

AD-A086 745

STANFORD UNIV CA EDWARD L GINZTON LAB
MICROWAVE TUBES.(U)
JUN 80 M CHODOROW

F/G 9/1

AFOSR-77-3351

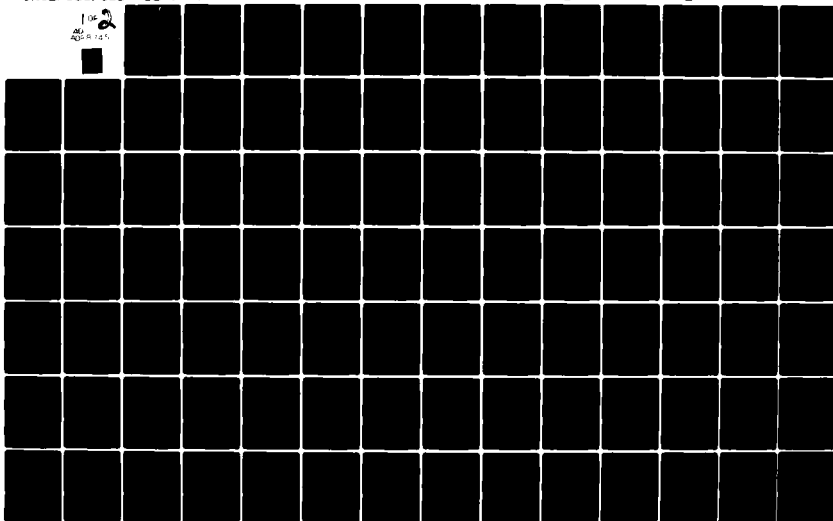
UNCLASSIFIED

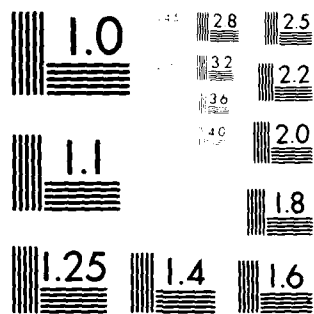
GL-3131

AFOSR-TR-80-0692

NL

100
AD-A086 745





MICROCOPY RESOLUTION TEST CHART
NATIONAL BUREAU OF STANDARDS-1963-A

AFOSR-TR- 8-0-0692

LEVEL

Edward L. Ginzton Laboratory
W. W. Hansen Laboratories of Physics
Stanford University
Stanford, California 94305

Handwritten: 12
Handwritten: A

AD A088745

AFOSR Grant 77-3351

"MICROWAVE TUBES"

Final Report for Period 1 June 1977 to 30 September 1978

Approved for Public Release
Distribution Unlimited

DTIC
ELECTE
S SEP 4 1980 **D**
C

Principal Investigator:
Marvin Chodorow
Professor of Applied Physics and Electrical Engineering
Stanford University

Prepared for:
AIR FORCE OFFICE OF SCIENTIFIC RESEARCH
Air Force Systems Command
Bolling Air Force Base
Washington DC 20332

DDC FILE COPY

2 June 1980

Approved for public release;
distribution unlimited.

80 9 2 065

AIR FORCE OFFICE OF SCIENTIFIC RESEARCH (AFSC)

NOTICE OF TRANSMITTAL TO DDC

This technical report has been reviewed and is
approved for public release IAW AFR 190-12 (7b).
Distribution is unlimited.

A. D. BLOSE

Technical Information Officer

REPORT DOCUMENTATION PAGE		READ INSTRUCTIONS BEFORE COMPLETING FORM
1. REPORT NUMBER AFOSR-TR-80-0692	2. GOVT ACCESSION NO. AD-A088	3. RECIPIENT'S CATALOG NUMBER 745
4. TITLE (and Subtitle) (6) MICROWAVE TUBES.		5. TYPE OF REPORT & PERIOD COVERED Final Scientific Report. 1 June 1977 - 30 Sep 1978.
7. AUTHOR(s) M. Chodorow		6. PERFORMING ORG. REPORT NUMBER G.L. Report No. 3131
9. PERFORMING ORGANIZATION NAME AND ADDRESS Edward L. Ginzton Laboratory Stanford University Stanford, California 94305		8. CONTRACT OR GRANT NUMBER(s) ✓ AFOSR-77-3351
11. CONTROLLING OFFICE NAME AND ADDRESS Air Force Office of Scientific Research ATTN: NE Bolling AFB, D.C. 20332		10. PROGRAM ELEMENT, PROJECT, TASK AREA & WORK UNIT NUMBERS 61102F, 2305132 (17) B2.1
14. MONITORING AGENCY NAME & ADDRESS (if different from Controlling Office) (10) Murvin / Chodorow		12. REPORT DATE 2 June 1980
		13. NUMBER OF PAGES 95
		15. SECURITY CLASS. (of this report) UNCLASSIFIED
		15a. DECLASSIFICATION/DOWNGRADING SCHEDULE
16. DISTRIBUTION STATEMENT (of this Report) Approved for public release; distribution unlimited.		
17. DISTRIBUTION STATEMENT (of the abstract entered in Block 20, if different from Report)		
18. SUPPLEMENTARY NOTES		
19. KEY WORDS (Continue on reverse side if necessary and identify by block number)		
Travelling Wave Tubes Crossed Field Devices Crossed Field Amplifier Free Electron Laser	Klystrons Magnetrons Gyrotrons Ubitrons	Microwave Tubes Cyclotron Maser Fast Wave Tubes Gyrocons
20. ABSTRACT (Continue on reverse side if necessary and identify by block number) The intent of this report is to describe in some technical detail the various kinds of microwave tubes, either now in use or projected, discuss their modes of operation, present characteristics, limitations and applications, and also to discuss in terms of their technical characteristics what may be anticipated for possible performance and/or applications of these various tube types. Since this report is intended for a range of readers of varying degrees of knowledge in the field, it will be organized in a fashion so that it will be		

UNCLASSIFIED

20. (Continued)

possible to read and get information at various levels of detail and sophistication.* Various separate sections which can be read independent of the main text provide further details of characteristics of the various tube types for those who might be interested. It is not intended, however, that any portion of this report would be suitable in itself for tube designers, and there will be categorical statements about results, methods, without necessarily giving derivations or all the logical connections between some of the statements. In many cases, plausibility will be used rather than detailed, logical arguments to prove the validity of particular technical evaluations.

We list first the various types of tubes which will be considered, with broad descriptions of their particular characteristics and possible applications. This will be followed by sections on each tube type, describing the tube in more detail. The tubes that will be described are the klystrons, the travelling wave tube, crossed field devices (magnetron, crossed field amplifier, etc.), the gyrotron (also known as the cyclotron maser), the ubitron (also known as the free electron laser), and the gyrocon.† Although these tube types are all listed here as being part of a single group of possible devices, it should be pointed out that various ones of this group have been investigated over quite different periods of time and are at very different stages of development and maturity. Therefore, the statements made about some tube types will represent the present status of a long history of research, development, and refinement and describe the current status of tubes which are actually in operation, have been produced in significant numbers and for which extrapolations are based on a great deal of experience and data. Others of these tubes have had a much more limited history, and have been explored only to a minor extent and predictions about these have to be properly qualified as representing uncertain extrapolations from somewhat meager results. In spite of this difference in the historical experience with these various tube types, we still anticipate that one might be able to make some useful comparisons about the relative merits, frequency and power ranges possible for the various tube types.

* It is anticipated that there will be a subsequent internal report issued by this laboratory which will contain specialized, detailed treatment of various tube topics. Copies of any such report would be available as a supplement to this report.

† We shall use the terminology ubitron and gyrotron since quantum effects are negligible in almost all practical cases and all the useful theoretical treatments neglect quantum effects.

UNCLASSIFIED

PREFACE

The following report is intended to be a review of the current status of microwave tubes, their modes of operation, and comparative performance, in terms of peak and average power, bandwidth, efficiency, high frequency limitations — all factors which might be of interest to either the user of such devices or to contract managers interested in procurement.

It was intended originally to include more detailed theoretical descriptions of these devices. This would of necessity have been quite mathematical, and not very illuminating, except possibly to already knowledgeable experts. Instead, there has been an attempt here to describe qualitatively the operating principles so as to demonstrate the contrasting modes of operation in as much detail as possible, as clearly as possible, which, hopefully, would indicate why various devices have different areas of application and why particular devices are superior for these various areas of application. In any case, it is true that for some devices, for example, and in particular, crossed field devices, there is no useful, i.e., mathematically simple analytic theory. One can understand the performance of crossed field devices, in all aspects quite well but, in practice, to get quantitative answers, one must use a computer which provides answers but not a great deal of generalized insight, except to an expert who is willing to examine trajectories carefully and try to understand the implications of such examination in terms of possible effects of changes in operating conditions, on performance. Even this is not too great an

aid in design. Obviously, most of this cannot be conveyed in a report of this kind.

The gyrotron and the ubitron do have a quite elegant mathematical theory which can predict results quite well, either analytically for small signal theory and without too great difficulty, by use of the computer, under large signal conditions. But this theory does not lend itself to any simple generalizations. At least for the moment, there has not emerged in the theory any simple set of parameters such as exist in klystron theory or in traveling wave tube theory which can be used to characterize a range of operation of the device. In both traveling wave tube theory and klystron theory there are impedances, beam parameters, gain parameters (Pierce's C) which can, at least indicate what the range of operation of the device will be in power, bandwidth, etc. This does not yet exist for the gyrotron and ubitron, and makes it somewhat difficult to give the same kind of general theoretical description, with useful parameters as it is for the other devices quoted. Therefore, for these two devices, the gyrotron and the ubitron, it was thought most useful to describe the basic physical concepts involved in the interaction, as to the character of the electron motion, and the relation of the motion to the interacting electric fields. as as to leave the reader with some idea at least of how these devices work, and also perhaps as to the reasons why, in certain ranges of frequency and power, they look to have advantages over the more conventional devices, i.e., the klystron and the traveling wave tube.

As to the gyrocon, this is conceptually a rather simple device. There is at the moment a few isolated experiments by the Russians under somewhat poorly defined conditions, and some numerical design calculations at Los Alamos which confirm the possibility of high efficiency.

Finally, the two largest sections in this report on the klystron and traveling wave tube, do represent the workhorses of microwave generation, for most (not all) applications, wide-band radar, countermeasures, for high peak and average power. For both of these devices there is a useful small signal theory which selects simple parameters for each device, which can characterize the performance quite easily. Such theory is easily available in most texts and, as has been stated, does work quite well in describing the possible range of operation, either in gain, bandwidth, etc.

There is an attempt in the sections in this report on these devices to say a little about their performance under large signal conditions which for many applications are actually their normal operating conditions. Also, there is an attempt to discuss some of the problems with these devices, and their limitations, and also to indicate why they begin to be less useful as one goes to higher frequencies (and why the gyrotron currently offers much better possibilities).

It should be stated, also, that there exists for both TWT and the klystron, quite straight forward theoretical approaches which can predict quite well (with the aid of a computer) their large signal performance, their bandwidth,

their large signal performance, their bandwidth, gain, output power, etc. This computer approach is used quite widely in design. Presentation of this theory would have gone somewhat beyond what seemed to be the useful scope of a report of this kind. However, some of the results of the theory and what kinds of problems it can deal with are indicated in the report.

It should be mentioned that a brief internal memo is under preparation in this laboratory which will treat the subjects of (1) gain, bandwidth, and efficiency characteristics of klystrons, and the design approach for achieving such characteristics and, (2) methods of large signal calculations for coupled cavity TWTs. Copies of this internal memo can be made available to any recipient of this report.

Accession For	
NTIS GRA&I	☒
DDC TAB	☐
Unannounced	☐
Justification	
By _____	
Distribution/ _____	
Availability Codes	
Dist	Avail and/or special
A	

MICROWAVE TUBES — A REVIEW

I. Introduction

The intent of this report is to describe in some technical detail the various kinds of microwave tubes, either now in use or projected, discuss their modes of operation, present characteristics, limitations and applications, and also to discuss in terms of their technical characteristics what may be anticipated for possible performance and/or applications of these various tube types.

Since this report is intended for a range of readers of varying degrees of knowledge in the field, it will be organized in a fashion so that it will be possible to read and get information at various levels of detail and sophistication.* Various separate sections which can be read independent of the main text provide further details of characteristics of the various tube types for those who might be interested. It is not intended, however, that any portion of this report would be suitable in itself for tube designers, and there will be categorical statements about results, methods, without necessarily giving derivations or all the logical connections between some of the statements. In many cases, plausibility will be used rather than detailed, logical arguments to prove the validity of particular technical evaluations.

We list first the various types of tubes which will be considered, with broad descriptions of their particular character-

*It is anticipated that there will be a subsequent internal report issued by this laboratory which will contain specialized, detailed treatment of various tube topics. Copies of any such report would be available as a supplement to this report.

istics and possible applications. This will be followed by sections on each tube type, describing the tube in more detail. The tubes that will be described are the klystrons, the traveling wave tube, crossed field devices (magnetron, crossed field amplifier, etc.), the gyrotron (also known as the cyclotron maser), the ubitron (also known as the free electron laser), and the gyrocon.* Although these tube types are all listed here as being part of a single group of possible devices, it should be pointed out that various ones of this group have been investigated over quite different periods of time and are at very different stages of development and maturity. Therefore, the statements made about some tube types will represent the present status of a long history of research, development, and refinement and describe the current status of tubes which are actually in operation, have been produced in significant numbers and for which extrapolations are based on a great deal of experience and data. Others of these tubes have had a much more limited history, and have been explored only to a minor extent and predictions about these have to be properly qualified as representing uncertain extrapolations from somewhat meager results. In spite of this difference in the historical experience with these various tube types, we still anticipate that one might be able to make some useful comparisons about the relative merits, frequency and power ranges possible for the various tube types.

*We shall use the terminology ubitron and gyrotron since quantum effects are negligible in almost all practical cases and all the useful theoretical treatments neglect quantum effects.

II. Overview

In all microwave tubes, the basic mechanisms for producing microwave power can be divided as follows:

1. Extracting a high density electron beam from a cathode and accelerating the electrons by d.c. voltages to relatively high velocities.* Such an accelerated beam must be controllable so as to have well-defined trajectories or class of trajectories, and usually requires some combination of focussing electrodes and magnetic fields. The magnetic fields, in some cases (klystrons and TWTs), serve merely to focus the beam, that is, confine the electron flow to a relatively narrow, straight channel, so that it can interact with suitable electromagnetic circuits along its trajectory.

The electromagnetic circuits are of various kinds, cavities, or propagating circuits and it is important to prevent interception of the beam by these circuits. In other cases — crossed field devices (magnetrons, etc.), the ubitron, gyrocon, and gyrotron, the magnetic fields affect the motion in a very particular fashion with specific trajectories involved which are intimately related to and essential for the specific mode

*Actually, in the free electron laser or "ubitron" the beam velocity very often required corresponds to such a high voltage, several millions of volts, that the acceleration of the electrons is obtained by means of a microwave accelerator. The microwave nature of the acceleration is not essential to the performance of the device as a microwave tube, it is just a convenient way of getting high energy electrons. In principle, one could use d.c. acceleration except for technical difficulties and actually using a microwave accelerator adds some complexity and possible disadvantage to the ubitron performance.

of interaction with electromagnetic circuit. This will be treated further in the discussion of the specific devices.

2. The steady state (d.c.) trajectories are then modified by modulating the electron motion, either in magnitude and/or direction of the velocity through the interaction of the electrons with time varying electromagnetic fields, produced by some particular circuit geometry, i.e., either a cavity resonator, a series of cavities, or a modified waveguide, such as a helix.
3. This time varying modulation of the velocity results in changes in the successive electron trajectories as a function of the entrance phase into the modulating field. The initially injected, uniform (constant in time) electron beam becomes nonuniform because of this time variation in the velocity, and one gets an electron beam with a time varying density (current). This conversion of a time varying velocity into a time varying current density is commonly referred to as "bunching." Some electrons overtake others because of this difference in velocity. The conversion mechanism differs widely among various tube types, and it is the particular conversion mechanism which is probably the most distinctive characteristic among the tube types. In some tubes (TWT, klystron), this is merely the more energetic electrons overtaking less energetic ones. In other tube types — magnetrons, ubitrons and gyrotrons — the presence of (steady) transverse magnetic fields possibly with spatial variation plays a key role in converting velocity variations induced by the electromagnetic field into density variations at later points in the trajectory.

In all cases, however, the resultant time varying current passing through electromagnetic fields associated with a circuit (cavity, coupled cavities, waveguide, etc.) transfers power to the field, the beam kinetic energy being transferred to electromagnetic energy in the circuit, and then to some transmission system.

4. Finally, the remaining kinetic energy of the electrons is converted into heat in some collecting electrode which may be part of the electromagnetic circuit (as in a magnetron) or a separate collecting electrode.

In many of the devices, except the magnetron type, it is in principle possible to operate the collecting electrode at a depressed potential, that is, a potential relative to the cathode lower than the potential of the main interaction regions of the device. This implies that the electrons will strike the collecting electrode with lower kinetic energy, and there is a saving in dc power. This actually occurs in practice, for example, in satellite traveling wave tubes and in many others as mentioned later. This is a general possibility for these kinds of devices.

Before discussing in detail the characteristics and performance of individual tube types in later sections, we shall describe each of them briefly here, indicating these characteristics and nature of the electronic interaction in each.

The klystron. A cylindrical electron beam, maintained at a roughly constant diameter usually by some kind of magnetic focusing, passes through a series of resonant cavities (see Fig. 1). Electromagnetic power is fed into the first cavity, so that there is a r.f. field (voltage) across the gap. This

voltage at the first gap produces a time varying velocity modulation of the electrons, their velocity as they leave the first gap, depending on the phase at which they cross the gap. In the drift tube region beyond, this velocity modulation results in bunching, that is, nonuniform motion of the electrons, with the fast ones tending to overtake the slower ones that preceded them, so that there is a time varying current at the second gap. This current drives the second cavity and produces a voltage across the gap which is higher than on the first gap, so that there is voltage gain. This voltage remodulates the beam velocity, producing greater modulation and therefore higher current than after the first gap. This process can be repeated through a number of cavities, with voltage gain at each successive step and greater current so that ^{at} the last cavity there is a high amplitude of r.f. current (the peak amplitude is of the same order as the dc current). This r.f. current drives the last cavity, producing a very high r.f. voltage comparable with the original dc voltage applied to the cathode. The last cavity is coupled to a transmission line of some kind, so that most of the power usually delivered to the last cavity is transmitted through the line to a useful load, an antenna, etc.

In the earlier cavities, the power delivered to the cavities is dissipated entirely in the resonator r.f. losses, but the sole purpose of the cavity beam interaction in these earlier cavities is to produce a higher voltage than in the previous stages, so

as to modulate the velocity of the electrons more heavily so as to produce larger currents beyond, through the bunching process.*

Finally, the beam after passing through the last cavity, passes into a (usually) cylindrical container whose size and properties are determined entirely by the amount of beam power it has to dissipate. Much of this description, with modifications, will apply to other devices which we shall describe briefly here. The klystron, however, is the simplest conceptually, largely because of the geometrical isolation of the various functions, that is, beam production (cathode anode), velocity modulation, bunching, power delivery, beam collection.

Traveling Wave Tubes. A second kind of microwave tube of great importance and which probably comprises the majority of devices in current use is the traveling wave tube (TWT). In this device, an electron beam is injected into an electromagnetic circuit in which there is a traveling electromagnetic field. Typically, the circuit consists of a periodic metallic structure in which the traveling wave is periodic and usually not simply sinusoidal. However, the wave can be represented by a sum of sinusoidal components. The geometry is designed, so that at least one

*At each cavity, the phase relations between the gap voltage and the r.f. current in the beam is such that there is an average loss of energy from the electrons. This, of course, is necessary to provide the power for the cavity excitation. The exact phase relation that exists depends on the detuning of the cavity resonator, i.e., the difference between the operating frequency and the resonant frequency of the cavity.

component of the wave traveling along the structure will have a low velocity, i.e., lower than the velocity of light and, therefore, an electron beam injected along the axis of this device can travel in approximate synchronism with this component. In these circumstances, if one injects some power into this waveguiding system, the electrons moving along the axis of the system will interact with the electric fields of that component of the traveling wave traveling in synchronism with the electron beam and there will be a cumulative interaction. Initially, the electric field of that component either accelerates or decelerates electrons depending on the particular phase of the field. The effect of this field then is similar to the gap field of a klystron. Electrons which are accelerated will tend to overtake electrons ahead of them which have been decelerated, both kinds tending to have a relative motion toward a field point of zero amplitude (zero phase). Thus, the field is continually producing bunches through velocity modulation. If the average (d.c.) electron velocity is actually slightly faster than the component producing this bunching, the bunches as they form will tend to drift into a decelerating region of the wave, they lose energy which is transferred to the electromagnetic wave. The net effect is that the amplitude of the wave grows. This is a continuous cumulative process of velocity modulation, bunching, and energy transfer which proceeds along the whole circuit, the wave grows so that at the output end where it feeds a transmission system, it has much larger amplitude than at the input. At the end of the circuit, the electrons pass into a collector just asⁱⁿ the klystron.

We can see that this device embodies the various functions common to all microwave tubes but, unlike the klystron, for example, all these processes are happening continuously along the whole device. The advantage of this as compared to the klystron is that a propagating circuit such as used for the interaction here can have a bandwidth much greater than a resonant cavity, such as used in the klystron. The bandwidth in this case is determined not by an impedance-frequency characteristic, but by the rate of change of the velocity as a function of frequency. If the velocity of the wave is approximately constant over a considerable range of frequency, one can get the interaction describe over this whole range. This characteristic is of tremendous importance. For example, for a particular type of circuit, the helix, bandwidth on the order of two to one, can be easily obtained. For other kinds of circuits, known as coupled cavity circuits, which are essentially klystron-like resonant cavities coupled electromagnetically to their neighbors through apertures, the bandwidth will be much less, on the order of ten or fifteen percent. The latter circuit, however, being all metal can handle much higher powers than the helix which is a fragile, coiled wire inside of a dielectric envelope, and cannot easily be cooled.

In spite of these limitations, however, for some applications, traveling wave tubes are very important. The coupled cavity circuit can operate at very high powers, comparable with klystrons, though their efficiencies are almost always not as good. The helix at lower power levels can have the bandwidth quoted and is still capable of operating at powers from several

watts up to several kilowatts of average power. Helix tubes, as an example, constitute all of the communication tubes used in satellites. Coupled cavity tubes represent a majority of the tubes used in airborne radar.

There are problems with traveling wave tubes not common to klystrons, for example, the the circuit extends from input to output, there is a possibility of oscillation due to reflections and, therefore, one has to provide attenuation along the circuit to prevent reflected signals from the output arriving at the input with large amplitude. This leads to deterioration in performance, and adds complication. It is a problem which has been solved, but not completely satisfactorily. Gains typically are in the range of 25-40 db. Much larger gains offer difficulties associated with reflection/attenuation problems.

Crossed Field Devices. This constitutes a class of microwave tubes often referred to as magnetrons, which are the most common devices of the class. The magnetron, which has the major characteristics of all the others, consists of a hollow cylindrical structure which constitutes the anode, and a central electrode concentric with it which is the cathode. The anode will have, typically, radial slots cut into it, equally spaced around the inner circumference which constitutes the electromagnetic circuit, and over some frequency range a wave can travel circumferentially around the internal slotted structure with the slot depth and spacing determining the velocity of this wave. Each slot can be considered as a shorted section of planar transmission line, open at the inner radius so that there is fringing electric

field at the open end extending out into the interelectrode space between the anode and cathode. See Fig.

The device is operated with a potential applied between the cathode and the anode, giving a radial electric field E and an axial magnetic field B_0 applied parallel to the cylindrical surfaces of the cathode anode. Therefore, the motion of the electrons is not purely radial as it would be if only the radial electric field were present. The axial magnetic field produces a transverse component of the motion, and one can show that the resultant motion is a superposition of a circumferential drift, $U_0 = E/B$ which is the important characteristic of the motion, plus some cyclotron motion characteristic of the magnetic fields. By proper choice of the voltage (electric field) and the magnetic field this average angular drift of the electrons can be adjusted so that it is synchronous with the wave traveling around the interior circumference of the anode, and there is interaction between them just as in the traveling wave tube and the statements made about the traveling wave tube would apply here. However, the bunching and motion of the electrons and the interaction with the r.f. field is more complicated because the effects of the r.f. fields are drastically modified by the presence of the axial magnetic field, and the subsequent motion is of the same kind as arising from the superposition of the static electric and magnetic fields.

Basically, we can state without proof that those electrons which are so placed in the electromagnetic field that they lose energy will tend to drift toward the anode because of the combined effects of this r.f. electric field which is retarding them, and

the axial magnetic field. So there is a gradual transfer of energy through the electrons in which they lose energy to the circumferential r.f. electric field and gain it by drifting toward the anode.

A net result of these combined effects of the axial magnetic field, radial d.c. electric field and the angular component of the high frequency electron field is that there is a very efficient transfer of energy to the electromagnetic field obtained from the motion of the electrons toward the anode. The resulting efficiency for such devices tends to be quite high, of the order of 80-90%. A disadvantage is that whatever energy is left in the electrons has to be dissipated on the anode so the anode which is also the r.f. circuit also acts as the collector and, therefore, there are limitations on average power. The device can be used for high peak power, but it is not as good for high average power.

What has been described, so far, is a completely closed cylinder as the anode, which means that the total circuit is a closed loop and obviously has built in feedback, and one has an oscillator. One can use the same principle for an amplifier by interrupting the anode in a particular way at some point along its circumference and isolating a separate input point and output point for the circuit, and arranging for the electrons to circulate continuously around so as not to provide feedback. Under these circumstances, one can get an amplifier also with good efficiency, but because of the inherent possibilities of the feedback via the electrons and other characteristics, one cannot get large gain and the small signal behavior is not very good.

There are also linear variations of this device where the anode faces a negative electrode called the sole, which is not the source of the electrons. The electrons are injected into the region between the anode and sole^{and}/propagate parallel to the circuit under the action of the static electric and magnetic fields. This provides isolation between input and output, and is not subject to the limitations mentioned above.

However, all of these devices, because of instabilities in the electron beam in the presence of crossed electric and magnetic fields, can be very noisy and can have oscillations occur in some cases which are not related to the electromagnetic structure. Therefore, their use as amplifiers has to be very carefully limited, although within those limitations they can be very useful.

Fast Wave Tubes. In the devices described up to now, the klystron, the traveling wave tube, and the magnetron, it is always the case that there is a requirement for electrons, with relatively low velocities, usually considerably less than the velocity of light, to interact with a synchronous electric field, and always in geometries where the electric field is the fringing field due to some electromagnetic circuit, i.e., it may be the field inside the cylindrical gap of a klystron resonator, or the fields in the interior of a helix, or the gap fields^{of}/a series of resonators, as in the coupled cavity traveling wave tube. In all these cases where the electric field has to have a slow wave component for synchronous interaction, the general nature of Maxwell's

equations require that that component have its maximum amplitude near the walls and, therefore, the electrons have to be near the walls, e.g., if one makes the tunnels of cylindrical regions too large, the fields on the axis where the beam might be, becomes intolerably small.

The net result is that for a slow electron beam which interacts with a slow electromagnetic wave, one is always faced with the problem of having to have a beam which is narrowly confined by its metal boundaries. Stated otherwise, one cannot afford to have very large tunnels with a small beam and still get very large interaction. All of these terms, large and small are all measured on a scale, proportional to the velocity of the electrons and the wavelength. The result is that if one is confined to reasonable velocities (up to beam voltages of, say, 100 kv) as one goes to shorter wavelengths, one is faced with a problem of having to get appreciable density of an electron beam through apertures of decreasingly smaller radii (proportional to wavelength), and this is the restriction which makes it very difficult to go to very high frequencies. At five or six millimeters, for example, aperture sizes become so small that the tubes really become mechanical marvels, and the amount of current one can transmit is small.

We would like to describe briefly in this section, and more extensively in a later section, a class of devices known as fast wave devices, which circumvent this problem. Basically, these devices employ modes of interaction in which the electron beam can have a velocity much less than the velocity of light

and still get a cumulative interaction with an electromagnetic wave whose phase velocity is greater than or equal to the velocity of light ^{get} and/amplification oscillation, etc.

The basic idea is to put a periodic transverse motion into the electron beam, either by using an array of alternating magnets or by injecting the electrons into an axial magnetic field, so the electrons rotate at the cyclotron frequency. The electromagnetic wave is moving past the electrons, but the frequency of the wave experienced by the electrons in that case will be less than the laboratory frequency, since the electrons are moving as well as the wave, and the number of waves per second seen by the electrons (the doppler frequency) is less and can be much less than the actual laboratory frequency. If the condition is met that this frequency, as seen by the electrons, coincides with the periodic frequency of the electron transverse motion, one can get a cumulative interaction. Whatever net interaction occurs in one cycle of the electromagnetic phase will be repeated many times because the phase relation of the e.m. wave to the periodic motion is always preserved. One can then get a cumulative action, electrons can undergo modulation in their velocity, they will undergo bunching, which may be axial or transverse, and, as a result, one eventually gets to a situation where these bunched electrons still undergoing their periodic motion can interact with the electromagnetic wave passing them, so as to deliver energy to the wave.

There has been no detail of the possible devices given here, that is given in a later section. Here we can only say that in the ubitron (the free electron laser) this transverse

periodic motion is produced by passing the beam along the axis of a periodic array of alternating magnets, and the electron undulates as it passes through the array. In the gyrotron the electrons are injected into an axial magnetic field and rotate at the cyclotron frequency. In either case, by picking the parameters properly, one can have an electromagnetic wave propagating in the same direction as the electrons with a velocity equal to or higher than the velocity of light interact in the same way with any given electron on each successive cycle of the electron motion, so that wherever the net effect of this interaction is (which will differ for different electrons) it accumulates and one eventually produces current in the beam, which delivers power.

These are the so-called fast wave tubes. They circumvent the difficulty mentioned above about having to go to intolerably small dimensions to reach millimeter waves. Both of these devices seem to be very promising candidates for going to higher frequencies than with slow wave tubes. As to data, the ubitron in its original form operated at about six millimeters, and a later version which was called the free electron laser, operated at several microns, with hundreds of kilowatts in the first case, and tens of kilowatts in the second, both cases being pulsed.

In the gyrotron, the best performance in the U.S. has been something over 200 kw cw at 30 GHz, using a structure which obviously can be scaled to much higher frequencies, perhaps 2-5 millimeters, without running into electron beam interception difficulties.

The Gyrocon. Finally, we would like to mention briefly a tube which has recently been suggested, and there have been some rudimentary, isolated experiments done primarily by the Russians which are worth mentioning.* It is a device which is particularly suited for relatively low frequencies and where one is willing to use very high voltages and low currents. What one gains possibly is very high efficiency. Theoretical values of about 90% have been quoted, and the few isolated tests have indicated perhaps a little over 80%. Simply stated, one produces an electron beam which then passes through a cavity which has a circularly polarized transverse electromagnetic field. This field produces a transverse deflection of the beam in which the direction of deflection rotates. This is the same kind of deflection one gets in a cathode ray tube if an electron beam passes through two pairs of deflecting plates which are run in quadrature and the pattern on the screen is a circle. Just as in that case successive groups of electrons move along trajectories which are straight lines drawn from the deflection system to the target. All of these trajectories lie on a cone with its apex at the deflector and its base at the target. It is possible by using a subsequent solenoid (the outside ^{of a} solenoid) to increase and magnify these deflections so that for a given amount of r.f. deflecting field (which can be either magnetic or electric) the subsequent total deflection of the beam can be increased, i.e., trace out a bigger circle at the target plane.

*There were some much earlier patents and some limited research in this country a number of years ago.

When the beam has reached a sufficiently large deflection (and it is to be remembered that one then has essentially groups of electrons whose point of impact moves in a circle), one puts at the locus of that circle a waveguide which is bent in a circle with a slot in the top of the waveguide. The phase velocity of the guide is matched to the velocity of the moving spot around the target circle. Note, this is not the velocity of electrons or any component of that velocity. This spot velocity can be greater than the velocity of light (and is so in this application). This is also true of the motion of the spot on a cathode ray screen. The electrons will pass through this slot, their point of entry moving at the spot velocity and can excite a running wave in the waveguide whose peak phase arrives at each point in the waveguide just as the electrons are arriving there. Therefore, there is a transfer of energy from the electrons to the wave, and in principle if one can bring the electrons to rest just as they leave the guide, one will be extracting all of their energy and one will get 100% efficiency (minus the cavity losses). If one puts in the various effects which have to be taken into account, one still gets perhaps 85%, but these numbers are achieved at voltages of about 300 or 400 kv, and perhaps 10 or 20 amperes which are quite different voltage to current ratios than for the kinds of tubes we have been discussing.

The principal merit of this device is its possible high efficiency. In some special applications where high power, particularly for cw applications is important, this could be a useful device. Present interests are largely for applications to accelerators where one requires very large average powers, at low frequencies.

III. Klystrons

The first tube type we shall describe will be the klystron. This was invented about 45 years ago and there was considerable activity on this tube type before and during WW II but, for various technical reasons, much of this activity was confined to a particular version known as the reflex klystron which has only special applications, either as a local oscillator at milliwatt levels for receivers, or as an easily modulated oscillator at power levels of the order of a watt for microwave relay links.

Major klystron development took place after WW II, and the klystron is now the best understood, most predictable in performance of all tube types, and is now preeminently the tube most commonly used in all cases where high average or high peak power are necessary. With the exception of some very recent work on a device which we shall discuss later (the gyrotron) all outstanding performance in terms of very high average power and high frequencies has been achieved by means of klystrons. For peak power, also, all the highest performance has been achieved with klystrons. The reason for this predominance of the klystron in power handling capabilities (as compared particularly to crossed field devices) is its geometry which is much superior for handling either high peak or average power because of dissipation capability, voltage breakdown, etc. It should be stated that the gyrotron, because of its geometry, will most certainly be able to reach higher average power at higher frequencies than the klystron.

Power. It would be of little value to list here all the possible tubes and all possible applications for klystrons. A typical klystron (amplifier) ^{See Fig. 1} is a multi-cavity device with 3 to 6 cavities, the number of cavities determining both the gain and bandwidth. We list some typical characteristics. In peak power, tubes have operated at 40 or 50 megawatts at about 3000 MHz. These have been pulse tubes with low duty cycle, perhaps 6×10^{-4} . The principal application for these has been for electron accelerators. Such peak power is exceptional, not because it is unusually difficult to achieve, but because most applications, other than accelerators, have not required this kind of peak power. Most standard performance for radar applications range from 0.5 megawatts up to 10 - 15 megawatts in pulse power and average or cw up to several hundred kw. The frequency range for such tubes extends from relatively low frequencies — 300 to 400 MHz up to 10,000 MHz. As an example, we quote a particular tube which has been used in the BMEWS system at 400 MHz. This operates at 1 Mw peak power, 75 kw average, 1 millisecond pulses. Parenthetically, tubes of this kind have lasted under continuous operation for upward of 15 years. Much higher peak and average powers are available in other tubes and at higher frequencies.

Another typical performance is that of a c.w. tube operating at about 450 kw at about 2000 MHz, which is used in the Goldstone transmitter. There has also been unusual performance in a laboratory at about 8 GHz of about 1 Mw c.w. This was only for a limited time period and it was never

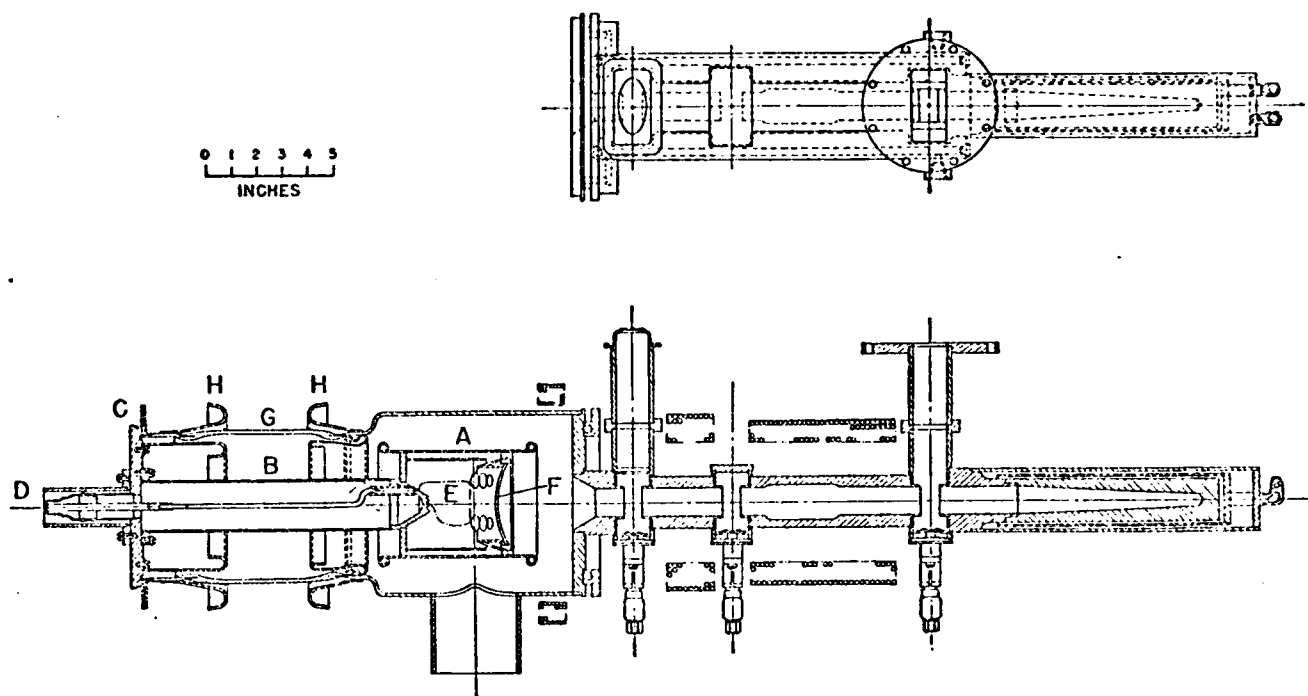


Fig. 1 Simplified assembly drawing of the pulsed klystron amplifier.

actually operated outside of a laboratory. It is some indication of what is possible. An equally interesting performance of a klystron which actually is in use is 250 kw c.w. at about 8 KHz, which is used in the Haystack transmitter for planetary radar exploration.

Gain and Bandwidth. Typical gain for a multicavity klystron ranges from 30 to 70 or 80 db, with bandwidths ranging from a fraction of a percent to 7 or 8%. Such bandwidths are achieved by suitable stagger tuning and spacing of the various cavity resonators of the tube.

In general, for any cavity resonator in a klystron (other than for the first where the gap voltage is produced by the power input from a transmission line) the cavity voltage is generated by the high frequency electron current arriving at that cavity, which was produced by modulation by the previous cavities. The amplitude and phase of the cavity voltage generated by the beam, relative to the phase of the r.f. current in the beam, is determined by the impedance of the cavity. In a precise way, the cavity can be considered to behave like a resonant circuit driven by a current generator (the electron beam), with the phase and amplitude of the voltage relative to the phase and amplitude of the current depending on the circuit impedance and resonant frequency. (See Fig. 2.)

To get bandwidth, one has to stagger tune all of the cavities preceding the last (output) cavity, just as one stagger tunes the various circuits of a standard wideband,

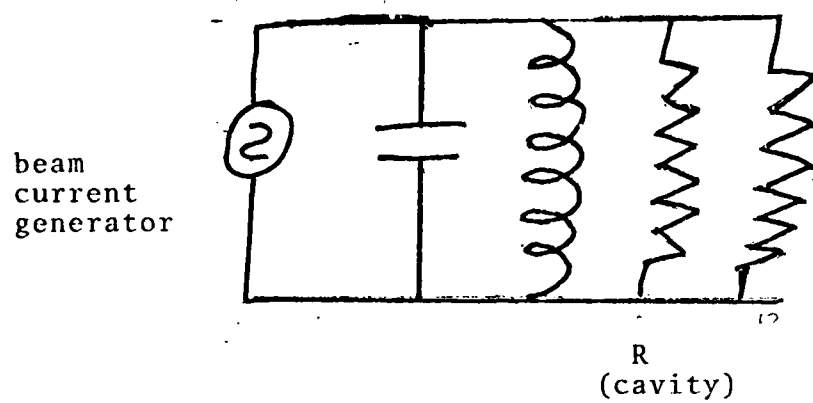


Fig. 2 Equivalent circuit for beam cavity.

multi-stage, lumped constant, video, or i.f. amplifier, where one gets bandwidth by proper tuning of the various stages. There is one major difference (and difficulty) for the klystron amplifier, that is, that the r.f. electron current driving any single stage (cavity) is the sum of the separate currents generated by all of the previous cavity voltages. This is the same behavior as if one had the input to any stage of an i.f. amplifier (as an example) being the sum of all of the outputs, with various phase shifts, of all of the previous stages, instead of the almost universal case of the input to any stage being the output of the previous stage.

This phenomena in the klystron, of the driving current for the nth cavity being the sum of the individual excitations (currents) from all of the previous cavities with various phase shifts (corresponding to different transit times between the nth cavity) and the preceding cavities, makes the gain bandwidth formula much more complex. In particular, the output characteristic is a much more sensitive function of the resonant frequencies of the various cavities. For example, by an improper choice of some resonant frequencies and cavity spacings, one can actually put a zero in the gain bandwidth characteristic. However, even though the calculation is more difficult, it is possible to select cavity parameters and cavity spacing so as to get satisfactory and predictable gain bandwidth characteristics.

As a general controlling parameter, it should be pointed out that all impedance levels in a klystron vary monotonically with (actually closely proportional to) a quantity associated with the electron beam current called the beam impedance, which is the ratio of the beam voltage to the beam current (V_o/I_o). In any such tube $I_o = KV_o^{3/2}$ where K, known as the perveance varies from about 0.5 - to 2.0×10^{-6} . A high voltage tube (100 kv) will have a lower value of this beam impedance than a lower voltage (10 kv), tube with the same K. Therefore, in a high voltage tube, all impedance levels will be lower and the corresponding Q's and, therefore, the bandwidths correspondingly greater. Typical examples

10,000 V - 1 A	$R_o = 10,000$ ohms
100,000 V - 30A	$R_o = 3,000$ ohms

therefore the corresponding bandwidths for the latter case will be about three times as great as for the former. These statements are not intended, obviously, as a complete design recipe, but provide some reasons for the various numbers to be quoted for typical bandwidths.

In general, one can increase bandwidths (and gain) by increasing the number of cavities. One will still be restricted by the beam impedance, so the range of bandwidths will vary from 1% or so for a 10 kv, 1 amp tube to 6 to 8% for a 30 to 50 amp, 100 kv tube.

One more comment is relevant. To achieve the high bandwidths quoted, say 5% or higher, it is often necessary to design the loading of the output cavity (since it must have a sufficiently broad band impedance), so that its impedance seen by the beam has a broader band than simply that of a low Q broadband, single resonant circuit. By suitable external loading, i.e., proper coupling structure and use of a suitable, filter network in the output transmission line, one can get the impedance (response characteristic) of the output cavity as seen by the beam, to be, for example, that of a critically coupled double tuned circuit or something more complex. This amounts to modifying the simple circuit shown in Fig. 2 so that the equivalent load coupled to the resonant circuit is something more complex than a simple resistance, e.g., a suitably terminated filter with its own double or triple resonant characteristic. This can be accomplished quite easily and results in the kinds of bandwidths for the overall device that have been quoted.

Currently, it would certainly be true that for any application ranging from about 300 MHz up to 10 GHz, requiring either high peak and/or average power and reasonable bandwidths (something less than 8% or so) a klystron amplifier would be

the most straightforward means of achieving such performance.

We should point out that the lower value of frequency quoted above is not a physical limit, but merely represents the fact that other physically smaller tubes, triodes and tetrodes, might be feasible alternatives. There is no physical frequency low limitation on klystron use. The average power and peak power become easier to achieve as one goes to the lower frequencies. We shall discuss this in more detail later.

On the other hand, there are high frequency limits on operations for klystrons (or any tubes) due to definite physical limitations in size. As one increases the frequency, the dimensions of the tube and, in particular, the tunnels through which the electron beam passes, become smaller, proportional to the wavelength and the design problems become more difficult because of possible beam interception. This will again be considered when we discuss other devices which are not limited in the same way, such as the gyrotron.

The average power capability, in general, on any device decreases because of the possibility of melting due to beam interception. It might be pointed out, for example, that the 250 kw tube referred to above at 8 GHz has something of the order of 600 kw of beam power passing through a beam tunnel which is of the order of 0.5 - 1 cm in diameter.

Efficiency. The efficiency of standard klystrons, depending on the application, etc., ranges from perhaps 40% up to 65%. The efficiency will be somewhat lower if one is trying to get a wideband tube of, say, greater than 4 or 5%

bandwidth, since the tuning of the cavities and, more important, the optimum ratio of beam voltage to beam current is not the same for high efficiency as it is for high bandwidth. Specifically, the theory and experiment indicate that for large bandwidths (and a given dc power input), one wants as large a ratio of current to voltage as possible, whereas, for the best efficiency, one wants as small a ratio of current to voltage, assuming in either case that one is keeping dc input power constant.

The best efficiency achieved, so far, has been 74% at 2 GHz. This efficiency was obtained using an additional cavity along the beam which was tuned to the second harmonic of the operating frequency. The theory (and experiment) indicates that such a cavity properly placed and properly tuned can enhance the current amplitude at the operating frequency, and produce better efficiency. Stated more precisely, the harmonic content in the modulated beam at the fundamental (operating) frequency will be higher if one can put a second harmonic modulation on the beam of the right amplitude and the right phase, and thus produce better efficiency. The experiments verify this theory quite well. The theory actually predicted 85%, but the discrepancy between the theory and the experiment (74%) are understood. It is probable that if one is willing to go to a higher voltage and lower current than in the tube which has just been described, one could get a better efficiency, 80% being quite plausible.

It should be pointed out that all of these efficiency numbers refer to a design in which the electron beam, after

passing through the final output cavity and having generated all of its rf power, is collected in a beam collector electrode, operated at the same dc potential as the body of the tube (relative to the cathode taken as zero). This statement refers to the fact that in many tubes, operation involves something known as a depressed collector, which enhances the overall efficiency, i.e., the collector electrode on which the electrons impinge after passing through the interaction circuit, is run at a lower potential than the interaction circuit. Therefore, electrons are retarded and lose some of their kinetic energy before striking the collector, and the efficiency of the overall device is improved. Such depressed collectors are common in relatively low powered helix tubes, and can prove overall efficiency from a normal 20% or so, which represents the interaction efficiency (conversion of dc electron beam power to rf power) to perhaps 40%. The extra 20% is retrieved by the depressed collector. This standard technique involves the complication of having two, or sometimes more voltages* applied to various parts of the tube (relative to the cathode potential). This means that, as is usual, the body of the tube (the rf circuit) is usually run at ground potential which makes it convenient to connect to waveguides, coax lines, etc., which are also at ground, the cathode is run negative, and the collector is run somewhat negative relative

*In some complex collector designs, one has several collector electrodes operated at different voltages so as to collect various groups of electrons at the lowest voltages possible.

to the interaction structure, but positive with respect to the cathode. Although this technique is very useful and can improve efficiency considerably, it becomes less and less effective as the interaction efficiency of the tube improves, the interaction efficiency being defined in terms of the actual conversion of dc to rf energy in the interaction circuit.

The whole problem of depressed collectors will be discussed again in the section on traveling wave tubes. We mention it here, to point out that in these high efficiency klystrons this technique is less applicable because there is less energy left in the beam after the interaction, the residual electron energies are spread over a wider range and it is not as easy to recover some of this energy by retarding (depressed) electrodes as in the case of lower intrinsic (interaction) efficiency. However, there have been some calculations made which do indicate that even in a high efficiency klystron, such as described above, operating at 75% or 80%, perhaps an additional 5 to 10% of the total energy could be recovered with a depressed collector electrode and, therefore, a 75% efficiency tube could be raised to perhaps a total of 85%. It should be emphasized that this has not been done in practise, but in terms of future development this might be of interest.

We would like to go back, briefly, to the question of klystrons at lower frequencies than have been referred to extensively above. It should be stressed that klystrons

would be just as efficient at such lower frequencies and their power handling capability would be even greater than at higher frequencies. There is no frequency limitation in going down in frequency, either in efficiency or power handling. Actually the converse is true as far as total power output is concerned. The problem is, as one goes to lower frequency (longer wavelengths), tube dimensions which are proportional to operating wavelengths become larger. A low frequency klystron just becomes large whereas triodes or tetrodes, in particular, do not have to be as large and these tube types do not have some of the limitations characteristic of them at higher frequencies. However, triodes and tetrodes typically have much less gain than klystrons, and therefore require more tubes in a chain to get a given gain and there is a trade-off involved in any possible applications. There is evidence that even in very low frequency applications for very high power, particularly for accelerator storage ring applications, even at 200 or 300 megacycles, there has been a preference for klystrons despite their relatively large size and despite the fact that presumably space charge control tubes would work quite well. The advantages, presumably, are efficiency and the much higher gain^{than} for the triodes and tetrodes, the high gain resulting in a much less complex system. If one wanted to use triodes or tetrodes with low gain, one would require a chain of such tubes to get from some nominal power to the megawatt range, and the connections between these various tubes and all the other circuitry

apparently adds enough complications to counterbalance the advantage of reduced size as compared to klystrons in which one can get 60 or 70 db quite easily in one envelope, with no problem of connecting tubes to each other, separate power supplies, etc.

The principal application of klystrons, such as have been described, are in high powered radar, both c.w. and pulsed, in c.w. ground transmitters of medium bandwidth for satellite communications, in almost all UHF TV transmitters and, in general, in any medium bandwidth application where high power, high gain, stability, freedom from noise, freedom from oscillation, are important. These properties of noise and oscillation, we will discuss in somewhat greater detail for other tube types in which these qualities are not as good. Specifically, the noise spectrum, for example, for a klystron is much better than any other existing tube in terms of noise power and frequencies near the carrier, either AM or FM.

IV. Traveling Wave Tubes (TWT)

The next device we would like to describe is the traveling wave tube (TWT), sometimes called the O type traveling wave tube, in contrast with the M type. In the M type there is a magnetic field perpendicular to the electron motion which affects its performance drastically. In the O type, there may be magnetic fields parallel to motion for focusing, and the beam configuration is very similar to that in a klystron. The difference between the TWT and the klystron is that in the TWT the beam passes along the axis of an electromagnetic circuit in which an electromagnetic wave is propagating together with the electron beam.* The circuit is so configured that there is some significant amplitude of the electric field of this electromagnetic wave in the same region which is traversed by the beam and there is interaction between the beam and the electric field of this traveling electromagnetic wave.

Before discussing the interaction between the electron beam and the traveling wave, and describing some of the details of the various kinds of electromagnetic configurations used, it may be worth stating some of the broad features of all traveling wave tubes. Typically, they are electromagnetic

*Of course, there are electromagnetic circuits (the resonant cavities) distributed along the beam path in a klystron, but these are uncoupled separate circuits and a signal is transferred from one circuit to others only by the electron beam. There have been applications in klystrons of the use of short-circuited lengths of propagating circuits as so-called extended interaction cavities, but this just constitutes a longer resonant cavity of a more complex kind, and coupled to other resonators only by the beam.

systems with some periodic structure which results in control of the propagation behavior of the waves, so that the propagation is different than in a uniform transmission line or waveguide.

The purpose of any particular configuration is always to slow down the propagation velocity of some component of the traveling wave so it is possible for the electron beam, traveling at a velocity less (usually much less) than the velocity of light, to interact cumulatively with the propagating electromagnetic wave. Since the typical construction is to introduce periodicity in the propagating circuit, one usually talks of a periodic circuit. The effect of this periodicity is always to change the propagation characteristics of the wave, so there is at least one spatial component in the region along the beam path (not a frequency harmonic) which propagates at something close to the velocity of the electron beam. We shall not go into mathematical details here.

Qualitatively, we can say that the introduction of a periodic structure in a propagating circuit results in introducing a periodic modulation in the propagating wave which is not simply related to the free space wavelength, but is related to the periodicity of the structure. This is, in some sense, an exact statement except that the modification of the waves by the periodic structure may be so gross that one cannot simply or adequately describe it in terms of a small perturbation of the waves on a nonperiodic structure. The most common examples of periodic circuits used for microwave tubes, namely the

helix and the so-called coupled cavity circuit, exemplify this statement in an extreme way.

A more general approach to describe a typical periodic circuit is to use the terminology of filter theory or periodic network theory. One talks about a periodic network of circuit elements (more broadly described as electromagnetic elements), in which there is a phase shift in circuit voltage and current and fields as one goes from element to element, the phase shift depending on the frequency and the geometry of the individual elements. The existence of this periodicity in the voltage and fields means that the electric fields inside of any given element (particularly the field of that element which will be exposed to the beam), will have characteristics which depend on the original frequency of the electromagnetic wave injected, the geometry of the element, and the periodicity of the structure. Mathematically, one can describe this field by a product of a factor which represents the phase shift from element to element and another periodic term which represents the characteristics due to the detailed configuration of each element.

The net result, and this is the important result, is that the total field along the line of motion of the electron beam, which is the one which is relevant for the interaction process can be described as a superposition of an infinite sum of terms which are characteristic of the periodicity of the device, multiplied by an exponential phase shift factor which varies with frequency and is the phase shift characteristic of that particular periodic structure.

All of this infinite set of terms are said to constitute a set of Space Harmonics. The mathematical expression for the field E_z at the location of the beam is written below,

$$E_z(r, z) = e^{-j\beta_0 z} \sum A_n(r) e^{-j\frac{2\pi n}{L} z} \equiv \sum A_n e^{-j\beta_n z} \quad \beta_n = \beta_0 + \frac{2\pi n}{L}$$

where the period of the circuit is L , and the exponential term in front of the summation depends on the frequency and the nature of the circuit. All of the terms together describe the electric fields propagating along the circuit. The important characteristic for traveling wave tube operation is that the propagation velocity of the electron beam must be such that it is synchronous with one of the terms of this summation. One can define the so-called beam propagation constant given by $\beta_e = \omega/v_0$ where v_0 is the beam velocity. As one can see from looking at the summation, there is an infinite number of propagation constants for the various space harmonics, and if one of these is equal to or close to the beam propagation constant, there will be approximate synchronism. It should be stressed here in using this summation to describe the fields on the circuit, that the whole summation constitutes the field which exists and propagates for any given frequency. The separate terms in the summation are not individual modes. They are merely a convenient mathematical way of describing the total field configuration which describes the propagating wave on the circuit.

The one term, the one Space Harmonic in this infinite sum which propagates at the same velocity as the beam is the handle, as it were, by which the beam affects the circuit and

vice versa. It cannot exist by itself, and all the other components described in the summation must exist with relative amplitudes as determined by the nature of the circuit. It should be pointed out that the summation is taken to represent an electromagnetic signal carrying energy in the same direction as the beam — say to the right — but the summation contains terms with both positive and negative phase velocities or propagation constants and, therefore, the field it describes includes components which propagate (phase) both to the left or to the right or, stated differently, in the same direction or in the opposite direction as energy flow (and the electron flow).

Obviously, the interacting component must be one in which the phase of the component propagates at the same velocity and in the same direction as the electron beam.* To be somewhat more specific about the nature of the circuits used we list two of the most common circuits used in traveling wave devices.

1) The helix is most appropriate as an interaction circuit for beam voltages up to about 10 kv. Tubes using helices are suitable for powers from a few milliwatts up to perhaps several kw cw and perhaps 10 to 20 kw pulse power. 2) Coupled cavity

*It may be mentioned usefully here that a wave on the circuit which is carrying its energy in the opposite direction to the electron beam flow will also have components described by a similar summation with propagation constants either in the same or the opposite direction to the energy flow. This will mean that there might be components whose phase velocity is to the right (the same direction as the electron flow), even though the energy flow (group velocity) of the total wave is to the left. Any interaction will be of a drastically different kind than for the case of the beam circuit power and electron flow in the same direction. This can lead to so-called backward wave oscillators or amplifiers, rather than traveling wave tubes.

circuits which consist of a series of adjacent cavities which are coupled to each other through common apertures in a common wall. These are more thermally rugged than the helix and for other reasons having to do with their electromagnetic properties are suited for use at voltages from about 20 to 30 kv up to 100 or 200 kv. It should be mentioned that the generic name of coupled cavity tube is used for a particular structure, which typically is operated up to about 50 or 60 kv and peak power outputs of perhaps 100 kw. For higher voltages the common coupled cavity configuration used is something known as the cloverleaf, and is suitable for power outputs of several megawatts peak.

We return to the problem of interaction between the beam and circuit. As stated, the beam velocity should be synchronous with the velocity of one of the space harmonics indicated in the equation. The fact that the field in any periodic device can be written in the form shown above, including a summation over an infinite set of terms is sometimes known as Floquet's Theorem, and is an application of a general theorem about the description of any physical property in a periodic structure. In this case, it is the electric field which can be written in this form quite generally. It is a common practise to describe graphically such a set of space harmonics by drawing what is known as a Brillouin diagram. This merely shows the existence and the relation between the various space harmonic propagation constants at various frequencies.

Thus, if the periodicity of the structure is L , then the Brillouin diagram would be something as shown in Fig. 3 & 4.

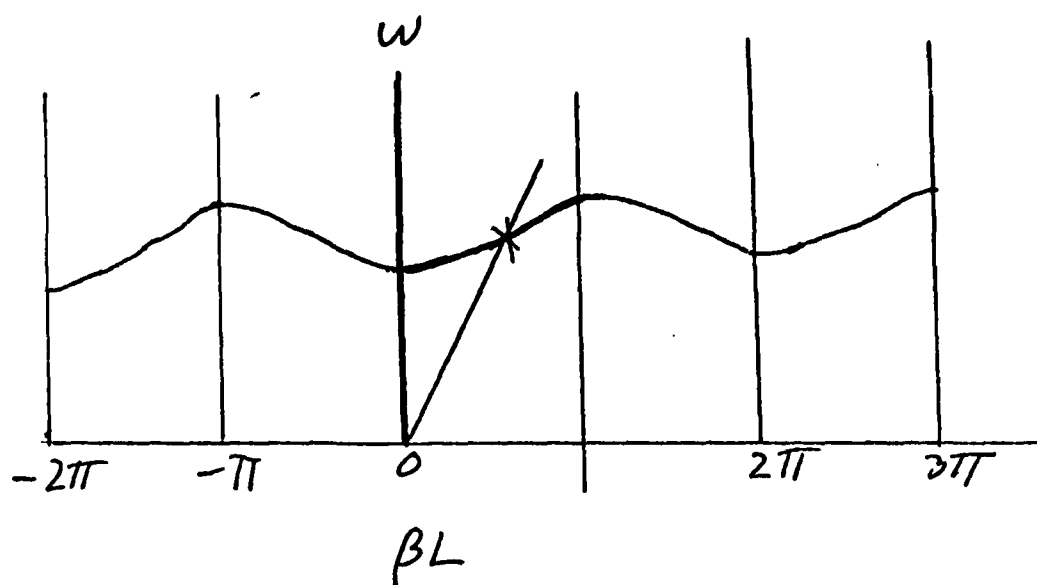


Fig. 3 Brillouin diagram for cloverleaf circuit.

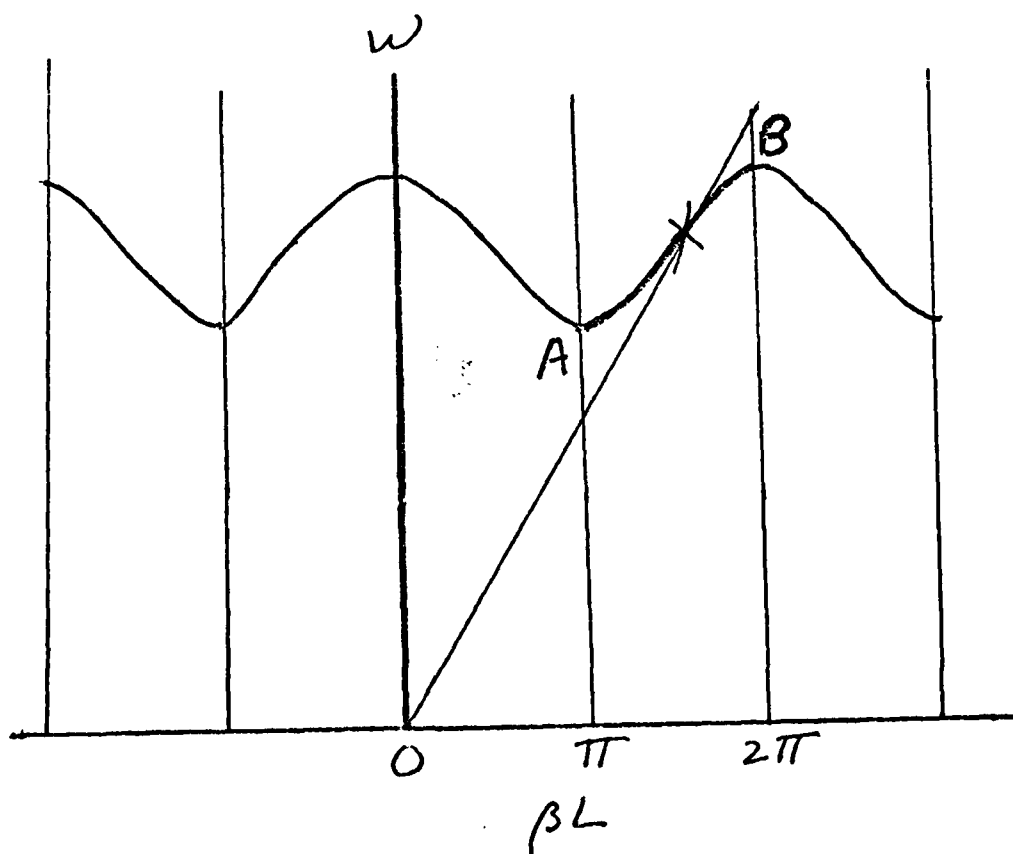


Fig. 4 Brillouin diagram for coupled cavity circuit.

where the propagation constants are plotted as a function of frequency. The only thing which is prescribed by Floquet's theorem is the spacing along the axis between adjacent space harmonics for a given frequency. This spacing is always $2\pi/L$ i.e., $\beta_{n+1} = \beta_n + 2\pi/L$. The shape of the curve, showing β as a function of frequency depends completely on the nature of the circuit. It can have a variety of shapes, as indicated in the figure which shows typical examples. The first branch for which there can be TWT interaction (positive group and phase velocity) is shown as a heavy line.* All the other branches shown as light lines are just displacements of this curve, moved by distance of $2\pi/L$ and $4\pi/L$, etc. Certain properties of this kind of plot are of great importance. The phase velocity v_p for any space harmonic at a particular frequency is given by ω/β , the slope of the line from the origin to the point on the dispersion curve and the group velocity $d\omega/d\beta$ is given by the slope of the tangent to the curve at that point.

It can be seen that the slope of the various branches at a fixed frequency are all the same. The slope $d\omega/d\beta$ at that frequency, represents the group velocity. Since this is the velocity at which the energy is carried, all (and each) of the space harmonics together, which constitute the wave, must be transporting their individual contributions to the energy at the same velocity. Note that the phase velocity for the various components, which is given by ω/β is smaller as one goes to higher values of $\beta_n L$ which corresponds to slower phase velocities.

* Fig. 4 would be typical of the coupled cavity circuit, Fig. 3 the cloverleaf circuit.

The corresponding portions of the curve shown in Fig. 3, 4 which have negative group velocity, also consist of an infinite set of segments spaced by distances of $2\pi/L$ (along the β axis) and represent the same mode carrying energy to the left. It is apparent also that there are some of these space harmonics for which the phase velocity is to the right, and these correspond to what has previously been referred to as backward wave components (namely, where the group and phase velocities are in opposite directions).

Now in any traveling wave tube the condition for significant interaction is that the electrons have velocity close to the phase velocity of one of these space harmonics of the circuit. Under these circumstances, we can ignore the other space harmonics — not because they are not present, but because they do not interact in a cumulative fashion with the beam. For the relevant (synchronous) space harmonic, depending on the phase of the electrons along this (sinusoidal) component, some electrons will be accelerated, and others will be decelerated. Since the electrons are moving at about the same velocity as that particular component, there is a cumulative effect of this field, and one will get bunching of the electrons around zero points of the space harmonic wave. An accelerating half cycle, together with the decelerating half cycle ahead of it will both push electrons toward the zero field point between them and there will be bunches of higher density of electrons at these zeros of the (space harmonic) wave. 180 degrees away from these points, that is at a zero in the field preceding an accelerating

half cycle and following a decelerating half cycle, the electrons will be pushed away from that point, and there will be a deficiency of electrons here.

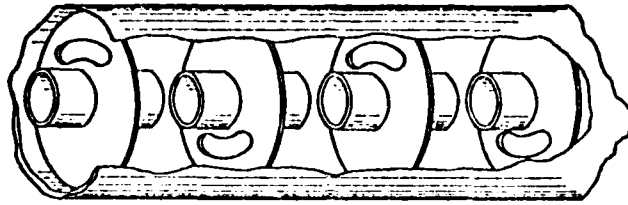
The periodic bunches of electrons (obviously occurring every cycle) can then transfer energy to the field if the beam velocity, instead of being exactly in synchronism, is slightly faster than the wave. The bunching process described above will still occur but, as the bunches form, they drift forward into a decelerating half cycle of the wave. Figure 4 indicates this qualitatively. By this process, one will get a transfer of energy from the beam to the electromagnetic wave, that is, to the circuit. This is the basic mechanism for traveling wave tube interaction and, to repeat, depends on an approximate synchronism between the electron velocity and a particular space harmonic (sinuosidal component) of the entire propagating field, which produces bunches and then because of the slight difference between the relative velocities, these bunches slowly drift into decelerating portions of the wave and lose energy.

This is the simplest description of the traveling wave tube interaction. It should be pointed out that/^{even though} the interaction is between a particular space harmonic because this is the only component which can have a cumulative interaction with the electron, as the as energy is transferred from the electrons to the wave by the process just described, the wave as it grows because of the added energy, has to distribute this added power in some specific ratio among all the space harmonics, since all of them have to be present to

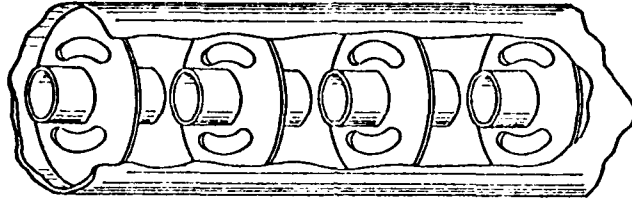
satisfy the conditions imposed on the field configuration by the metallic (or dielectric) structure which is propagating the entire wave.

Briefly, even though all the other components are present, they move either faster or slower than the electrons and have no cumulative interaction, but as the energy is transferred through the interaction with one component, this energy always has to be distributed to maintain the same relative proportion of all the components because that condition is imposed by the circuit or, more mathematically, by the boundary conditions of the metallic or dielectric structure. It should be stated here that under some conditions one cannot ^{completely} ignore the interaction with the other components.

Having discussed the generalities of TWTs, it is useful to discuss in somewhat greater detail the various common circuits used. The most common circuit used, at least in terms of numbers of tubes, is the helix, but this requires some special consideration and we will discuss it later. We first discuss the generic class of traveling wave tube circuits and their performance which are called generally coupled cavity TWTs. As the name suggests, and as was mentioned briefly in an earlier section, the coupled cavity TWT consists of a series of cavities which individually look something like klystron cavities, though the similarity is not always too great. Fig. (5) shows two versions — Fig. 5a is commonly referred to as the Hughes circuit or, sometimes, the Chodorow-Nalos circuit, since these were the

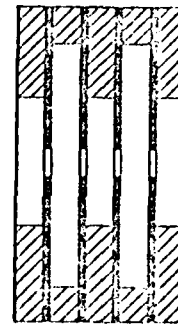
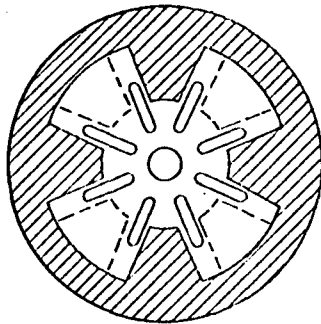


The "Hughes" structure (inductively coupled klystron cavities with staggered slots).



Chodorow-Nalos structure (inductively coupled klystron cavities).

Fig. 5a



The "clover-leaf" structure.

Fig. 5b

first people to demonstrate the usefulness of this kind of circuit.^{1,2} The tubes using this circuit, which have been built at a large variety of frequencies; voltages, and curriencies, typically operate in voltage ranges from about 20 kv to perhaps 60 or 70, and powers ranging from about 10-15 kw up to perhaps 100 kw. This circuit is the most commonly used for high power traveling waves at the ranges of powers mentioned here and it is probably true that most airborne radar uses this circuit.

For higher peak powers than listed above say on the order of a megawatt, one has to go to a somewhat different circuit Fig. 5b called the cloverleaf.³ Fig. 5b shows the outlines of two adjacent cavities. Alternate cavities are rotated 45° relative to their neighbors. The dotted lines then indicate the outlines of (say) the odd cavities, the cross-hatched region indicates the metallic exterior of the even cavities, and the slots in the common wall of odd and even cavities have the same relative position to the four sectional protrusions into each cavity. For this (cloverleaf) circuit, the cavities have no reentrant posts as in the standard coupled cavity circuit. This is a design technicality concerning with optimizing the interaction with electrons.

Let us consider the behavior of the standard coupled cavity circuit. The beam, of course, passes through the cavities along the axis and interacts with the electric field in each cavity in the gap shown very much as it would in a

¹ Chodorow & Nalos, 1956, Proc. IRE 44, p. 64.

² Chodorow & Nalos, Otsuka and Pantell, 1959, IRE Transactions on Electron Devices, Vol. ET-6, p. 48.

³ Chodorow and Craig, 1957, Proc. IRE, Vol. 45, Aug., p. 1106.

klystron with similiar cavities. The major difference in this circuit is that the coupling slots shown, provide electromagnetic coupling between adjacent cavities, so that any power that is injected into any one cavity is then transferred successively to the following cavities through the slot. The slot's dimensions and cavity dimensions determine the propagation properties of the circuit.

We can state here, without proof, that one can describe in detail the electromagnetic properties of such a string of coupled cavities by representing each cavity by an equivalent lumped constant resonant circuit, with the capacity being that across the gap and the rest of the cavity constituting an inductance. The coupling slots are also represented by a resonant circuit. (In this case, the resonance corresponds to a frequency for which the slot is approximately a half wavelength long.)

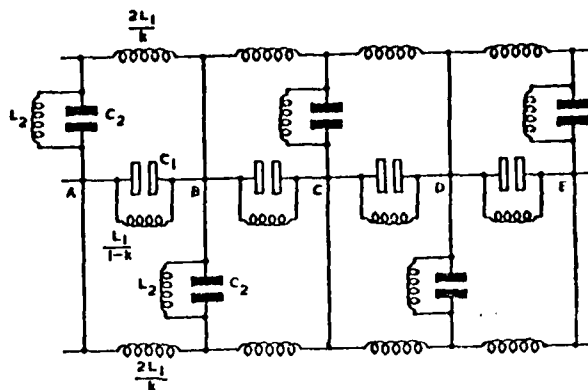
If one puts all of these statements together and tries to describe the structure in terms of lumped constant elements, one can show that it can be represented by the top circuit of the diagram (shown in Fig. 6i). C_1 , L_1/k and $L_1/1-k$ represent the cavity capacity and various parallel portions of the cavity inductances, and C_2 , L_2 represent the circuit behavior of the slot. By geometric transformations which do not change the circuit character, one can go through a series of steps shown in the figure which gets to a final equivalent circuit, shown in the bottom line, Fig. 6iv. It is not intended to go through the derivations or justifications of the use of this

circuit. It is sufficient to say that the representation by an equivalent circuit of this kind works very well and by suitable measurements on passive cavities one can find the parameters which appear in this circuit and use these to calculate the operating properties of the TWT with considerable accuracy.

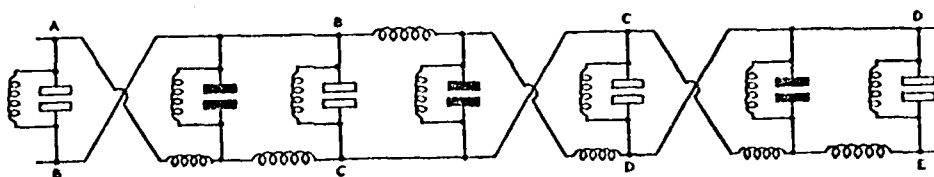
We consider now the interaction of the electron beam, it is apparent either from looking at the cavity structure itself or its equivalent circuit that it is a periodic array of electromagnetic elements with particular phase shift characteristics as a function of frequency. If one injects an electron beam with a voltage (velocity), such that as it went from one gap to the next it encounters the same phase of the field at each gap there would be cumulative interaction. It is also apparent that there would be a series of voltages for which this could occur since aside from some maximum velocity of electrons (and corresponding voltage), so that the electrons just stay in step with the phase shift, there are also lower voltages for which the electrons take extra cycles, one, two, etc., to go between gaps and, again, one would have a condition for proper phasing between electron motion and the field of successive gaps. It is obvious that there is an infinite set of such voltages each of which would permit the electrons to see the same phase of voltage at successive gaps.

It is not to be implied that these are all equally useful since there are other conditions for useful interaction, e.g., the electron has to get across the gap in something less than

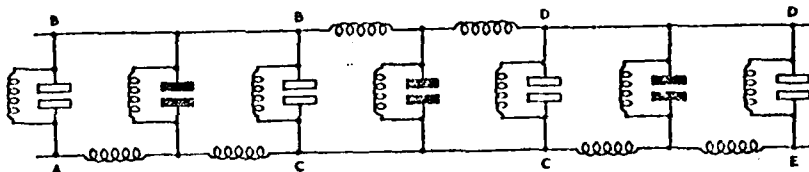
(i)
CIRCUIT IN
MICROWAVE
FORM



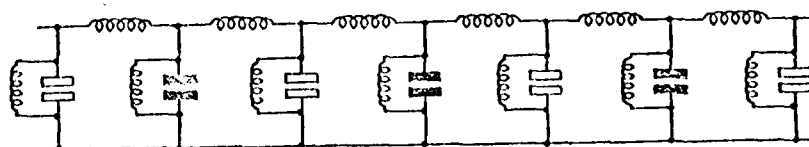
(ii)
FIRST STEP



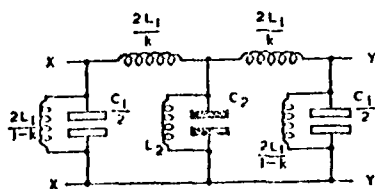
(iii)
SECOND STEP
NB. CONDENSERS
AB, CD, ETC. HAVE
BEEN TURNED OVER



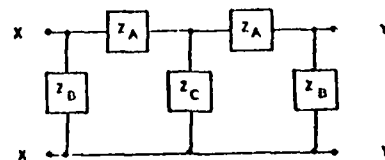
(iv)
FINAL FORM



(v)
BASIC CELL



(vi)



Equivalent lumped circuits for the coupled-cavity structure shown in

Fig. 6

a cycle to see some net effect of the modulating voltage across the gap. Therefore, even though one can get an infinite variety of synchronism conditions as described above, most of these are not useful because they would correspond to the transit time across the gap being many cycles and, therefore, the average effect of the voltage from the gap would be very, very little, essentially due to any residual part of the transit time across the gap beyond an integral number of cycles. All of what has been said can be said in an equivalent manner by going back to the Brillouin diagram for this kind of circuit which is shown in Fig. 4. It is to be recalled that at any point in this diagram the value of β/ω gives the velocity of a particular space harmonic and if that velocity corresponds with the electron velocity, one gets synchronism. This is the alternative way of describing what was stated just above. Again, however, the higher space harmonics (larger values of β) are not useful for the reasons just stated.*

The tubes using this circuit are almost always operated at the lowest value of β which corresponds to positive phase and group velocity (shown in the diagram with an x). It should be noted that, in the diagram which is suitable for this circuit, the smallest value of positive β for a given frequency corresponds to negative group velocity, that is, the direction of energy flow for this particular space harmonic (and its associated other space harmonics) corresponds to the energy being propagated in the direction opposite the electron flow and, therefore, it is not suitable for interaction with an electron beam for traveling

*In the context of space harmonics one would find that these have very small amplitude.

wave tube use. One could use this for a backwave oscillator, though this particular circuit is not commonly used this way.

Digression. Comments on efficiency, oscillation, feedback.

We have given perhaps more detail than is useful for the reader who is not designing such devices, but there is a significant physical consequence about backward and forward space harmonics, etc., and that is in order for this circuit to be used in the manner just described, the transit angle between gap centers must be approximate $3\pi/2$ (3/4 of a cycle). If one did not make these cavities reentrant as shown and assume that these dividing walls were very thin, this would make the transit angle across each gap also $3\pi/2$ (3/4 of a cycle), which would give relatively inefficient interaction with the field in the gap. (at 2π transit angle across a gap, the effect of the gap is zero.) It is therefore necessary to have the reentrant form for the cavity as shown in one of the previous diagrams, so that the actual transit time across the gap corresponds to something of the order of a $3/8$ of a cycle $(3\pi)/4$.

It can be shown that this division of the total electron path between gap and drift tube, is not the most efficient one in terms of optimum interaction with the electron but, if one is going to use this circuit, this is necessary.

The requirement for synchronism as the electron goes from gap to gap determines spacing between gaps for this kind of circuit. On the other hand, the requirement for optimum interaction within a gap determines the gap width. These two are different then, therefore, one gets the choice of dimensions which have just been described, and the price one pays is some loss of efficiency as compared to having the electrons see a field most of the time that they are passing along the circuit. This implies transit time between gap centers of $3\pi/4$ which is not possible for the coupled cavity circuit we have been discussing. It is possible with the cloverleaf which has a quite different geometry, but has a much narrower bandwidth.

The reentrant structure effects the efficiency of the interaction, but there is a more serious matter, of instability. This is related to the fact that if one is going to operate in the range of frequencies roughly lying between A and B on the Brillouin diagram, and let us say that one is going to use something like 50-75% of that region, one cannot avoid having the velocity of the electrons such that it is close to synchronism with a particular space harmonic at those frequencies centered around .B. As long as one is interacting only with the

component between A and B, this is desirable. There are, however, components which correspond to backward waves, that is, waves which are carrying energy to the opposite to electron flow which are somewhat to the right of B on the diagram, and are also close to synchronism with the electron beam, and interaction with these can lead to a great deal of difficulty, feedback with gain, possible backward wave oscillation, and there is a problem of possible oscillation in this frequency region because of the interaction with the backward wave. (In this context we are always referring to components of a wave which are carrying energy in the opposite direction of the electrons.) These components can cause difficulty as has been stated, oscillation or, in some cases even with no oscillation energy can get back toward the beginning of the circuit where it can interact with the electrons in an inappropriate manner, i.e., providing fields with undesirable phase relations and which can cause fluctuations in gain and phase shift as a function of frequency, etc., all of which are problems which are common to this class of device and have always been troublesome.

There is a related problem which leads to somewhat the same difficulty. At frequencies close to the edge of the pass band at B the impedance matching between the propagating circuit of the tube and the external waveguide through which the power is going to be transmitted is poor. This is a consequence of a well known property of the characteristic impedance of filter circuits near pass band edges. This means that there may be unusually large reflections from the output end in this frequency range, and this can lead to similar problems to those just described above due to interaction of the electron beam with reflected components.

All of this has been discussed in perhaps excessive detail. These are somewhat complicated effects which may not be completely comprehended from the rather brief description given here. They are mentioned because they are an important problem, and the reader may come across this difficulty about backward waves near pass band edges, and it was thought worth discussing here even at the cost of digressing from the main line of description.

A few other comments can be made about this TWT. First, it lends itself to periodic focusing, since ferrites and iron can be incorporated in this geometry to provide a periodic magnetic field which is the same as the period of the structure and leads to a rather compact design. The device can be run at relatively high average power, since it is an all metal structure

and cooling can be readily provided close to the region through which the beam is passing. Average powers on the order of tens of kilowatts are not at all unusual. This is to be contrasted with the helix which will be described later in which the metallic structure with which the beam interacts has to be suspended in a vacuum with dielectric supports. Cooling is not easily applied and this limits the average power to much lower values than are mentioned here.

The bandwidth for such tubes is typically of the order of 10 to 15%*. It is not easy to discuss the reasons for limitation on bandwidth, but it is related to the fact that if one changes the design parameters so as to increase the bandwidth of the cold circuit, i.e., the frequency spread between points A and B, many of the problems which were mentioned earlier about possible oscillation at the band edges, excess backward wave interactions, etc., become accentuated. (This is not a complete description of the difficulties. That would take us beyond the kind of description we are trying to provide here.)

It should also be stated here that treating the electron interaction as only with that space harmonic which lies between A and B, and ignoring the effects of the fields of the other space harmonics which are present is not a completely accurate description, even though the electrons are synchronous only with that space harmonic. Such a simple theory (which is much

*By various special design techniques including frequency selective loss near band edges, etc., bandwidths of the order of $\pm 25\%$ around some center frequency can be achieved.

more appropriate for the helix tube) leaves out the fact that the electron does "see" the other space harmonics even though their effects are not cumulative in the same sense as the interacting space harmonic. These other fields cannot be ignored completely, in particular for a structure such as is shown here where the electron really interacts with the field only in the gaps, and the notion of talking only about one space harmonic is an idealized one (the relevant space harmonic which extends everywhere along the circuit and is traveling close to synchronism with the beam). There is a treatment which is somewhat similar to that of a klystron where one talks about the electrons interacting with the total field at the gap, drifting on to the next gap, and interacting with the total field, etc., which is more accurate, and also includes the backward waves automatically.

For small signals in the beam and many gaps, it would turn out that considering the electrons to interact only with one traveling space harmonic continuously would provide a correct answer. The klystron-like approach with interaction at each successive gap but using the fact that the cavities are coupled, gives a more complete treatment. It is really the best treatment for this kind of device, and does give much more realistic answers. It also lends itself to extension of large signal theory, which we have not discussed here at all.

Severed Circuits. We might add a few more parenthetical but important comments. In any traveling wave tube it is necessary to provide attenuation along the circuit to prevent oscillation

due to reflections from the output to the input, so that reflections of this kind are attenuated in going from the output to the input. It is difficult to get adequate attenuation for this kind of circuit by putting lossy coatings on interior surfaces, although such coatings exist. Typically, one does get the attenuation in a very drastic fashion. One interrupts the circuit completely at one or more intervals, so that the power flowing along the circuit is absorbed at the interruption commonly called the sever. The electron beam goes on from that section of the tube into the next section. There the excitation on the beam launches a new growing wave on the succeeding circuit, and this, in turn, is amplified by interaction with the beam. If the lengths of successive separate sections of the circuit are properly chosen, etc., one can get overall gain for the whole length, even though one is throwing away all the circuit power at the severs. This is equivalent to complete loss being put in at these points, but there is net gain for the whole device.

Severs do provide adequate attenuation between the output and input but complicate the structure. Furthermore, there is a price to pay. The last (output) section of the tube, which has to be driven at a large signal — for maximum output power, is usually restricted in length (and gain) to prevent oscillation — just ^{for} as the previous sections. However, this means that the excitation on the beam entering this section has to have such a large amplitude, almost at saturation that that section is operating under adverse conditions and cannot reach the power

levels (and efficiency) that would be possible if the injected current were not already in a seriously saturated condition. It has been shown that to get optimum efficiency this output section should have some minimum gain — just to avoid having a saturated beam at injection. Because of the oscillation problem, one cannot afford to have the required length of circuit.

Helix. The other most important circuit used in traveling wave tubes is the helix. It is ideally suited for relatively low powers. In operation ranging in voltages from several thousand volts up to perhaps 10,000, at average power levels of a few watts to several hundred watts of average and 10 kw or so pulsed, its performance is very good. It is a very wideband circuit, and has many applications for communications, for countermeasures, etc. Its disadvantages have to do in large part with dissipation. Since the circuit is entirely enclosed in a dielectric envelope, the problems of cooling such a circuit are difficult and one cannot run very high beam powers through it. There are other restrictions which are of a more sophisticated nature concerning the electromagnetic properties of the helix which make it unsatisfactory at very high voltages. These will be discussed in somewhat greater detail below after we have described the circuit.

The helix, as implied by the name, is a coiled, single wire (though more complicated sets of wires are possible) which is wound in a cylindrical spiral so that there is a region inside the helix through which the beam can be transmitted. Its propagating properties can be in large part understood by

considering it, the helix, as the center wire of a coaxial transmission line with the outer conductor at considerable distance away, and if this center wire instead of being straight is coiled so it has the final form of the helix, then the field lines which normally would go from the center conductor to the outer conductor now become very drastically distorted so there are now electric fields between adjacent turns along the axis of the helix due to the relative voltages on the various turns. One gets a very distorted field picture and the particular feature which is important is that there is axial field in the interior of the helix and, as stated, fields between adjacent turns or between neighboring turns. The field actually on the axis can be considered as the average field due to the voltage variation due to different instantaneous voltage of the various turns. See Fig. 7.

This is obviously a highly qualitative picture. One can calculate the actual field, but what has been stated is really true, i.e., that the distortion of the center conductor into a coil results in axial fields which are the ones which will be interacting with the electron beam injected on the axis. An equally important feature of this description is that one can calculate with reasonable accuracy the velocity of the wave along this coiled transmission line purely by assuming that there is a wave propagating along the wire at the velocity of light. One gets a propagation along the axis of the system which is much lower from purely geometrical considerations and, specifically, that the ratio of the axial velocity to the velocity of light is

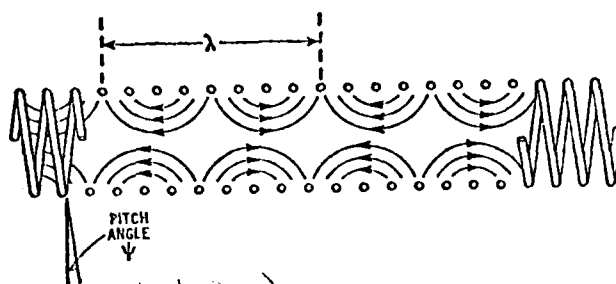
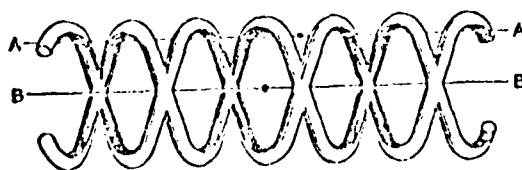


Fig. 7a



Twin cross-wound helix.

Fig. 7b

approximately equal to the ratio of the pitch of the helix (the spacing between turns) to the circumference.

This description does not strictly apply as one gets to lower frequencies when the velocity of the helix does go up to much higher values. Lower frequencies, of course, is a relative term which depends on some scale which is related to the operating wavelength and the diameter of the helix and the number of turns per wavelength. It is sufficient to say here that for most applications it is true that one can get an approximately constant velocity (by the geometrical consideration) over frequency ranges of 2 to 1 or better, and for dimensions which will give good interaction with the electron beam, for velocities not much higher than perhaps 0.2 the velocity of light (a 10 kv helix).

As stated earlier, the limitations on the use of this circuit derive in great part from the fact that it is a relatively fragile structure (a coiled wire) suspended in a vacuum where dielectric supports which are usually not very good thermal conductors and, therefore, the amount of dissipation permitted for a helix tube is much smaller than for a coupled cavity tube. The dissipation, of course, arises from interception of electron beams and r.f. losses due to circulating currents on the helix. With improvements in technology, the use of dielectric supports of BeO or BN it has been possible in recent years to extend the average power handling capabilities considerably over earlier performance and so that one can get powers of several hundreds of watts or even kilowatts, the higher powers being confined to lower frequencies which correspond to larger dimensions and

more adequate cooling. There is another and more sophisticated limitation on high peak power. To get high peak power one needs to go to higher voltages for operation, that is, to a faster helix. The helix becomes a more open structure, i.e., a larger pitch to circumference ratio to get such higher phase velocity. In any open structure, there are fields outside the structure and for higher velocities the fields extend further out and much more of the total energy is stored outside the helix. These external fields add to the power flow so that there is a greater power flow for a given field on the axis. Therefore, the interaction effectiveness (the interaction impedance) is reduced as one goes to higher beam voltages and helices designed for such voltages.

There are some high voltage alternatives to the simple helix, for example, the cross-wound helix*, which consists of two helices wound in opposite directions (sometimes also called the ring-bar circuit). In this circuit, for somewhat complex reasons, the deterioration of performance because of excess electric energy stored outside the helix is not as severe. Such structures can go to higher voltages with adequate interaction, say 20 to 30 kv, and can run at peak power levels, for example, on the order of 50 kw, but the average power is still limited because of thermal dissipation. There are applications, however, where pulse power of the kind just listed and relatively low average power are adequate and cross-wound helices are used for this purpose.

In general, however, the helix is the dominant circuit for power levels of a few kw pulsed, and/or 1 or 2 kw or lower average because of its fine bandwidth, very good interaction

*Chodorow & Chu, J.A.P. 26, pp.33-43, 1955.

impedance, simplicity of construction. Appropriate attenuation also can be applied relatively easily (to prevent oscillation), this attenuation usually consists of carburizing the dielectric support rods, and this can be applied in such a way that attenuation and oscillation problems encountered with coupled cavity tubes do not occur.

Nothing has been said so far about a space harmonic description of the helix. A helix, like any other periodic structure, does have space harmonics, but these have some peculiar characteristics. A helix has certain symmetries beyond merely periodic translational symmetry. In particular, there is a relation between its angular symmetry and its translational symmetry. In translational symmetry if a structure is periodic with some period L , two points displaced by an axial distance nL are not distinguishable where n is an integer. In a helix — two points displaced by some fraction f of its period p (the pitch) and displaced in angle by $f (2\pi)$ are indistinguishable even for non integer f . This has important consequences for its space harmonics — basically it connects particular values of axial propagation constant ($\beta_n = \frac{2\pi n}{p} + \beta_0$) with particular angular variations of corresponding fields. It turns out that there is only one important space harmonic with electric field on the axis (β_0). This is the harmonic which interacts with the beam in a normal traveling wave tube. There are other space harmonics which have the characteristic that their fields are concentrated very close to the helix itself (usually a tape or a wire), with zero field on the axis and in which the field amplitude varies with angle around the helix. These space harmonics are noteworthy for two

reasons. First, if one has a beam which is so formed that a considerable portion of the current passes close to the helix (or a hollow beam), then one of the space harmonics, which is a backward wave, (that is, a forward phase velocity and backward group velocity) can be used very effectively in a backward wave oscillator and was so used for many years for very wideband oscillators where one could get tuning by changes in the voltage of the electron beam. This is no longer as commonly used because of the existence of alternative sources.

Another characteristic of this backward wave space harmonic is that for a certain frequency range it will be synchronous with the beam at twice the frequency of the signal on the beam and can lead to harmful interaction with the beam at this higher frequency. This can be troublesome. To confuse things further, even the fundamental space harmonic with field on the axis which is used for ordinary amplification is also in approximate synchronism with the second harmonic current in the beam. This is the concomitant price of having a very wideband circuit. This coupling of the second harmonic of the beam can be troublesome. It is not possible here to say a great deal about these interactions, but merely to point out that they exist and can be a complication in certain ranges of performance.

In its range of applicability, however, the helix is a very useful circuit. It is the tube used as the satellite transmitter in satellite communications. It is widely used for any application requiring relatively low power and wideband. It is also used for countermeasures, etc. In the past, helix tubes of special design were very good low noise amplifiers used for receivers but, to

some extent, this function has been displaced by the development of very low noise solid state amplifiers in the microwave range.

V. Crossed Field Devices (Magnetron)

There is a whole class of devices which are related in some ways to traveling wave tubes, but have a fundamental attribute which makes them quite different. In general, they tend to be much more efficient as oscillators, but as amplifiers they have limitations, do not operate very well at small signals, can easily break into oscillation, are subject to a great deal of noise, etc. But within the restrictions of their characteristics, crossed field devices, because of their great efficiency and compactness, can play a vital function in certain applications.

There is a variety of such devices with different geometries, cylindrical, planar, etc., oscillators, amplifiers, their general characteristics can be described somewhat as follows. The electrons interact with a traveling electric field, which is maintained by a traveling wave circuit, just as in the traveling wave tube. The circuit velocity is determined by the geometry of the circuit. Just as in the traveling wave tube, there are gaps in the boundaries which have electric fields, and adjacent gaps are coupled to each other by various electromagnetic means which may be different in detail than in a traveling wave tube, but serve the same function. As seen from an electron passing the gaps, the environment looks very much as it would look in a traveling wave tube, including the fact that the fields decrease in amplitude as one moves ^{away} from the metal wall.

The major difference between these devices and the traveling wave tube devices, is that the electrons move in a region in which

there are constant (in time) electric and magnetic fields perpendicular to each other — crossed fields. It is known that in such a crossed field, the average motion of the electrons consists of a drift perpendicular to both fields at a velocity (in certain units) equal to the ratio E_0/B_0 . This is the average drift velocity. Depending on how the electrons enter this region, at what angles, velocities, etc., there is usually superimposed on this drift motion a cyclotron motion corresponding to the magnetic field, so that the general motion is usually described as a cycloid, which is shown in Fig. 8, which is really a superposition of a drift motion, plus motion in a circle moving at the drift velocity. All crossed field devices have this characteristic and the velocity (determined by E_0/B_0), is selected so that it is synchronous with the circuit velocity, i.e., with the velocity of the electromagnetic wave along the surface. The typical electron motion then is this drift with various oscillations around this average velocity, due to the superimposed cyclotron motion. One would like to arrange it in some way that the superimposed cyclotron motion is small in amplitude. This is a function of how the beam is injected into the field. It is not always achieved, and really does not effect behavior too greatly.

We will start by describing the simplest crossed field tube. It is not the most common but everything that can be said about it can be said about the other kinds of devices, although the other devices will have complexities superimposed by their mode of operation. In the tube shown in Fig. 9

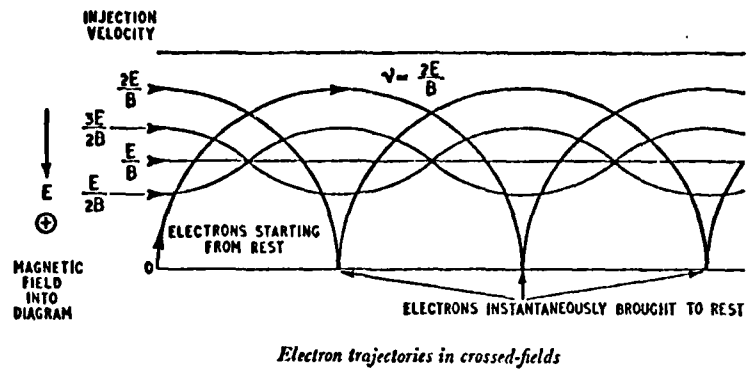


Fig. 8

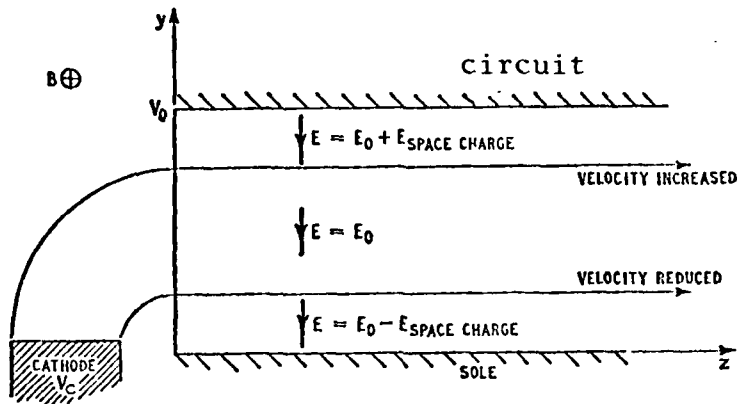


Fig. 9 Injected beam crossed field amplifier

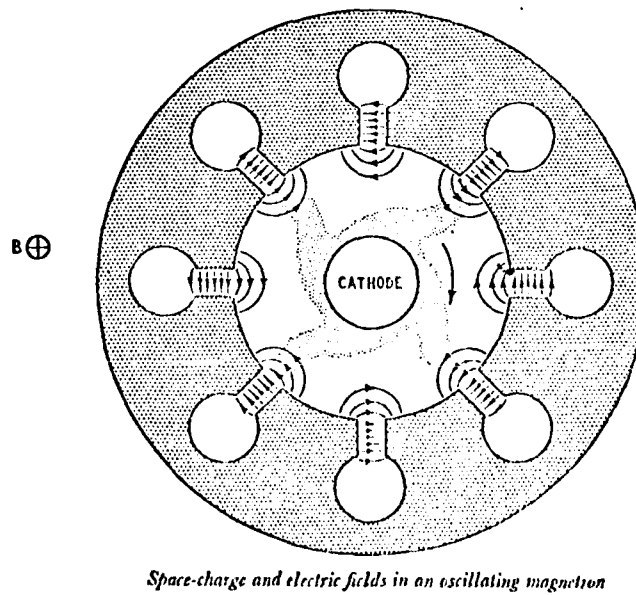


Fig. 10

one has a cathode external to the crossed field region, and one injects the electrons from this cathode (at zero potential), into the interaction region between the propagating circuit (the anode) at positive voltage and another electrode (the sole) which is at zero voltage or negative. The injection geometry is designed to produce flow parallel to the anode (circuit) along an equipotential (V) such that its kinetic energy just corresponds to the drift velocity $u = E/B$ ($\frac{m_w^2}{2} = eV$). This drift velocity then is chosen so it is synchronous with the waves on the circuit (at anode potential), and one will get interaction with the circuit.

The subsequent behavior, however, under the action of these r.f. fields is different than in an O-type TWT. As the electrons lose energy or gain energy from the field, their velocity no longer satisfies the E_0/B_0 condition, the electrostatic and magnetic forces no longer balance and the consequence is that the electrons will tend to drift either closer to or further from the anode. Actually, those electrons which have lost energy to the circuit fields will drift toward the anode, those which have gained energy, away from the anode. In this ^{moreover} tube, accompanying the circuit electric fields parallel to the average direction of motion of the electrons (perpendicular to E_0) there are also circuit fields perpendicular to this motion (parallel to E_0). (This statement about fields transverse to main drift, also applies in the gap of a klystron or the gap of an O type TWT, but the transverse fields in those devices have very little effect.)

In the crossed field devices, however, ^{the} steady (as seen by the moving beam) transverse circuit fields, in combination with the constant magnetic field produce longitudinal drift motion (perpendicular to E_0), and parallel to u_0 — the average drift velocity. The direction of this drift depends on the sign of the transverse circuit field and one gets bunching in the longitudinal direction (parallel to u_0 the average drift velocity). The phasing is such that these bunches tend to be concentrated in the region of maximum retarding longitudinal circuit fields. Since this appears as an approximate constant field to the synchronous electrons this, together with B_0 causes a drift toward the anode.

Without going into a more detailed description of all the fields and forces with signs and directions, one can say that from this combination of static electric fields, static magnetic fields, and synchronous circuit fields, with both transverse and longitudinal components, one gets bunching and a drift of the bunches toward the anode, picking up energy from the dc field, but their phase position relative to the circuit is such that they are transferring energy to the traveling electric field (the circuit). The electrons which are unfavorably placed initially, tend to drift away from the anode and eventually arrive at the sole.

Of all the crossed field devices, this injected beam amplifier is the closest in its overall behavior to an O type TWT. It is important to remember, however, that the dynamics of the electron motion is quite different because of the presence of both a constant magnetic field and a constant

electric field, with their consequent effect on electron motion. The principal similarity to an O type TWT arises from the fact that one has an injected electron beam which moves parallel to the circuit, and eventually can be collected by a collector after the interaction of the circuit has been completed.

All the other types of crossed field devices to be discussed below have cylindrical geometry with a central cylindrical cathode and the anode (circuit) as an outer cylinder with the circuit facing the cathode. The electrons then are accelerated from the cathode, go through a complex trajectory involving superposition of the drift motion and cyclotron motion, and many of them then end up with a motion around the cathode parallel to the anode surface, with a drift velocity equal to E_0/B_0 , and the interaction is very much as we have described it for the linear device. The rejected electrons, the ones which do not reach a favorable phase, presumably return to the cathode. But, since the circuit is closed, that is, the circuit is arranged around the circumference of the circle and the electrons move with average paths which go around the cathode, there are built in feedback mechanisms via the electrons which have to be considered as drastically effecting the behavior of the device. More about this will be said below.

In the linear device just described, the absence of this feedback makes it possible to treat this as an ordinary amplifier, with small signal performance, due to a small injected signal on the circuit, etc. Many of the things said about an O type TWT are possible here, also. One can use a severed circuit, where the circuit consists of isolated sections with a common beam

passing from one section to the next, just as in an O type device. Backward wave oscillators are also possible as in O type. By and large, these linear crossed field devices are not as efficient as the circular versions, but they can be operated as true amplifiers. One can use a severed circuit, as in an O type device, with relatively large gains, etc.

These devices, as with all crossed field devices, tend to be quite noisy, in particular, there is an effect which is most easily detected in these injected beam devices known as Diocotron effect, which is an instability of an electron beam in crossed electric and magnetic fields. It can be shown theoretically that given a beam moving through crossed fields, with the current density, velocity and other parameters appropriate to these fields, so that one has (at least in principle) an equilibrium electron flow, then any perturbation of such beam will grow in amplitude and will result in large distortions in shape and velocity from the originally assumed parallel flow. Experiments have shown that such instabilities will exist in a beam with no circuit. Under the best of circumstances, such instabilities will lead to excessive noise being generated because the beam is no longer a dc steady beam. It is known that output such crossed field amplifiers, both linear and cylindrical, have excessive noise. Also, it is known that, under certain conditions, presumably because of such instabilities, one can get electron velocities quite different from what would be permitted by the dc conditions. There are large fluctuations and electrons with excess velocities much beyond those appropriate

to the applied voltages will actually strike the sole with high energy. It has been shown that one can get very large current to the sole even when the sole is negative with respect to the cathode. At best, such devices used as amplifiers will have excessive noise near the carrier.

It should also be pointed out that one can make such injected beam devices in a cylindrical format where the anode (the circuit) and the sole are bent into the form of coaxial cylinders, with an electron beam being injected into the region between them just as in the linear case. This merely makes the device more compact in terms of applied magnetic fields, circuits, etc., but it is subject to all the problems listed.

In cylindrical geometry crossed field devices which are the most common, one has ^acentral cylinder as the cathode (see Fig. 10) an anode concentric with the cathode with some suitable spacing between them. The magnetic field is applied parallel to the surfaces of the concentric cylinders, and the appropriate dc voltages applied between the cathode and anode. The electrons coming off the cathode then initially will undergo a somewhat complex motion, being accelerated by the dc electric field, and deflected by the magnetic field, under conditions of no circuit fields. In principle, all electrons should return to the cathode. Under actual conditions, with r.f. fields on the anode (the circuit), space charge effects, initial emission velocities, etc., one gets a certain fraction of the electrons with the characteristic E_0/B_0 drift velocity around the cathode and parallel to the electromagnetic wave

velocity around the circuit. The anode, of course, does contain the circuit, which typically will consist of a series of radial slots suitably spaced.

The most common device is the simple magnetron which is an oscillator, has coupling between the slots, both through their electric fields at the surface of the anode, and through magnetic fields at the back of the slot threading through the corresponding regions in adjacent slots (see Fig. 10). This is entirely equivalent to having a series of coils lying next to each other in a common plane so that one would get mutual inductive coupling between such coils. Since the whole circuit is closed upon itself there are conditions on what phase shifts between slots are possible. It can be easily seen that the appropriate condition is that the total phase shift around the whole circuit must be an integer number of 2π . It will be seen that one can have phase shifts of π/S where $S = M/2k$ where M is the number of slots and k is an integer. Typically, the magnetron oscillator is run with $S = 1$ corresponding to π phase shift between adjacent slots. This corresponds to a particular frequency and in terms of space harmonics which is appropriate with this kind of circuit as with an O type device, one chooses a frequency (and corresponding phase shift) so the circumferential drift velocity of the electrons (v_0/B_0) is synchronous with that of a space harmonic.

Though commonly, magnetron oscillators are run at a frequency corresponding to π phase shift per slot, it is sometimes possible for mode jumping to occur where the magnetron, when it is turned on, does not take off on the proper mode or other interfering modes can cause some problems. More serious is the fact that

as a pulsed oscillator, depending on initial conditions, initial current flow, etc., the oscillating frequency may not be completely repetitive, and there is some jitter in oscillation frequency. Many of these effects, for example, can depend on how abruptly the voltage is turned on. It should be stated that the whole question of starting conditions, and the exact nature of the electron flow under oscillation conditions, is still a very controversial one. Even the problem of what sort of flow exists under purely dc conditions, with no circuit, is not a simple one.

One can get solutions under dc conditions of the electron flow in the magnetron by means of a computer, but even this is not easy. In an oscillating tube, one must take into account secondary emission from the cathode which will occur since some electrons which initially gain energy from the circuit, are driven back to the cathode with some finite energy and produce secondaries.

Cylindrical magnetrons, the most common sources, however, are very efficient. There is also^a version, something called the coaxial magnetron, in which the interaction circuit is locked to an external high Q cavity, which locks the interacting mode in such a way that the spectrum one gets from an oscillator of this kind is quite clean, and the frequency does not jitter. The magnetron oscillator, in general, can run at relatively high peak powers and efficiency. There is also a class of amplifiers using the cylindrical geometry in which one has both an input and an output, and the closed (reentrant)

circuit is somehow interrupted so that there is direct electromagnetic coupling between the input slot and the output slot only along one angular segment. The power then, fed into one slot, will circulate around the anode and will be amplified by interaction with the beam. The circuit interaction may be a forward or a backward wave, in either case the phase velocity of the wave must be synchronous with the circumferential electron velocity, but the power flow may be in the same direction or in the opposite direction to the phase (and beam) velocity. Input and output are defined relative to direction of power flow. It is also possible, in addition to severing the circuit, to actually remove a few of the slots, so there is actually an angular drift space between the output slot and the input slot and the electrons' circulation/^{path}around the cathode includes this drift space. There may even be an extra control electrode in this drift space to affect the angular electron motion. This presumably reduces the amount of feedback due to the electrons, and permits a somewhat different kind of operation. Typically, a device using a backward wave circuit, with no drift space is called the amplitron. Those devices which have a drift space are usually called crossed field amplifiers, and one must further designate whether they are using a forward or a backward wave for the interaction. These various devices all differ somewhat in their operating characteristics but, typically, they have something on the order of 10 to 12 db of gain, although a forward wave device with a drift space may have slightly more, - 13 - 15 db.

If a backward wave is used, the bandwidth typically is of the order of 7 or 8%, for a forward wave device, it may be 10 or 12%. The gains quoted are with relatively large input signal so that the tube, although driven as an amplifier, is really running saturated. With lower drives, one can get more gain, perhaps 20 db in some types of devices, but such tubes become extremely noisy if operated with small signal input, and with no input they generate a great deal of noise. This presumably is related to the Diocotron instability and possibly to some electronic feedback. In terms of noise, a typical amplifier of this kind, even at large signal, will have noise about 30 db down from the carrier at frequencies 1 MHz away from the carrier. This can be somewhat better, particularly in a forward wave device with a few slots missing (a drift space), and the noise can be reduced to 40 or 50 db. Efficiencies for these devices is in the range of 50 to 75%.

In the corresponding magnetron oscillators, the efficiencies can be comparable. For the oscillators, typical power output can be as high as 5 megawatt peak at 3 GHz, with perhaps 4-5 kw of average power, and a megawatt at 1 GHz, with 1 kw of average power. Amplifiers typically run at lower peak powers. These numbers typify both the advantages and disadvantages of such devices. One can get quite high peak power ^a in/relatively compact device, but because the electrons finally land on the anode and dissipate their power on the anode, there are limitations to the duty cycle and, therefore, to the average power.

It has been possible, particularly at lower frequencies by optimum design, to go as high as perhaps 80% efficiency in a cw device, running at 25 kw, at 1 GHz, and comparable efficiencies have been reported at higher frequencies, but it is not standard performance, and requires liquid cooling of the anode (the circuit) to take care of dissipation due to the electron bombardment.

The major advantage of crossed field amplifiers is that they are much more compact for a given power level than the corresponding O type device. While one can get much higher average power from an O type device, it is usually more bulky because of magnetic focusing required, separate collector geometry, etc. On the other hand, the O type device can have more gain, more bandwidth, as well as the higher average power. Also, as a true amplifier it can handle more complex signals, and has a better spectral response.

VI. Fast Wave Devices

In this section we would like to consider a class of devices which are quite different in their mode of operation than the devices we have discussed previously, namely, the klystron, the traveling wave tube, crossed field amplifiers, and/or oscillators. The particular virtue of these fast wave devices is that, because of their mode of interaction, they can attain frequencies and powers not possible with the kinds of tubes we have discussed up to now which we shall refer to as slow wave devices.

The characteristic of the slow wave tubes is that always there is some component of the electromagnetic field which is traveling at about the same velocity as the electron beam, so that there is cumulative interaction between the beam and the field. This applies both to the field modulating the motion of the electrons to get bunching, but also after the velocity modulation and bunching produce r.f. currents in the beam, the r.f. current being carried along at the beam velocity can also interact cumulatively with the electromagnetic field, and transfer power to the electromagnetic field.

In the traveling wave tube case, one could analyze the total field as a set of so-called space harmonics, and the requirement for interaction with the electron beam was that one of these space harmonics traveled at the same velocity as the electron beam. The same kind of analysis would apply to a crossed field device. For the klystron, a mathematical analysis would indicate that the interaction in a single gap can also be

written in a similar way, except that one uses an integral over a continuous set of space harmonics with various velocities rather than a summation over a discrete set and it is only the one term in the integral, synchronous with the electron velocity, which provides the interaction between the beam and cavity, either for velocity modulation of the beam or to induce r.f. voltage in the cavity for energy transfer. The nomenclature of slow waves refers to the fact that this component of the field which is synchronous with and interacts with the electrons must obviously have a phase velocity less than the velocity of light, hence "slow wave."

The basic difficulty in going to higher frequencies with "slow wave" tubes arises from a physical constraint on the characteristics of electromagnetic waves imposed by Maxwell's equations. Basically, any field component which has a slow axial velocity (so that it can be synchronous with an electron beam traveling at less than the velocity of light) will also have an amplitude which decays very rapidly as one moves from the boundaries toward the axis of the device. Stated differently and perhaps more usefully, if one has some slotted metallic structure surrounding the electron, such as exist in magnetrons, traveling wave tubes, or the resonant cavities of klystrons, an imposed voltage at the slot results in a set of traveling or standing waves in the region of the beam, and the field amplitude of that component which interacts with the beam decays very rapidly as one moves away from the metallic boundary (location of the slots). This means that the effective inter-

action is reduced, and that the useful portion of the beam is that close to the wall, that is, close to the slots in the wall which produce the interacting fields in the region traversed by the beam. However, all of these field variations are measured on a scale proportional to the wavelength, and the velocity of the wave (and electrons).^{*} A close distance will be measured in units scaled to the wavelength. Thus, a one centimeter diameter cylindrical beam in a device operating at 3 GHz, with adequate clearance, say .25 cm from the walls and strong interaction, will become at 10 GHz a 3.3 millimeter beam with a clearance of 0.8 mm from the wall and a 0.33 millimeter beam, with the same relative clearance .08 mm at 100 GHz.

These are all rigorous arguments based on scaling and put severe restrictions on the size of the aperture for which one can still get useful interaction. Essentially this means that for a 25 Kv beam, the aperture diameter cannot be appreciably larger than 0.10 of the wavelength or .065 at 10 Kv, etc., in each case the allowed size of the aperture being directly proportional to the velocity of the electron. If one is going to maintain some adequate beam clearance, this always

^{*} To state this numerically — the field amplitude in a cylindrical region will vary as

$$I_0 \left(\frac{2\pi r c}{\lambda v_0} \right) = I_0 \left(\frac{\omega r}{v_0} \right) \quad \text{where } I_0(x) \text{ is a modified Bessel}$$

function and where ω is the frequency, λ the wavelength, r the radial position, and v_0 the velocity of the active field component (and the electron). I_0 varies from 1 at $x = 0$ to 5 at $x = 3$.

For planar (rather than cylindrical geometry) the transverse variation is proportional to $\cosh(x)$ and in both cases (radial or planar) the field varies approximately exponentially for x greater than about 3.

results in intolerable beam focusing requirements as one goes to very high frequencies, and also the problem of very fragile structures, so that one soon runs out of the possibility of scaling to very high frequencies or ends up with mechanical marvels or monstrosities depending on one's prejudices. In any case, it is always very difficult and very expensive. With the current state of mechanical construction ^{one} /is limited to no more than perhaps 50 GHz.

The devices to be discussed here overcome this problem by permitting cumulative interaction between electrons and electromagnetic waves, even if the electromagnetic waves are not synchronous with the axial electron motion and, as a matter of fact, have much higher phase velocities. Hence the name, fast wave tubes.

The basic principle involved in all of these is that one introduces a periodic motion in the electron beam either by a set of periodic deflecting magnets or the rotational motion in a uniform axial magnetic field, and then one provides an electromagnetic structure which can have a propagating electromagnetic field such that there is synchronism between the periodicity of propagating electromagnetic wave and the periodic motion of the electrons. Stated more precisely, one tries to get a relation between the various quantities, the electron velocity, the electromagnetic wave velocity, the electromagnetic frequency and the frequency of the periodic electron motion so that, as the e.m waves slip past the electrons at their relative velocity, at corresponding portions of the electron transverse trajectory, the electron sees the same phase of the electric field in each cycle of its periodic motion.

In slow wave tubes, the electrons see an electromagnetic wave moving at the same velocity as the beam, so as far as the electron is concerned, it sees a more or less steady field. In fast wave devices, the electron sees a fluctuating electromagnetic field, but in each cycle of the electron's periodic motion it interacts with one cycle of the electromagnetic wave. There is some net effect whose magnitude and sign depend on the relative phase of the periodic motion and the electromagnetic wave. This net effect is repeated each cycle because of the synchronism imposed between the slip velocity, the periodic frequency and the electromagnetic frequency. Therefore, there will be a cumulative effect of this interaction even though it is not steady.

Some of these general statements will become clearer when we describe the two principal devices, the free electron laser (or ubitron as it was originally named when first discovered), and the gyrotron (or cyclotron laser, as it was originally named when first discovered). In both devices, there is a quantum mechanical name and a classical name. In one case, the classical name came first and the quantum mechanical name came later, and just the opposite for the other case, but in both cases the classical description is all that is necessary, and is more comprehensible. We will try to describe here the physical mechanism involved in the interaction, and some of the performance achieved or potentially possible. Again, we should stress the importance of these devices as they can dispense with the narrow transverse dimensions of metallic circuits around the

beam which is imposed in slow wave devices as described earlier in this section. In both of these devices, one can use normal uniform wave guides or large open resonators with suitably spaced mirrors around some region of space with free space propagation to produce a properly controlled and confined electromagnetic wave to interact with the electrons.

For both devices, and for either an open region with end mirrors or a waveguide (we shall use the term waveguide for either case), we introduce the notion of doppler shifted frequency. The electrons are moving axially with some velocity u_0 in a waveguide; an electromagnetic wave is propagating through the waveguide. The frequency experienced by the electrons will not be the same as the frequency of the electromagnetic wave. As the electron moves parallel to the electromagnetic beam motion, the number of cycles it experiences per second, the so-called doppler shifted frequency, will depend on the relative velocity of the electron beam and the electromagnetic wave. An electron moving in the same direction as the electromagnetic wave obviously sees fewer cycles per second than a stationary electron, and more cycles for an electron going in the opposite direction. If the electron is moving at the same velocity as the electromagnetic wave it sees a steady (constant) field even though to a stationary observer there is obviously an electromagnetic wave propagating along the waveguide or transmission system. This last case is typical of the conventional slow wave tubes.

The frequency experienced by the electron beam in general, then, is the "doppler frequency," it is given by the simple

formula

$$\omega_e = \omega \frac{(u_o \mp u_e)}{u_o} \quad \text{where } \omega \text{ is e.m. frequency and } u_o, u_e \text{ are the wave and electron velocity.}$$

Thus, the electron experiences the doppler shifted frequency due to the presence of the electromagnetic wave, and has a periodic motion of its own at some frequency determined by the axial magnetic field or the spatial periodicity of a deflecting structure and its own axial velocity. In either case, if the periodic motion of the electron is at the same frequency as the doppler frequency of the electromagnetic wave, there can be cumulative interaction. The motion of the electrons, because of this cumulative interaction can have a steady drift in both velocity and position. Different electrons will have different histories, depending on what the relative phasing is between their own periodic motion and the doppler shifted frequency due to the electromagnetic wave. The fact that different electrons have a different history means that their motion under the action of these two contributing effects will be different. One can get bunching and ultimately the bunches can transfer energy to the electromagnetic wave. This is a general and somewhat unspecific description but is really applicable to both devices to be discussed here.

The easiest device conceptionally possibly, is the gyrotron. Here the electrons moving axially rotate in a d.c. magnetic field at the cyclotron frequency. An electromagnetic wave passing through the same region with a synchronous (doppler) frequency will have a cumulative effect on the electrons. The circular orbits of the electron due to the d.c. magnetic field

will be such that, if in one-half cycle when the electron is moving down it experiences an electric field due to the electromagnetic field in one direction, on the second half cycle when it is moving up, the electromagnetic field has reversed sign and, therefore, the phase relation between field and motion is the same in both halves of the cyclotron orbit. So, if the electron is moving against the field on one-half of its orbit it is moving against the field on the other half, both field and direction of motion having reversed, and there is net energy transfer between field and electron. This does not give a net energy transfer mechanism averaged over the beam, because for every electron which loses energy to the field there are other electrons 180° out of phase for which the opposite is true — they take energy out of the electromagnetic field. Another mechanism is necessary beyond synchronism and this arises from the fact that the cyclotron frequency of the electron in a magnetic field is not a constant. Because of the relativistic mass, the cyclotron frequency is a function of the total energy of the electron. The cyclotron frequency ω_c of an electron is given by

$$\omega_c = \frac{e B_0}{m_0} \sqrt{1 - v^2/c^2}$$

where m_0 is the rest mass of the electron and v is the total velocity. Therefore, the cyclotron frequency increases as the energy (velocity) decreases.

It is this dependence of the cyclotron frequency on the relativistic mass of the electron which provides the means for getting power out of a gyrating beam in a magnetic field,

interacting with a transverse electromagnetic field. Those electrons whose initial phase relative to the field is such that they gain energy from the electromagnetic field will tend to have their cyclotron frequency decrease. Those electrons whose phasing is such that initially they lose energy to the electromagnetic field will have their relativistic mass decrease and their cyclotron frequency will increase. The frequency of the electromagnetic field applied to the electrons is such that it is somewhat higher than the relativistic cyclotron frequency of the injected electrons. Then, those electrons which initially have a phase to gain energy will shift to a lower cyclotron frequency and will drift out of that initial phase relation because the increasing difference between their cyclotron frequency and the electromagnetic frequency.

Correspondingly, those electrons which lose energy initially to the electromagnetic field have their cyclotron frequencies rise, so that they are closer to the electromagnetic frequency, and the rate of phase spreading between the electrons and the signal will decrease, since they are approaching each other in frequency. Therefore, these electrons will continue to stay closer to their initial phase and will continue to lose energy to the field. This is the mechanism which makes the gyrotron operate, and it is essentially a bunching phenomena in which the favorable electrons (which deliver energy) tend to stay in the same favorable phase, since their frequency of rotation tends to get closer and closer to the electromagnetic frequency, whereas the unfavorable electrons drift away from their initial phase relative to electromagnetic wave.

It should be remembered in talking about the cyclotron frequency here, that the frequency seen by the electron beam is the doppler shifted frequency which depends on the difference in the axial beam velocity and the phase velocity of the wave. If, as is very often the case, we are operating very close to the cutoff of the waveguide, where the phase velocity is close to infinity, frequency experienced by the electrons is the actual frequency of the wave. In any case, the important characteristic being pointed out here is that there is no slow wave structure required, and one can use a normal waveguide. As a matter of fact, for various reasons sometimes one uses higher order modes of an oversize waveguide so that the electron beam is actually quite far from the wall. A common field mode used in the waveguide is the so-called circular electric mode TE_{010} or TE_{020} in which the electromagnetic loss of the walls can be kept very small.

There has been no attempt here, obviously, to write down anything like a complete theory. Details of the motion, the drifting and the relative phasing between various electrons in the electromagnetic wave require detailed analysis. There are also problems having to do with distribution of the energy of the injected electrons between axial velocity and circumferential velocity, which effects the relative phasing. These are all important, but beyond what is intended in this description.

There are various important problems which cannot be discussed in detail here. For example, the electrons all come from a cathode at zero voltage, and all have the same total energy when injected into the interaction waveguide. However,

injection must be into a magnetic field region, or at least from one value of a magnetic field into another and the specific means of injection, that is, the trajectories from the cathode into the waveguide may give different partitions of the total energy between the axial and transverse velocities. The total energy for all electrons is the same but, depending on how they get from the cathode into the waveguide, the rotational velocity they may acquire may differ.

An excessively large spread in axial velocity means that the doppler shifted frequency for various electrons will be different, and even though only one generated frequency results, i.e., there is only one wave produced in the waveguide, the effective interaction between the electrons and the field is reduced because of phase slippage by some electrons which do not have the desired axial velocity. It is mentioned here because it is a problem, currently, and probably will continue to be a problem, but is not something serious enough to reduce the importance of these devices. It may mean that one does not get the desired efficiency.

As to performance achieved with these devices, it has been outstanding. In the U.S., the outstanding performance has been 215 kw c.w. at about 28 GHz in a single cavity oscillator with about 35% efficiency, 250 kw pulsed power with high duty cycle has been achieved at the same frequency. It should be pointed out that the effort in getting to this power did not relate to the interaction mechanism which was proven to be effective as theoretically predicted a number of years earlier. Much of the effort since has been on important but peripheral technical problems about output windows or r.f. power getting into the

wrong regions of the tube, parasitic modes, and the whole question of handling such powers, getting it from a tube into an output waveguide. At no point has there been really a problem typical of slow wave tubes, such as interception of the beam by the r.f. circuit, the delicate problems of focusing to get the beam through the circuit, etc. It is important to stress this because there are going to be attempts to go to much higher frequencies, and one can, on the basis of this data quoted, have confidence that similar powers will be achieved at higher frequencies and presumably with the same annoying but peripheral problems. The problem of dissipation and focusing of high density electron beams through too small an aperture will not be a problem, just as it was^{not}/in this particular tube.

There has also been considerable effort on traveling wave devices using this kind of interaction, in which case the electromagnetic wave is a traveling wave in an unloaded waveguide. Bandwidths of the order of 6 to 8% at about 25% efficiency and 120 kw have been achieved, but this was at a lower frequency (about 6 GHz), with the important feature again that in principle such device could easily go to c.w., and could be extrapolated to much higher frequencies. It would seem that the present intention to try to achieve very high power in the 2 or 3 millimeter range with this kind of device seems to be not at all unreasonable.

Ubitron. (Free electron laser.) This device is also a fast wave device. In this case a periodic motion of the beam is obtained by having the beam pass through a periodic series of magnetic deflectors (sometimes also called an undulator), in which the polarity of the magnetic deflecting field reverses periodically with some designed spacing. Here also, as in the gyrotron, the condition for interaction is that the doppler shifted electromagnetic wave in the waveguide, through which the electrons are passing, must be equal to the periodic frequency of the motion of the electrons passing through the magnetic array. This frequency, of course, is the velocity of the electrons divided by the magnet spacing.

The first device of this kind was built and a theory developed perhaps seventeen or eighteen years ago at around 6 millimeters, and put out something of the order of a megawatt of power. Interest in the device seemed to disappear and it was later revived in a completely different form. The theory was redone in a different way, and because of the frequency at which it was first demonstrated (in the micron wavelength range), it was named the Free Electron Laser. It was also true that in this latter case the circuit was not a waveguide, but a resonator which consisted of two spherical mirrors which provide a so-called Gaussian mode between the mirrors, of the kind commonly used in most lasers. The present effort on this device is directed toward the optical range, but it should be pointed out that the voltage required for the electrons are in the range of several tens of megavolts.

AD-A088 745

STANFORD UNIV CA EDWARD L GINZTON LAB
MICROWAVE TUBES. (U)
JUN 80 M CHODOROW

F/G 9/1

UNCLASSIFIED

6L-3131

AFOSR-77-3351

AFOSR-TR-80-0692

NL

2 of 2

201 8/1/80



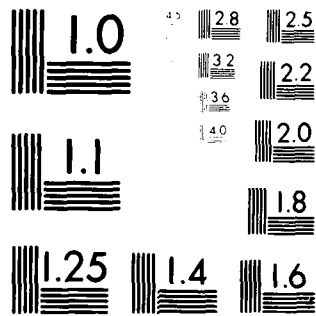
END

DATE

FORMED

10-10

DTIC



MICROCOPY RESOLUTION TEST CHART
NATIONAL BUREAU OF STANDARDS-1963-A

The interaction in a ubitron is more complex than in the gyrotron. Again, one requires the equality of the frequency of the periodic motion and the doppler shifted frequency of the electromagnetic. If one analyzes the equation stating this, it then becomes apparent that this means that while the electron is moving through one cycle of its periodic motion, which is the distance between two equivalent deflecting magnets, one cycle of the electromagnetic wave slips by the electron. The effect on any particular electron depends on the phase of the electromagnetic signal relative to the phase of its motion through the magnet. All electrons have the same parameters in that one cycle of the electromagnetic wave slips passed the electron during one cycle of its periodic motion, but different electrons obviously have different phase relations between these two periodic effects. This phase relation between the electron periodic motion and the e.m. wave continues over many cycles with a cumulative effect on the average motion of the electron, and the net effect for all the electrons is a consequence of taking this into account.

The mechanism involved in this joint action of the e.m. wave and the magnetic fields in some ways is related to the behavior of electrons in a magnetron-like device. That is, the important mechanism for bunching of the electrons depends on the common effects of the transverse magnetic field which is doing the deflecting, and the transverse electromagnetic field. If one examines the phase relation between these two separate deflecting fields, one can see that for any electron over each cycle of periodic motion, there will be an average

value of E/B where the electric field is due to the r.f. field and the magnetic field is due to the periodic magnets and this causes an axial drift of the electrons whose magnitude and direction will depend on the phase.

To state this perhaps more explicitly, some electrons will see a peak magnetic field and a peak electric field at the same time and a half-cycle later in the periodic motion of the electrons because of the slippage of the electric field will see again peak fields with both fields reversed in sign and, therefore, the average value of the ratio of E/B will have some non zero value because of this common sign reversal. This results, just as in the crossed field device in an axial motion of these electrons, on the average. Some later electrons will go through a similar behavior except that they will see the electric field a half cycle later at each point in their periodic trajectory. The statements made above will still apply. There will still be an average value of E/B , but this average value will have a reverse sign and, therefore, there will be an axial motion in the opposite direction for those electrons.

The electrons in-between these two sets see similar effects, but of smaller magnitudes because the phasing between the B and E are not quite optimal. This produces axial bunching and if all the parameters are correct, one can show that these bunches have been formed at such phase points in the electron beam that they have maximum transverse velocity in the deflecting magnets, when there is maximum transverse E

field, and, therefore, there is a transfer of energy from the electrons to the field due to this transverse motion of the electrons. But the bunching has been produced by an E/B interaction which has produced axial bunching.

In this ubitron, also then, as in the gyrotron, we have an interaction with a fast electromagnetic wave. The electron beam does not have to be small compared to the wavelength, it interacts with a wave in a waveguide (in the case of the earliest work at around 6 millimeters), or with a transverse electric field of a free space wave (in the case of the later work in the micron range). It should be stated that, in both cases, the efficiencies were on the order of a few per cent. Current theories seem to indicate that without some major changes in design or some new ideas, that the efficiency, at least for a single pass of the electrons through the electromagnetic field, will be quite low. Modifications have been suggested to recirculate the electrons through the same magnet structure many times restoring whatever energy they have lost in each pass by some other means. This may be a way of getting to higher efficiencies, but certainly, currently, at least in the microwave range (defined as perhaps up to 300 GHz), the gyrotron seems the more efficient device).