacts were flash annealed at 480°C for 15 seconds in a 4% H_2 forming gas. This process created ohmic contacts by diffusing in large dopant concentrations, making the depletion region so small that high tunneling currents dominate.

The chip was then diced into several 1 mm square pieces so that diode parameters could be measured for a known area. Fig. 5.6 shows the I-V trace for one of these small diodes, and fig. 5.7 the forward-biased portion on a logarithmic scale. Although the complexity of the double heterostructure does not lend itself to a simple characterization in terms of a few physical parameters, these data do agree reasonable well with the published results for similar structures, implying that the PIN structure is functioning correctly.

More important, however, is the fact that the diode emits light at the cleaved edges when forward biased. This means that free-carriers are being formed at the junction, and that a significant number of these free-carriers are undergoing radiative recombination. Fig. 5.8 shows the radiated power (as measured by a silicon PIN detector) versus the forward bias current. The high-slope region is associated with defect-dominated recombination, and the low-slope with radiation-dominated recombination. The crossover point occurs at about 5 amperes/ cm^2 , which is one to two orders of magnitude below the current densities that will be used in practical implementations. Defect recombination, therefore, can effectively be ignored.

9127-03-045

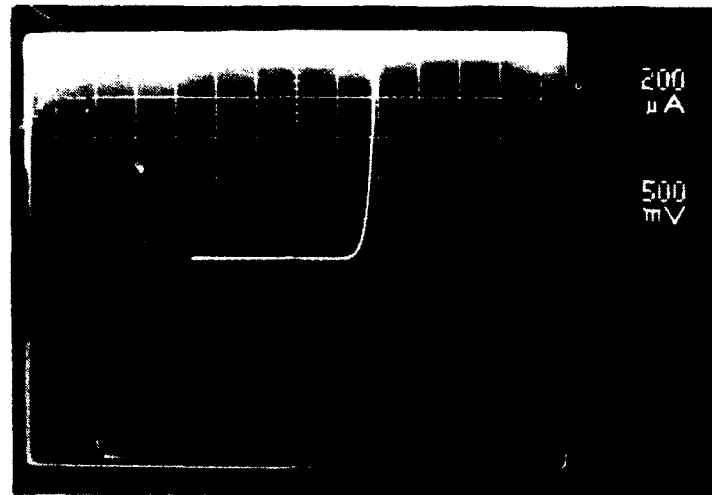


Fig. 5.6 I-V trace for diodes fabricated from the material of fig. 5.5

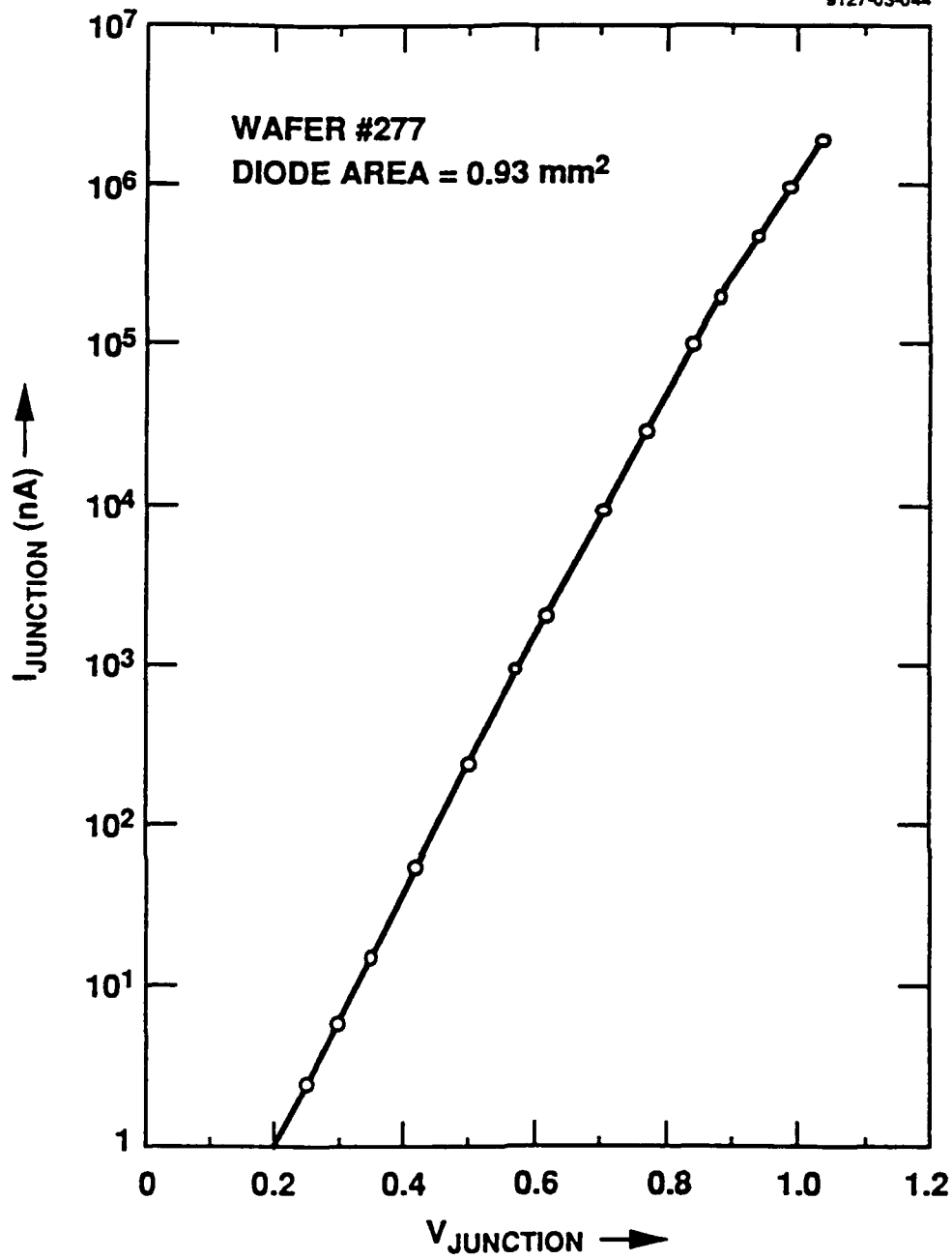


Fig. 5.7 Logarithmic plot of the forward-biased characteristics of the same diode.

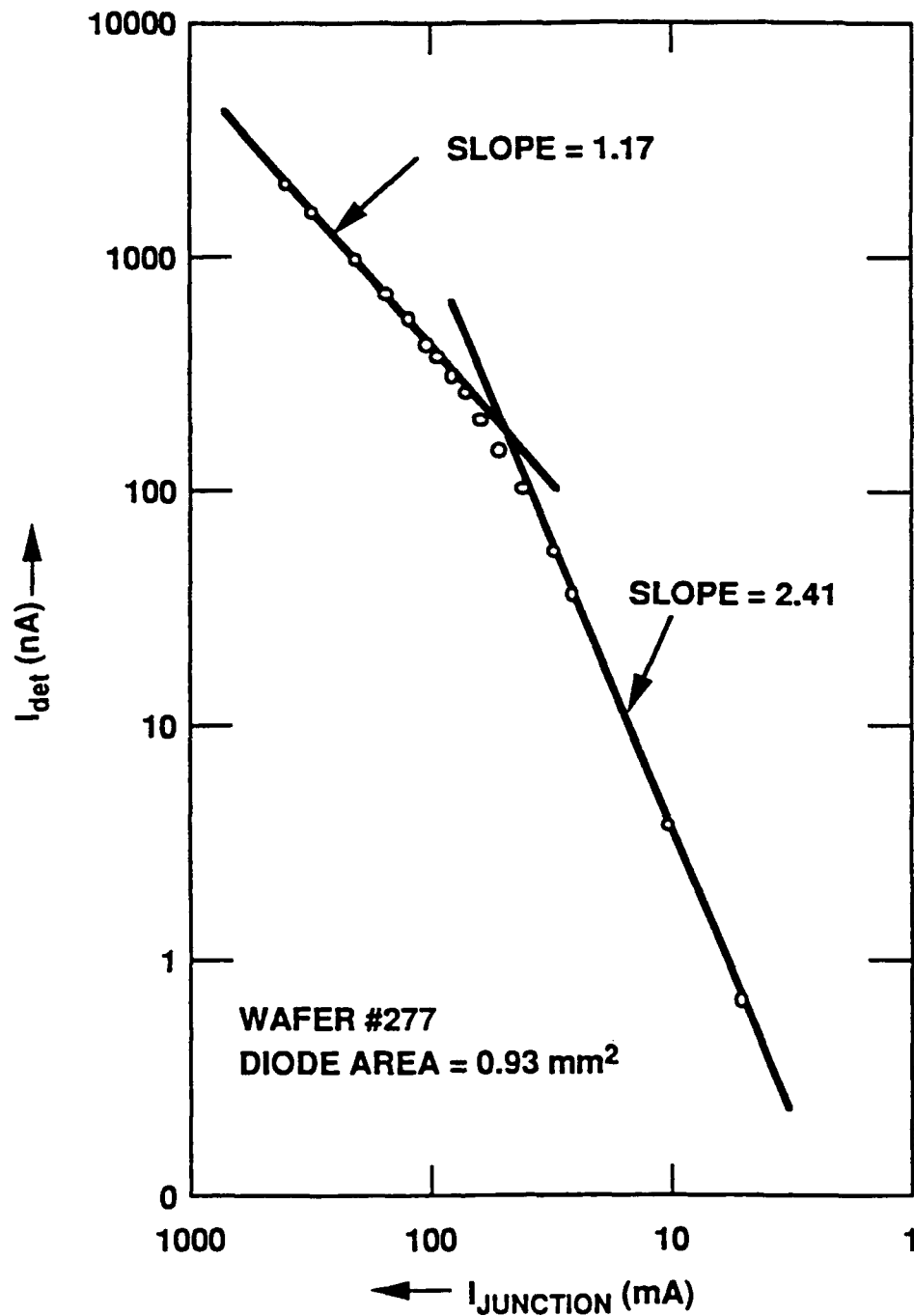


Fig. 5.8 Plot of the emitted light intensity (as measured by a Silicon photodiode) as a function of the injected current. The high-slope portion of the curve is associated with defect-dominated recombination, the low-slope with radiation-dominated recombination.

5.3 PIN Guiding-Antiguiding Modulator

After characterization of the material, a set of waveguides was fabricated using the photolithographic techniques described in section 5.1. The patterned resist was developed and then postbaked to harden the long thin lines defining the waveguides. An $\text{H}_2\text{SO}_4:\text{H}_2\text{O}_2:\text{H}_2\text{O}$ etch solution with a 2:16:320 ratio was used to etch that portion of the surface not covered by resist, thereby forming the waveguide rib. A matching photomask was subsequently used to form the resist pattern for the top ohmic electrode and probe pad for each guide. These contacts were then sputtered and annealed as described earlier. The resulting guides were 5 μm wide, had a 0.6 μm rib height and, after cleaving, were 2 mm long. An ohmic contact was also formed on the bottom surface, so that the PIN waveguide sandwich could be forward biased (fig. 5.9a).

The original intent was to use these guides to determine the level of free-carrier injection by measuring the change in transmission due to free-carrier absorption. The expected decrease in transmission for a 2 mm length and a carrier concentration of $10^{18}/\text{cc}$ was about 60%, a reasonably large effect. However, when 1.3 μm laser light was fiber-coupled into one end, and the emerging light observed with an IR vidicon, it was found that the light was completely attenuated when the PIN guide was forward biased with 40 ma of current! Using existing carrier lifetime data, it was estimated that the free-carrier concentration for this level of injection was about $6 \times 10^{17}/\text{cc}$. This level would lower the index of refraction within the guide by about 4×10^{-3} , which, as it turns out, is just a tad more than the 3.8×10^{-3} lateral index step due to the etched rib. It would thus appear that injecting free-carriers into the guide cancels the lateral confinement, allowing the light to spread laterally and creating, in effect, a "guiding/anti-guiding" modulator.

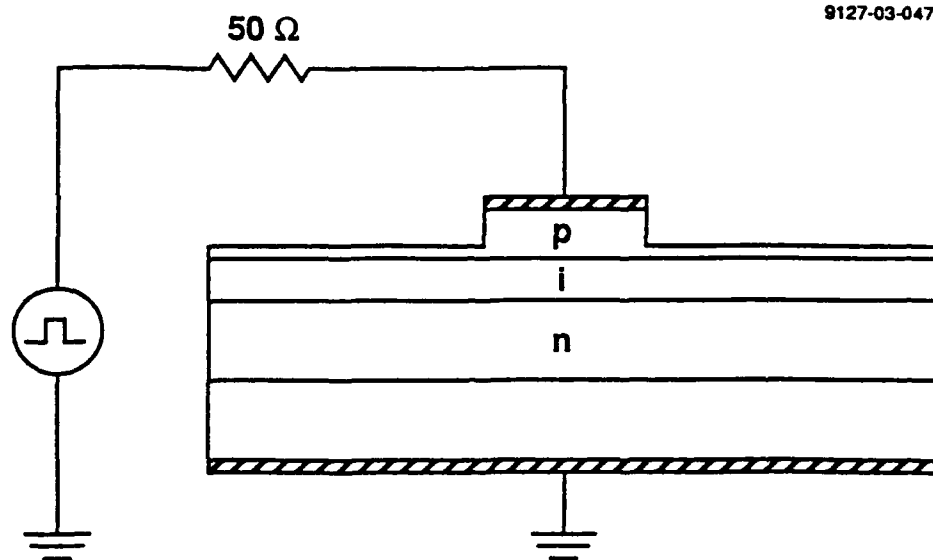
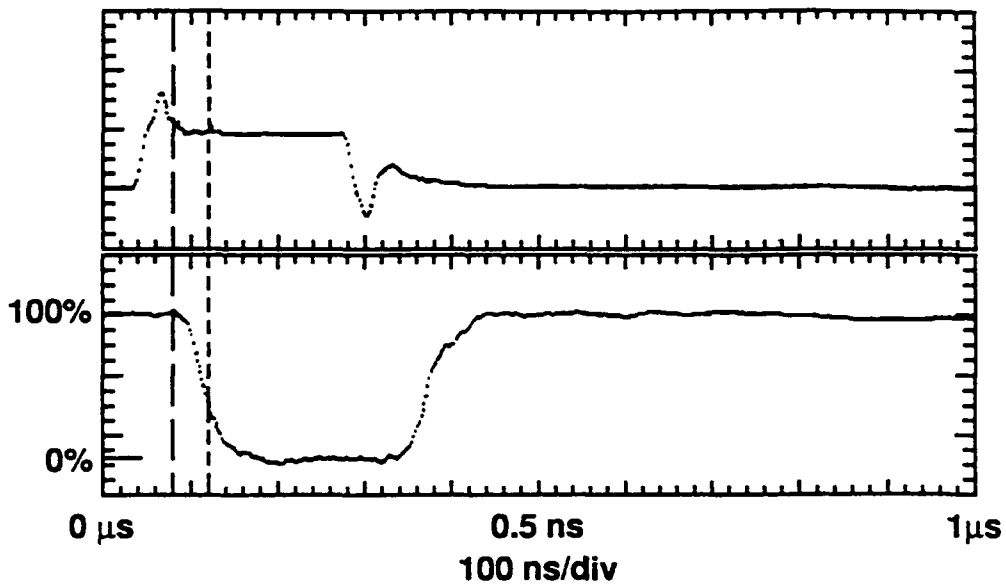


Fig. 5.9a PIN "anti-guiding" modulator

The response of this device to a step current pulse of 40 ma is shown in fig. 5.9b. The measured rise time is 38 nanoseconds, which is close to this particular detector's rise time of 32 nanoseconds. Although these times are not terribly meaningful, since they are limited more by the experimental setup than by the properties of the material, they are consistent with free-carrier effects, and are much faster than the known thermal response times (see section 5.5).

This principle of changing the guide index to affect the degree of confinement has been used by others in the past, and is thus neither new nor particularly original. It is, however, attractive because of its simplicity and the ease with which devices of this type can be fabricated. It is interesting to note that this approach has recently been proposed as one way of achieving 60 GHz modulation in semiconductors.



STOP MARKER: 120 ns
START MARKER: 82 ns
 ΔT : 38 ns

Fig. 5.9b Change in transmission (bottom trace) versus applied voltage (top trace) for the anti-guiding modulator. Note that the transmission goes to zero.

5.4 PIN Directional Coupler Modulator

Another way of utilizing the large change in the index of refraction associated with free-carrier injection is to incorporate the PIN structure into the directional coupler switch described in 3.2.1. Fig. 5.10 shows how the interguide index could be changed using the PIN sandwich discussed earlier. An ohmic contact at the top of the PIN stack allows one to inject current into the region between the guides. The free-carriers created by this current lower the index of refraction in the guiding layer between the guides, thereby decreasing the coupling and increasing the transfer length. For the proper injection current, light will switch from one waveguide to the other. A second ohmic contact on the bottom of the substrate allows current to flow out of the switch and back to the source.

The actual electrode geometry used, shown in fig. 5.11, has been changed to accommodate the high currents necessary for carrier injection. The majority of the current is carried laterally by the lower part of the electrode, which feeds current through bridging fingers into the center injecting strip as needed. This use of a wide parallel feed electrode greatly reduces resistive losses, assuring even current injection along the device. The plating bus shown in the figure is used during fabrication to up-plate the electrode to a 3 μm thickness, and plays no role in the operation of the device.

S-bends were employed to separate the guides (in this case by 40 μm) so that light could be unambiguously launched into and recovered from the four ports of the device. The photomasks were written with an e-beam spot size of 0.1 μm to minimize quantization roughness and the associated scattering losses in these S-bends. A bend radius of 2 cm was chosen to hold bending losses to a minimum without requiring excessively high rib

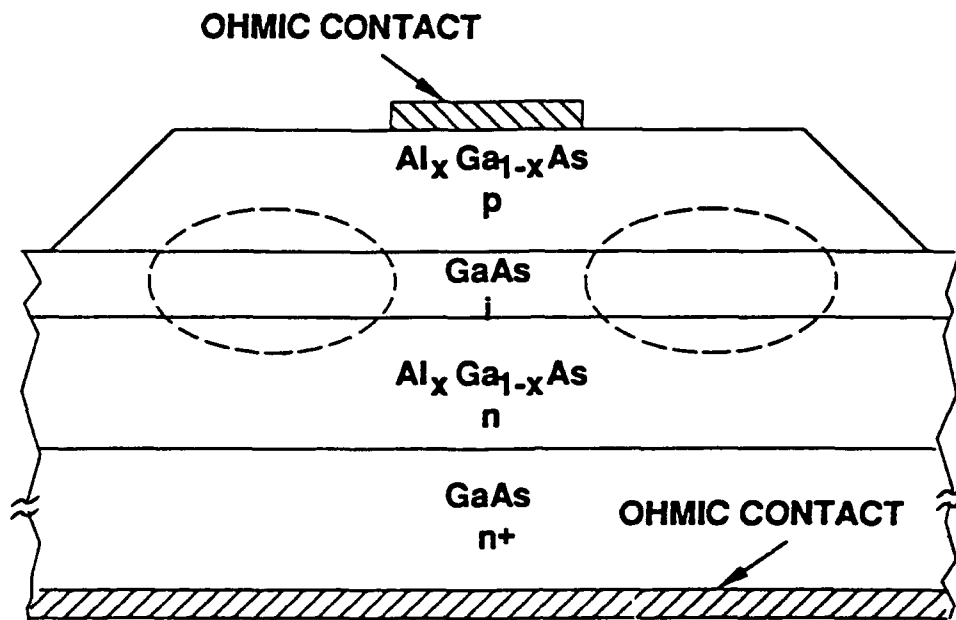


Fig. 5.10 Cross-sectional view of PIN directional coupler modulator.

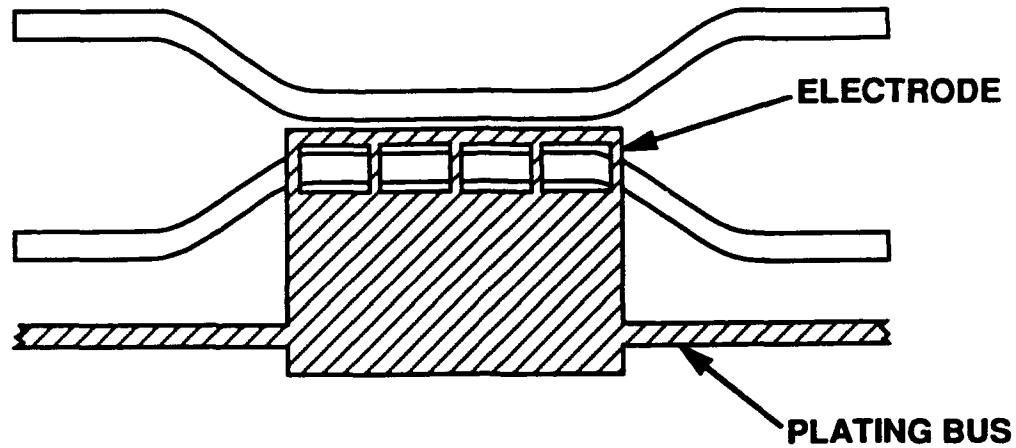


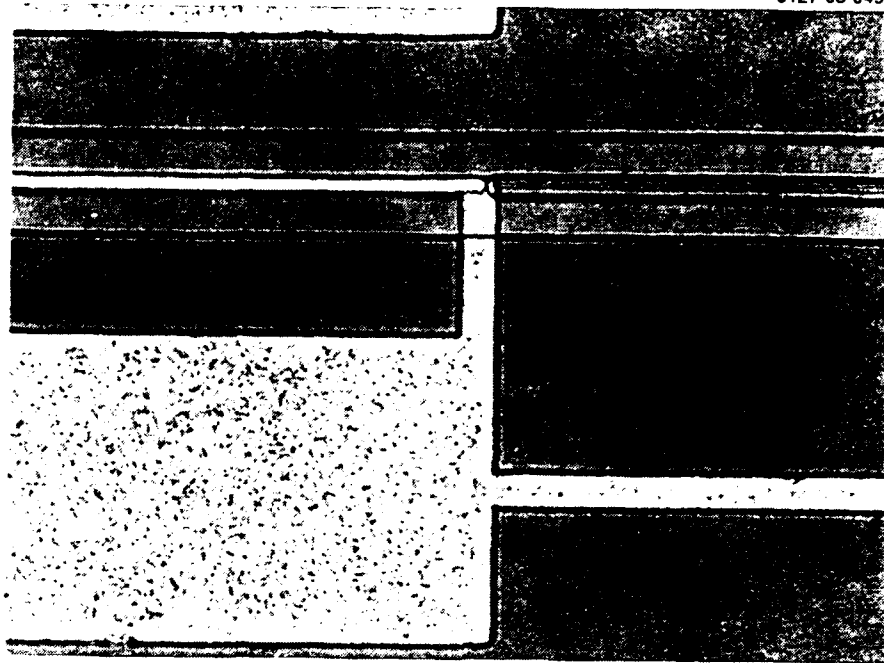
Fig. 5.11 Top view of PIN directional coupler modulator showing the wide electrode used to insure uniform current injection.

heights. Calculations using WKB⁽¹⁶⁾ and Bessel-expansion⁽¹⁷⁾ techniques showed that lateral index differentials of 0.003 should reduce the total bend loss to less than 0.1 db.

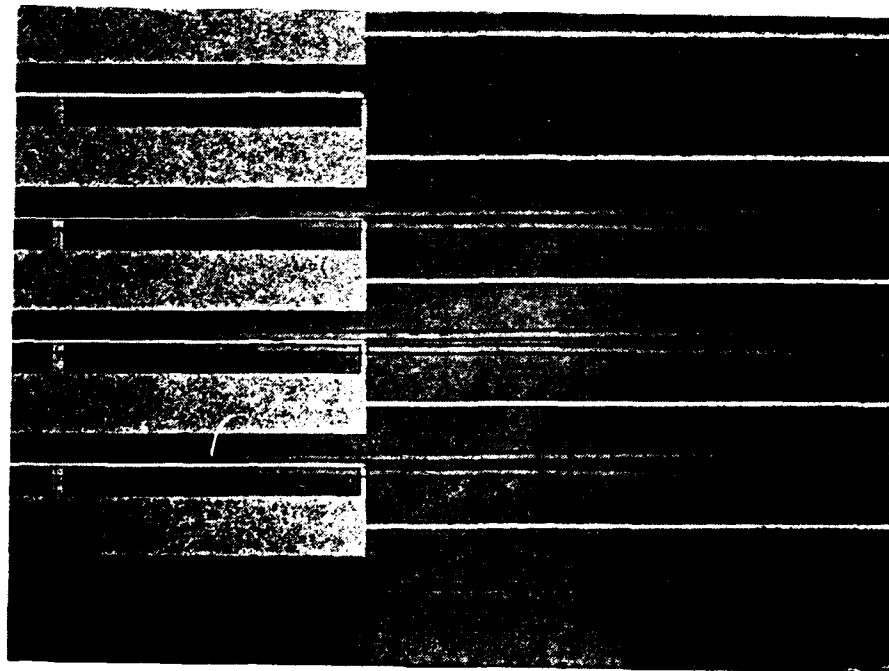
The actual devices were fabricated from the material shown in fig 5.5 using a four-layer process. The waveguides of the directional coupler and S-bends were first formed by photolithography and wet chemical etching. A 0.3 μm layer of SiO_2 was then sputtered onto the top of the wafer, and a second resist pattern formed that defined slots along the center of the couplers. These slots were etched through using either a wet buffered oxide etch or, on subsequent devices, reactive ion etching, so that the top AlGaAs layer was exposed in the slots. A third resist pattern was then formed that defined the electrodes, and these electrodes were then sputtered and annealed as in the case of the guiding-antiguinding modulator. The same resist pattern was then re-applied, and the electrode thickness increased by electroplating. The devices, before this final up-plating, are shown in fig. 5.12.

Electrical contact to the device is made with tungsten probes that touch the large portion of the electrode. The SiO_2 layer electrically isolates the most of the electrode so that current is injected into the center of the directional coupler only. The injected current is recovered at the ohmic contact on the bottom of the substrate.

The completed set of devices was cleaved and installed in the test setup of fig. 5.2. 1.3 μm laser light was coupled into one port of one device, and the intensity at the two output ports monitored with the IR vidicon. Fig. 5.13 shows the intensity profiles for the case of no injected current, and an injected current of 133 ma. The length of this particular device was 2 mm, which was just a bit short of one transfer length (about 2.5



(a)



(b)

Fig. 5.12 Photographs of the PIN coupler modulator before up-plating. Note the tapering effect of the S-bends in the bottom photograph.

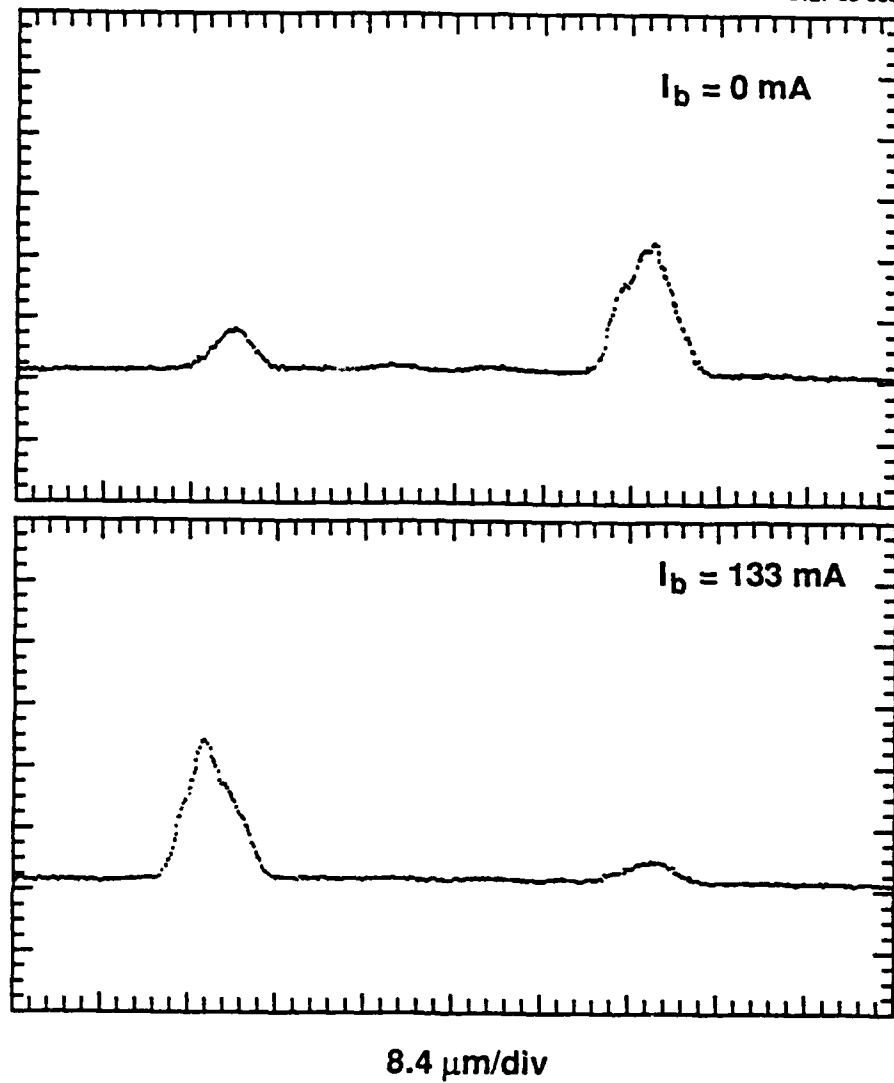


Fig. 5.13 Profiles of the optical intensity at the output facets of the PIN coupler modulator for two different injection currents.

mm for this geometry), so that there was still some light in the input arm (the left profile). Current injection between the guides greatly reduced the coupling between the guides, so that most of the power returned to the input guide (bottom figure).

Complete switching from one port to the other could not be achieved because the device tested was much less than two transfer lengths long. Unfortunately, the longer devices on this wafer that had the required two (or more) transfer lengths also had very poor transmission, for reasons not yet understood, so that they could not be tested. The fact that considerable switching took place for the short device that was tested is due primarily to the very large index change associated with free-carrier injection. By implication, longer devices should give complete switching for much lower values of injected current. Unfortunately, the program had already been stopped by this time, so that further investigations were not possible.

The transient response of this device is shown in fig. 5.14, together with the response time of the detection system. The delay between stimulus and response is just due to cable lengths, and has nothing to do with device performance. In order to compare this response with the theory of section 4.2, we have plotted (fig. 5.15) both the density profiles and the density differential between the center region and the center of the waveguiding region, as denoted by the small crosses shown in the cartoon above the density profiles. One sees that, although the total density is still increasing for times greater than 10 nanoseconds, the density differential is essentially flat. Because it is the difference in index that is responsible for switching, and not the absolute value, the response plotted here should mimic that of the device. The rise time obtained from this curve is 4-5 nanoseconds, which is in rough agreement with the 8 nanoseconds measured experimentally.

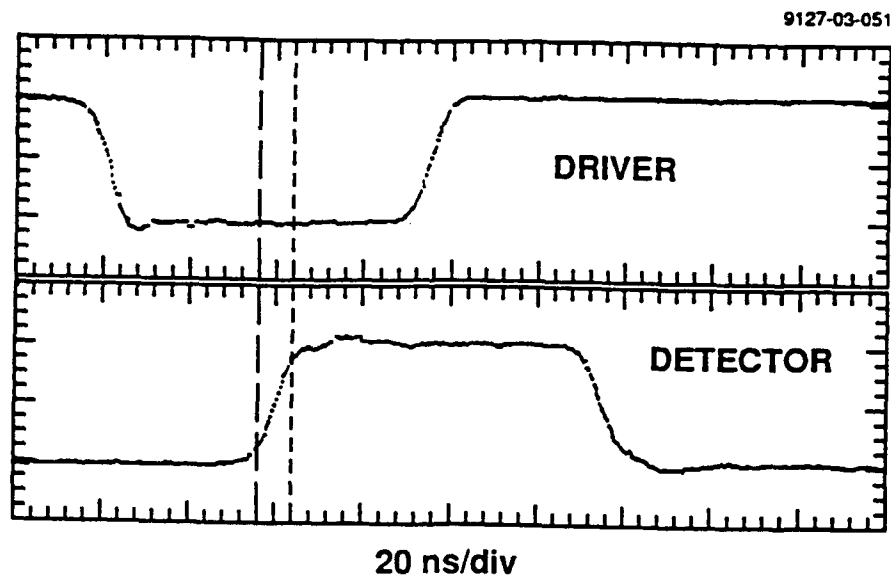


Fig. 5.14a Change in transmission (bottom trace) of the electro-optic switch of section 5.5 driven by a drive pulse (top trace). This measurement was made to determine the response time of the measurement system (driver + detector).

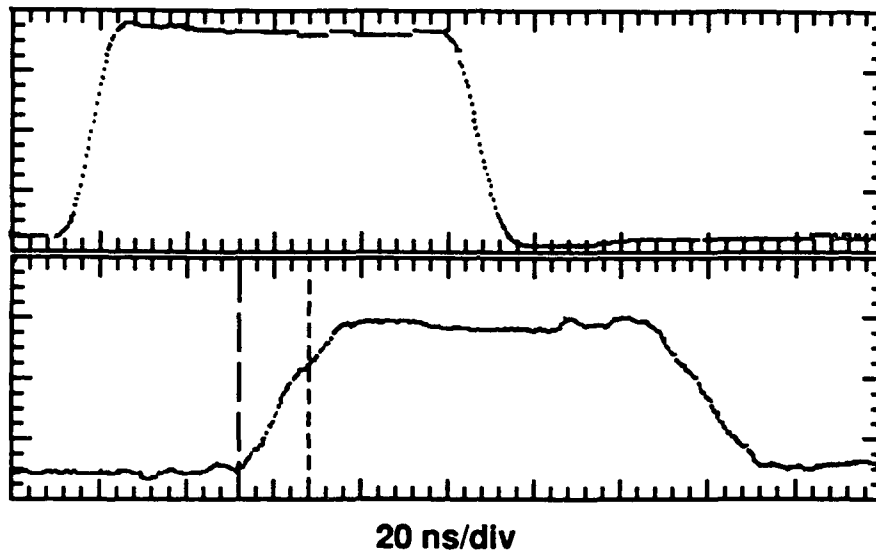


Fig. 5.14b Change in transmission of one channel of the PIN directional coupler modulator (bottom trace) for a 100 ma injected current pulse (top trace). The deconvolved rise time is approximately 8 nanoseconds.

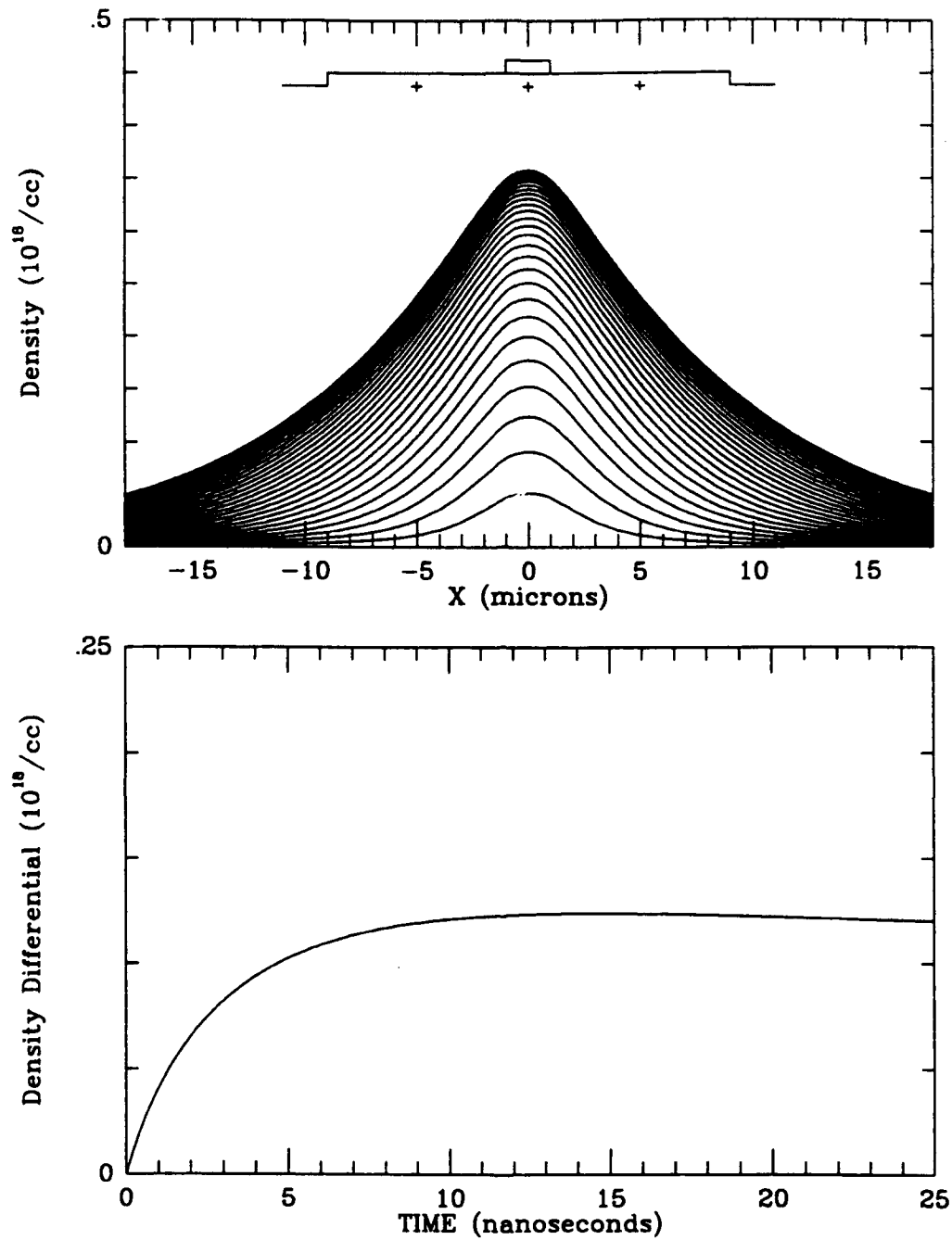


Fig. 5.15 Computed density profiles (top graph) for the PIN coupler modulator described in the text (1 ns between traces). This family of curves is used to calculate the device's dynamic response (bottom curve).

5.5 Electrooptic and Thermal Modulators

Historically, the first modulator fabricated on this program was an electrooptic switch that used the configuration of figure 5.16. The original intent was to fabricate a prototype Schottky diode modulator that could be driven either with free-carrier injection or by the well-understood and well-measured electrooptic effect, so that a quantitative comparison of free-carrier effects with known electrooptic coefficients would be possible.

A Schottky contact is a majority carrier device, i.e., the carriers responsible for conduction are primarily (but not exclusively) the dopants in the region below the metal contact. Rectification occurs because of the uni-directional hopping of majority carriers over the depletion barrier formed at the metal-semiconductor interface. If, for example, one had a gold contact on n-type material, then thermally-activated electrons in the semiconductor would, under forward bias, jump over the lowered barrier and into the metal to give a forward current.

For a sufficiently large forward bias voltage, however, the barrier that prevents holes from leaving the metal and entering the semiconductor can be lowered enough to produce a finite hole current, so that both electrons and holes can coexist in the region near the electrode, thus raising the free-carrier concentration above the dopant level. The electrons and holes would eventually recombine, just as they do in the PIN diode, but would nevertheless change the free-carrier concentration while they existed.

This effect has been studied in silicon, where it has been shown that minority carrier currents approaching 20% of the majority carrier current are possible⁽¹⁸⁾. However, the ratio of p to n

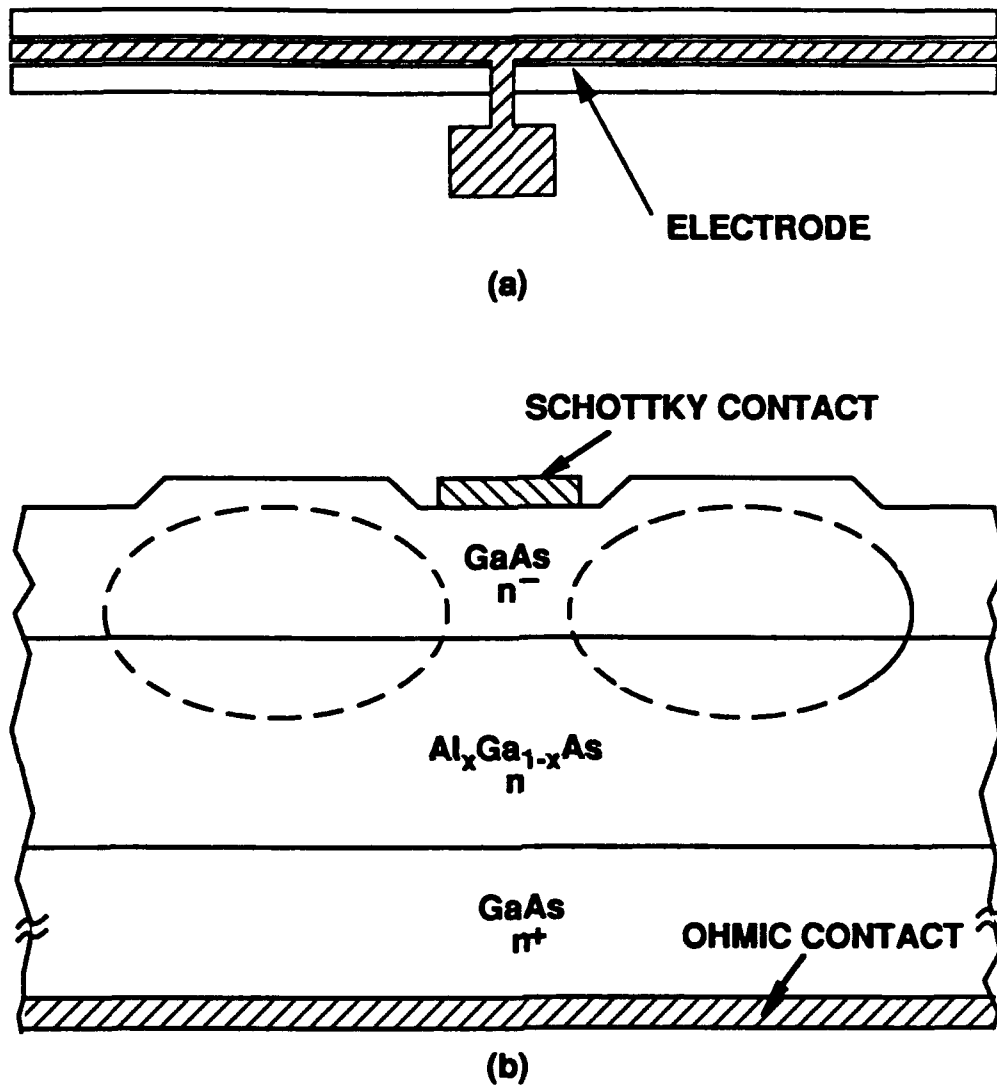
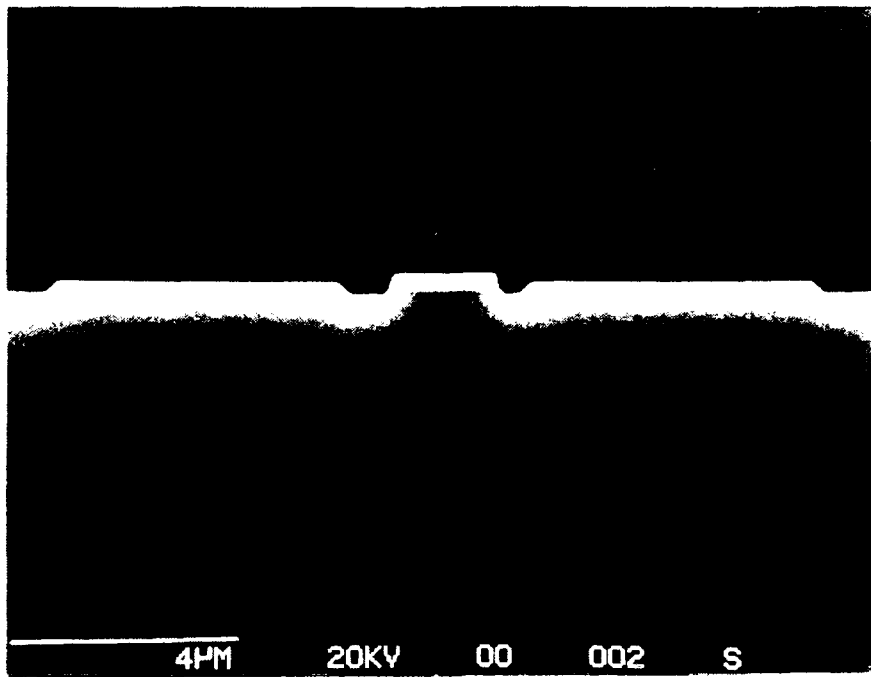


Fig. 5.16 Top (a) and cross-sectional (b) view of directional coupler switch with an interguide Schottky electrode.

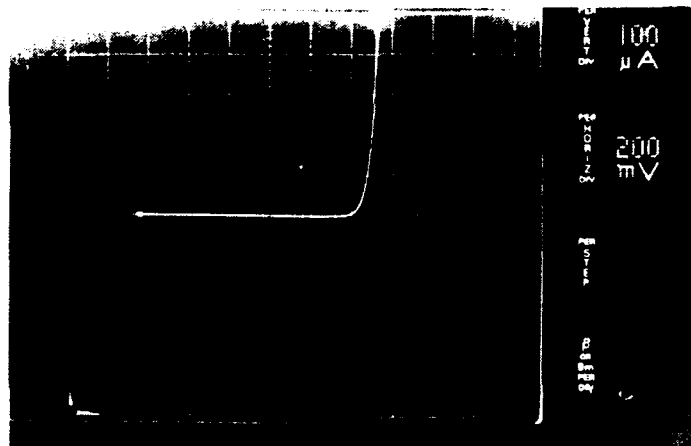
currents is a strong function of the energy gap and the barrier height of the system. Generally speaking, if the electron barrier height and band gap were the same, equal electron and hole currents would flow under forward bias. However, for most semiconductors, the band gap is larger than the barrier height, so that the electron current dominates. In GaAs, the difference between the barrier height (0.8 eV) and the band gap (1.42 eV) is so large that the hole current is vanishingly small for all values of forward bias voltage. No measurable free-carrier increase is thus possible in a forward-biased GaAs Schottky diode.

Ignorance is bliss, however, and we blissfully proceeded with the fabrication of a GaAs device, with, as we shall see, interesting consequences. The single heterostructure of fig. 5.4 was used to fabricate closely spaced waveguides. A Ti/Pt/Au metal electrode was evaporated in the interguide region to form a Schottky diode, and an ohmic contact formed on the back side of the wafer. An end view of one of the resulting devices is shown in figure 5.17, together with the measured I-V curve for the junction.

When the junction is reverse-biased, a strong electric field is formed in the depletion region beneath the electrode. This field changes the index of refraction, which in turn changes the coupling between the guides, thereby effecting switching just as it does for the device of section 5.4. Fig. 5.18 shows the measured power at the output facets of a 7 mm long device with 8 μm wide guides, a guide separation of 2 μm , and a rib height of 0.2 μm . With no applied voltage, all of the power leaves the original guide, as one would expect for a device that is two transfer lengths long. Application of an 19 volt reverse bias does switch roughly half the power into the other guide, but is not enough to cause complete switching. Application of a forward



(a)



(b)

Fig. 5.17 (a) SEM photograph of cleaved facet of the structure of fig. 5.16., showing the two guides and the interguide electrode. (b) The I-V response of the Schottky diode.

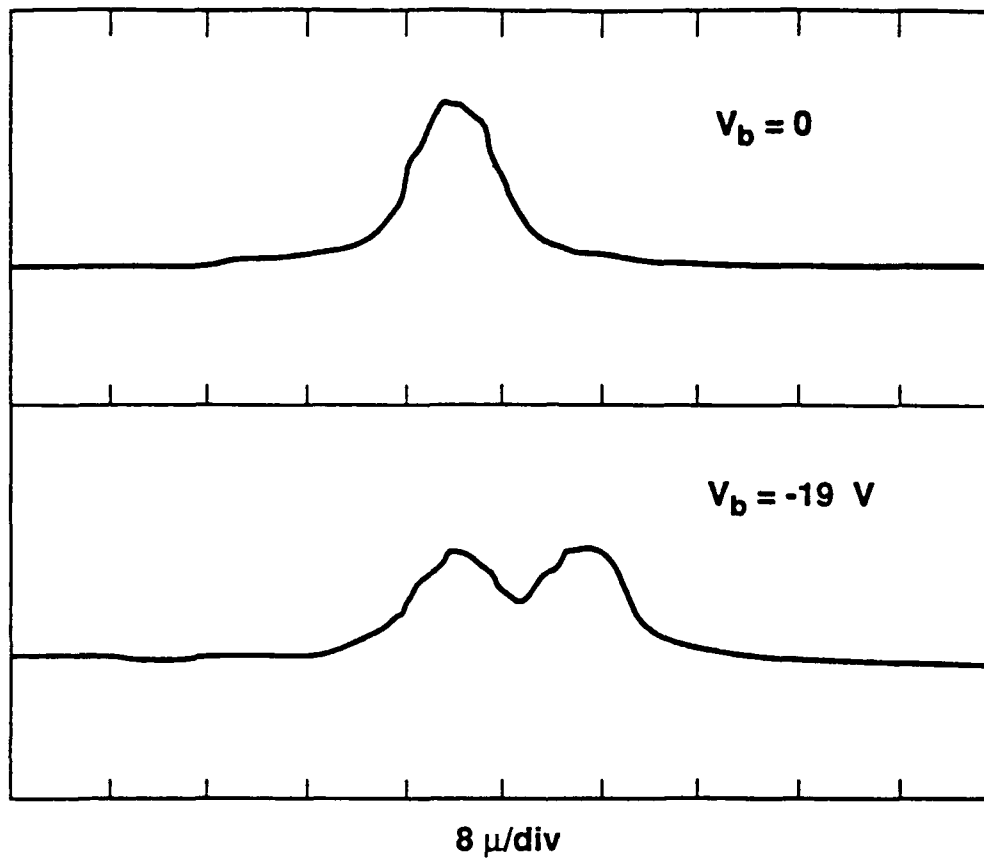


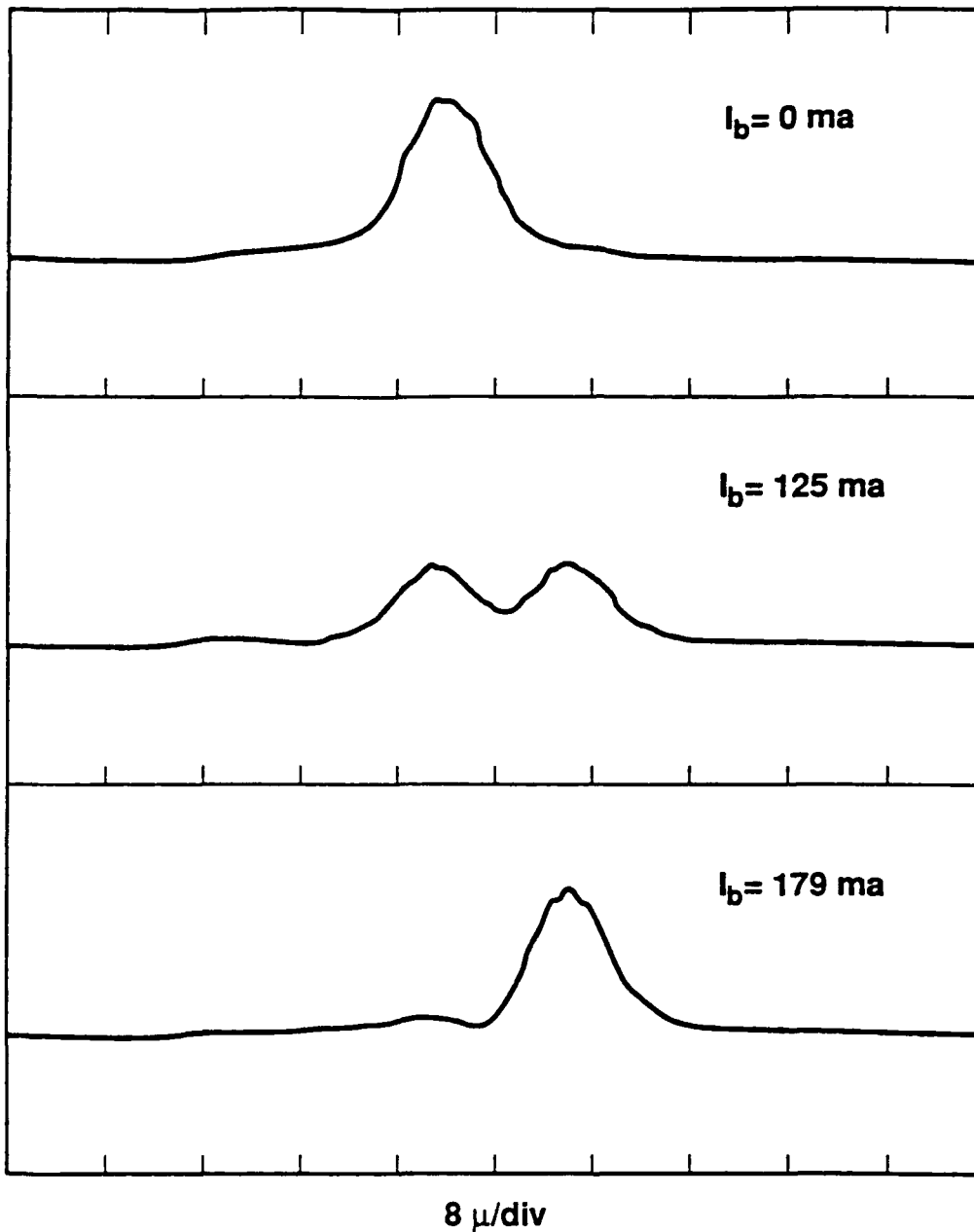
Fig. 5.18 Intensity profiles at output facet of the device of fig. 5.16 as a function of reverse bias voltage. Note that because this particular coupler does not have an S-bend feed section to separate the waveguides, the two modes overlap.

bias, on the other hand, switches all of the power from one guide to the other, as shown in fig. 5.19.

It was first thought that this complete switching was due to free-carrier injection. However, a measurement of the transient response of the device, shown in fig. 5.20, showed that the rise time was on the order of 0.3 microseconds, much slower than the expected 3 to 10 nanoseconds characteristic of free-carrier recombination. Optical measurements also showed that no 0.8 μm light was being generated, implying that no free-carriers were recombining, and hence, none were being generated.

Another possibility was that the heat generated in the interguide region under forward bias was changing the index of refraction. This Joule heating, which is the sum of the IV heating in the junction and the I^2R heating in the electrode, was considerable - almost a watt over a 7 mm device length. Although 0.3 μs at first seemed too fast for thermal effects, we nonetheless proceeded to calculate the expected thermal response due to a heat pulse. The thermal diffusion equation was integrated numerically in cylindrical coordinates using techniques similar to those used in section 4.2. The results, shown in fig. 5.21, show that the response is indeed quite rapid, with a rise time of about 0.3 μs for the temperature differential between the center of each guide and the interguide region. The small cartoon at the top of the figure shows, to scale, a cross-section of the two guides and the region in which heat was (analytically) applied.

This response time is much faster than other thermal switches, which have rise times that are typically 1 ms or slower⁽¹⁹⁻²¹⁾. The short rise time of this switch is due to the much larger thermooptic coefficient and the much higher thermal conductivity of GaAs relative to the materials normally used (glass, LiNbO_3 , and polymers). The switch was not terribly efficient, requiring



ig. 5.19 Intensity profiles for three different values of forward bias current.

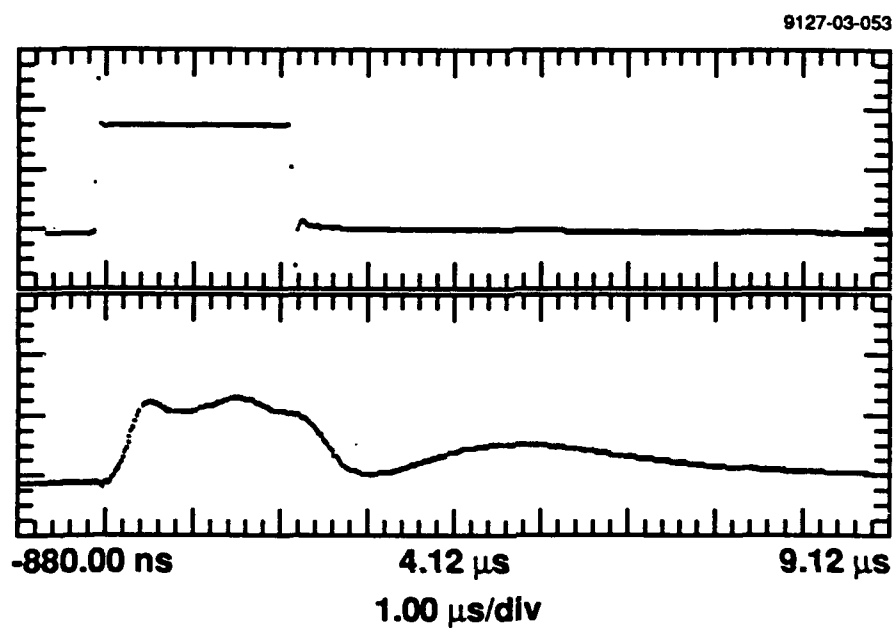


Fig. 5.20 Change in device transmission (bottom trace) for a 180 ma forward-bias current pulse (top trace).

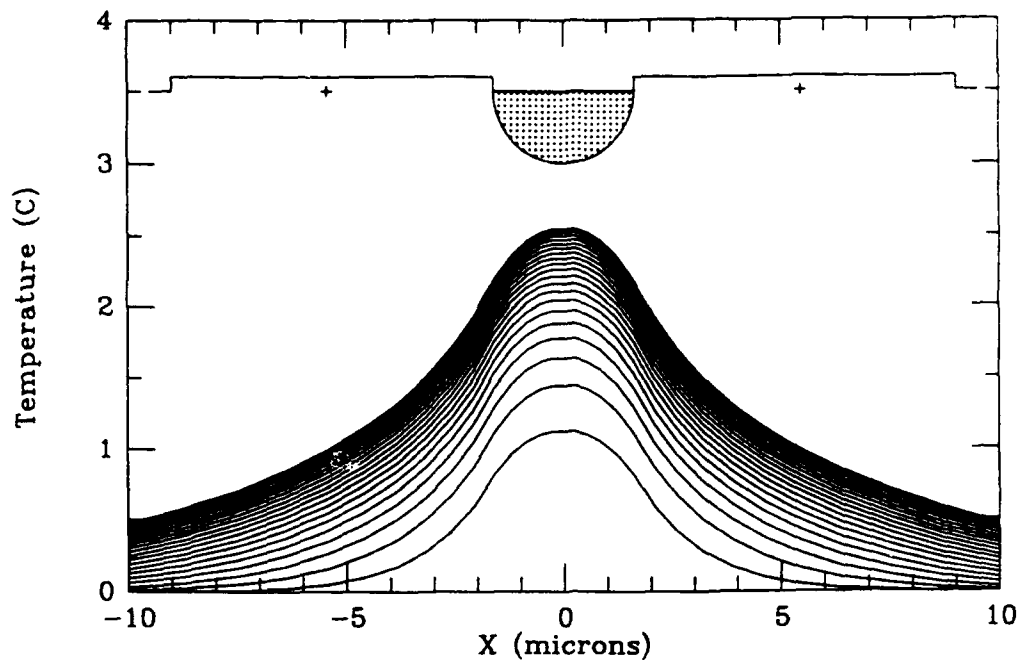


Fig. 5.21a Lateral temperature profile for a 0.8 watt heat step confined to the inter-guide (dotted) region. The time spacing between curves is 0.2 microsecond.

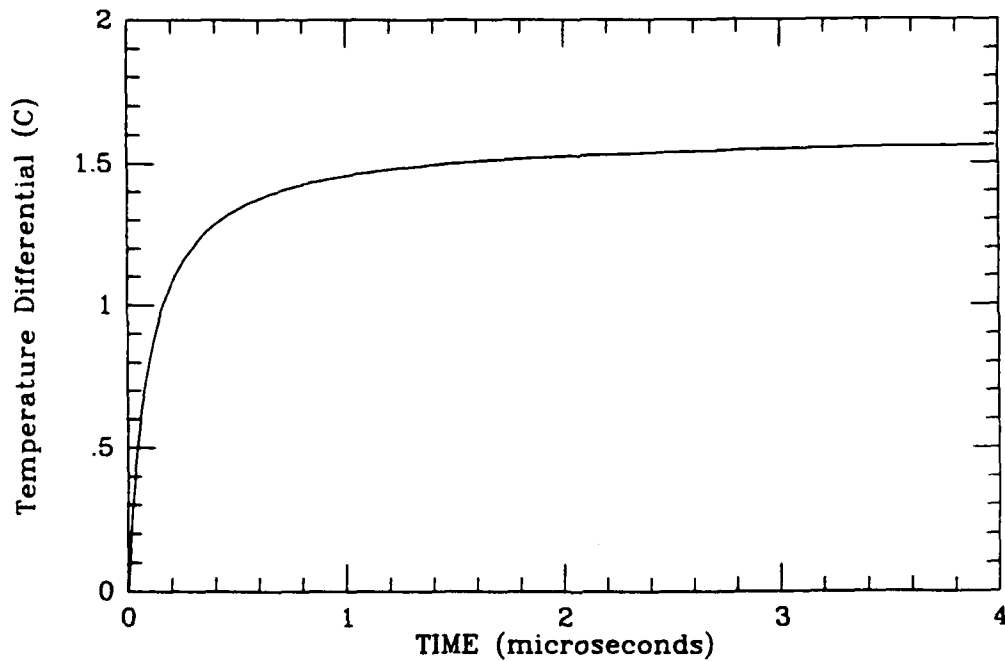


Fig. 5.21b Temperature difference between center of dotted region and middle of each guide (the small crosses at the top of fig. a).

about 0.8 watt to operate, but could nonetheless be useful in applications where moderate switching speeds are adequate. The advantages of thermal switching is that, like its PIN counterpart, it is polarization insensitive, and can produce large index changes over large waveguide dimensions.

SUMMARY

In summary, the use of free-carrier effects seems to create more problems than it solves. Large optical losses due to absorption, microwave losses due to conductivities that are too low (or too high!), finite recombination, transit and diffusion times, and the problems associated with low-impedance devices at high frequencies all seem to conspire against the free-carrier effect in such a way as to neutralize its primary benefit, the high index changes possible. It is our belief that these limitations are fundamental, and that any future device predicated on the free-carrier effect must first grapple with and overcome these limitations if it is to be successful.

REFERENCES

1. F. Stern, "Dispersion of the index of refraction near the absorption edge of semiconductors", Phys. Rev. 133A, pp. 1653-1664, March 1964
2. J.G. Mendoza-Alvarez, R.H. Yan and L.A. Coldren, "Contributions of the band-filling effect to the effective refractive-index change in double-heterostructure GaAs/AlGaAs phase modulators", J. Appl. Phys. 62, pp. 4548-4553, December 1987
3. W.G. Spitzer and J.M. Whelan, "Infrared absorption and electron effective mass in n-type Gallium Arsenide", Phys. Rev. 114, pp. 59-63, April 1959
4. C.H. Henry, R.A. Logan, F.R. Merritt and J.P. Luongo, "The effect of intervalence band absorption on the thermal behavior of InGaAsP lasers", IEEE J. Quantum Elec. QE-19, pp. 947-952, June 1983
5. S.M. Sze, Physics of Semiconductor Devices, 2nd edition, John Wiley & Sons, N.Y., 1981, pp. 76-79
6. M.J. Adams, S. Ritchie and M.J. Robertson, "Optimum overlap of electric and optical fields in semiconductor waveguide devices", Appl. Phys. Lett. 48, pp. 820-822, March 1986
7. R.G. Walker, I. Bennion and A.C. Carter, "Low-voltage, 50 Ω , GaAs/AlGaAs travelling-wave modulator with bandwidth

exceeding 25 GHz", Electronic Lett. 25, pp. 1549-1550,
November 1989

8. R.G. Walker, "Broadband (6 GHz) GaAs/AlGaAs electro-optic modulator with low drive power", Appl. Phys. Lett. 54, pp. 1613-1615, April 1989
9. K.C. Gupta, Ramesh Garg and I.J. Bahl, Microstrip Lines and Striplines, Artech House, Norwood, MA, 1979, chpt. 7
10. J.E. Zucker, M. Wegener, K.L. Jones, T.Y. Chang, N. Sauer, and D.S. Chemla, "Optical waveguide intensity modulators based on a tunable electron density multiple quantum well structure", Appl. Phys. Lett. 56, pp. 1951-1953, May 1990
11. Y. Okada, R. Yan, L.A. Coldren, J.L. Merz and K. Tada, "The effect of band-tails on the design of GaAs/AlGaAs bipolar transistor carrier-injected optical modulator/switch", IEEE J. Quantum Elec. QE-25, pp. 713-719, April 1989
12. H. Yonezu, I. Sukuma, K. Kobayashi, T. Kamejima, M. Ueno and Y. Nannichi, Jpn. J. Appl. Phys. 12, p. 1585, 1973
13. Dietrich Marcuse, Theory of Dielectric Optical Waveguides, Academic Press, Boston MA, 1991, p. 330
14. E. Kapon and R. Bhat, "Low-loss single-mode GaAs/AlGaAs optical waveguides grown by organometallic vapor phase epitaxy", Appl. Phys. Lett. 50, pp. 1628-1630, June 1987
15. H.C. Casey, Jr. and M.B. Panish, Heterostructure lasers, Part A, Academic Press, N.Y., 1978

16. M. Heiblum and J.H. Harris, "Analysis of curved optical waveguides by conformal transformation", IEEE J. Quantum Elec. QE-11, pp. 75-83, February 1975
17. E.A.J. Marcatilli, "Bends in optical dielectric guides", Bell Syst. Tech. J. 48, pp. 2103-2132, Sept. 1969
18. D.L. Scharfetter, "Minority carrier Injection and Charge Storage in Epitaxial Schottky Barrier Diodes", Solid State Electronics 8, pp. 299-311, 1965
19. M. Haruna and J. Koyama, "Thermooptic effect in LiNbO_3 for light deflection and switching", Electron. Lett. 17, pp. 842-844, Oct. 1981
20. M. Haruna and J. Koyama, "Thermooptic deflection and switching in glass", Appl. Opt. 21, pp. 3461-3465, October 1982
21. M.B.J. Diemeer, J.J. Brons and E.S. Trommel, "Polymeric optical waveguide switch using the thermooptic effect", J. Lightwave Tech. 7, pp. 449-453, March 1989

**MISSION
OF
ROME LABORATORY**

Rome Laboratory plans and executes an interdisciplinary program in research, development, test, and technology transition in support of Air Force Command, Control, Communications and Intelligence (C³I) activities for all Air Force platforms. It also executes selected acquisition programs in several areas of expertise. Technical and engineering support within areas of competence is provided to ESD Program Offices (POs) and other ESD elements to perform effective acquisition of C³I systems. In addition, Rome Laboratory's technology supports other AFSC Product Divisions, the Air Force user community, and other DOD and non-DOD agencies. Rome Laboratory maintains technical competence and research programs in areas including, but not limited to, communications, command and control, battle management, intelligence information processing, computational sciences and software producibility, wide area surveillance/sensors, signal processing, solid state sciences, photonics, electromagnetic technology, superconductivity, and electronic reliability/maintainability and testability.