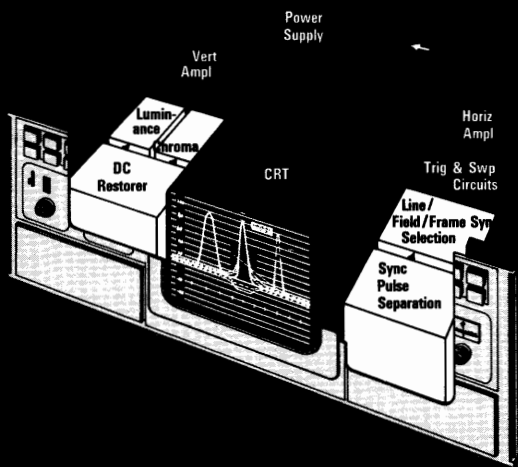




Circuit Concepts



**Television
Waveform Processing Circuits**

TELEVISION WAVEFORM PROCESSING CIRCUITS

BY

GERALD A. EASTMAN



CIRCUIT CONCEPTS

FIRST EDITION THIRD PRINTING MARCH 1969
062-0955-00
PRICE \$1.00

© TEKTRONIX, INC.; 1968
BEAVERTON, OREGON 97005
ALL RIGHTS RESERVED

CONTENTS

PREFACE

1	INTRODUCTION	1
2	BASIC TELEVISION SYSTEMS REQUIREMENTS	3
3	TELEVISION WAVEFORM COMPONENTS	11
4	DC LEVELS IN VIDEO SYSTEMS	19
5	DC RESTORER CIRCUITS	23
6	FEEDBACK DC RESTORER CIRCUITS	35
7	WAVEFORM SYNCHRONIZATION	43
8	SYNCHRONIZING PULSE SEPARATION	47
9	SYNCHRONIZING PULSE PROCESSING	77
10	BASIC COLOR CONCEPTS	97
11	COLOR TELEVISION SYSTEM REQUIREMENTS	103
12	COLOR MODULATION AND DEMODULATION	109
13	CHROMINANCE SUBCARRIER REGENERATORS	123
14	CHROMINANCE SYNCHRONIZING	145
15	SUBCARRIER PROCESSING CIRCUITS	151
16	CIRCUIT CONCEPTS IN THE FUNCTIONAL BLOCK SYSTEM	175
	INDEX	185
	CIRCUIT CONCEPTS BY INSTRUMENT	187

PREFACE

The use of cathode-ray oscilloscopes in the television equipment manufacturing and broadcast industry has resulted in the design of special application oscilloscopes generally called "waveform monitors." As a result, circuits not generally found in the conventional laboratory oscilloscope are designed and used in the waveform monitor to facilitate the display and measurement of various video waveform parameters.

The purpose of this book is to familiarize the reader not only with the operation of the specialized circuitry used in Tektronix waveform monitors, but also to outline the general and specific purpose of the circuit. As a result, the reader will find, in addition to the circuit concept information, sufficient television theory to analyze the intended circuit operation.

The television theory presented in this book is *not* intended to be complete, but only to provide sufficient information to illustrate the circuit concept of operation. The circuit description then provides an illustration of the concept as applied on a practical basis.

It is hoped that the additional circuit insight will enable the reader to extend the measurement versatility of the test instrument as well as facilitate the periodic maintenance of the test instruments.

INTRODUCTION

The electronic communication system called television is designed to reproduce the sensation of direct vision. The technical aims and compromises of such a system determine the requirements of the ideal waveform monitor oscilloscope. The system concepts of television necessary to artificially reproduce direct vision are, of course, based on the natural visual process of the eye.

visual
perception

The eye perceives six basic pieces of information when directly observing a scene:

1. Brightness -- the visual perception of the average background illumination.
2. Contrast -- difference between light and dark areas.
3. Detail -- structural information about particular geometric shapes.
4. Motion that may exist in the scene.
5. Chromatic content -- the colors of objects.
6. Stereoscopic content -- relative three-dimensional position between objects.

The fact that the eye can analyze and interpret these six different pieces of information about any scene testifies to the sensitivity and complexity of the sense of vision.

The communication system is intended to match the performance of the eye to some degree by recreating the six different pieces of information about a scene. All six pieces of information may be conveyed to some degree in a stereoscopic color-television system, but

in practice the most important information about a scene can be transmitted in a communication system without color or stereoscopic content. Since color transmissions provide optional information to increase the degree of realism, the concepts of transmitting color information and the circuit techniques involved will be treated together in a separate section. The sixth piece of information, stereoscopic content, is almost never used and will not be discussed.

electron-
ically
reproduced
picture

As in photography, technical and economic considerations indicate that compromises are necessary in the transmission of the five pieces of visual information now commonly included in the television system. The first and most important technical consideration is the fact that any electronic system is capable of transmitting only one bit of information at a time. Therefore, the picture will have to be broken down into small elements, transmitted, and then reconstructed at the receiver. All the elements of an electronically reproduced picture must appear simultaneously to the eye, so two alternative possibilities exist:

1. Use a separate electronic circuit for each element of the reconstructed picture. If any image detail is required, a minimum of 100,000 circuits will be required to simultaneously reconstruct the picture --- somewhat impractical for a public communication system.
2. Transmit all the elements of the picture in rapid succession over one electronic system. Due to persistence of vision, the impression to the eye will appear to be simultaneously reconstructed. Thus, direct vision can be recreated artificially.

Of the two possibilities, the second is the most practical and economical. However, several requirements must be met to make the system workable.

2

BASIC TELEVISION SYSTEM
REQUIREMENTS

simultaneous appearance	The first requirement is that the entire picture must be sequentially reconstructed in a short enough time -- less than one-tenth of a second -- to appear simultaneous to the eye. If the illusion of motion is also required, many complete pictures appearing in rapid sequence are necessary. Experimental tests have indicated that each picture should be completely displayed in one-thirtieth of a second or less, to create the illusion of motion.
uniform linear scanning	The next requirement is the reconstruction of the picture itself. Many different methods of reconstructing the picture are possible, but both the transmitter and receiver must use exactly the same system. The system actually adopted is called a "uniform linear scanning system" where the camera "sees" the scene at a fixed uniform rate according to an established pattern. The receiver is then made to reconstruct the picture (scene) at exactly the same rate and pattern as the camera. The scanning pattern is essentially the same as the action of the eye when reading this book -- starting at the upper left-hand corner, moving downward a line at a time, and finally terminating at the lower right corner. The scanning process from top to bottom forms one complete image.
scanning pattern	The scanning pattern adopted in North America consists of 525 horizontal lines laid down in a vertical sequence. The complete scanning of 525 horizontal lines form what is called a picture frame. The horizontal lines are one-third longer than the height of the image, forming a picture frame aspect ratio of 4:3.
aspect ratio	

flicker

The illusion of motion can be created by producing a picture frame consisting of 525 lines in one-thirtieth of a second, but image flicker from a television picture can become quite noticeable at a 30 cycle frame rate. When compared to movies, the problem of image flicker requires special attention because:

1. Television images are usually adjusted by the viewer to be brighter than movies, because of the inconvenience of darkening rooms in the home.
2. The viewer is closer to the image than in a movie theater.
3. The picture is scanned progressively rather than displayed simultaneously as from a movie projector, causing the ratio of illumination time to non-illumination time to be very high.

interlaced
scanning

The minimum picture repetition rate at which the eye perceives flicker increases with the logarithm of the brightness of the overall scene. For example, a television image scanned at the rate of thirty frames per second will appear to flicker when the image brightness level is about 1 ft-lambert. A brightness limitation of 1 ft-lambert is not acceptable because the resolving power of the eye is reduced when the image brightness is lower than 3 ft-lamberts. When the image repetition rate is increased from 30 to 60, flicker is not noticeable until the brightness is about 115 ft-lamberts. Typical television image brightness is greater than 20 ft-lamberts, so the scanning rate will have to be greater than 30 frames a second -- without increasing the number of scanning lines. The technique used to increase the scanning rate is called *interlaced scanning*. The action is somewhat similar to movie techniques. Movie film runs through the projector at the rate of twenty-four picture frames a second. This rate is enough to create the illusion of motion, but each frame is projected on the screen twice, so the eye sees 48 images a second. By increasing the number of images a second, the illumination to nonillumination time ratio (flicker) is reduced.

The television frame is also divided into two parts as in movie systems, but the method is different because the television image is reproduced a line at a time. The following is an example of interlace:

1. The image repetition rate is increased
 5. to the eye without having to increase
2. by scanning all the odd numbered lines
 6. the system bandwidth. The only purpose
3. and then scanning all the even numbered
 7. of the interlace scanning is to
4. lines. Sixty images are then presented
 8. effectively reduce image flicker.

field The television image divided into two parts in an
 frame interlaced pattern forms a "field"; two interlaced
 fields form one complete frame.

synchro- The scanning beam of the receiver must be at the same
 nization location in the reconstructed picture as the camera
 scanning beam, so some method is needed to synchronize
 the receiver to the camera. To simplify the receiver
 circuitry, synchronizing pulses are transmitted at the
 beginning of each line and at the beginning of each
 field to insure that the receiver scanning beam is at
 the same vertical and horizontal location as the
 scanning beam of the camera.

The last and most difficult requirement is providing sufficient picture detail. To reconstruct a picture with a reasonable amount of detail, the picture should have a minimum of 100,000 elements. If thirty pictures a second are to be completely scanned, three million bits of active picture information a second will have to be transmitted.

resolution The image detail perceived by the eye is determined by the "resolution" capability of the image reproducing system -- the number of basic picture elements that can be reproduced and discerned. In the printing process, resolution is determined by the number of "halftone" dots available in a given area (usually millimeters). In photographic film, the number of line-pairs per millimeter determines the resolution. (A line-pair consists of one adjacent black and white line.) In the television system, adjacent black and white lines are counted separately

so the number of resolution lines reproduced on a television cathode-ray tube is twice the number of photographic lines required to present the same amount of picture detail. Unlike film which has no line structure, the detail perceived from a scanned TV picture must be determined in terms of both vertical resolution *and* horizontal resolution.

vertical
resolution

The scanning lines are laid down on the CRT phosphor horizontally, so the resolution of the television picture on the vertical axis is limited to the number of horizontal lines used to completely scan the scene (Fig. 2-1A). The number of scanning lines, for television systems in North America, has been set at 525 lines, but 35 of the scanning lines are used for non-picture-element information, leaving about 490 lines for active picture information.

If the resolution on the vertical axis is less than the width of the scanning beam, the element, when reproduced, will be the light-average of the original element (Fig. 2-1B).

Due to the problem illustrated in Fig. 2-1, experience and calculations indicate that about 350 resolution lines, or 70% of the active scanning lines, can actually be resolved. This resolution is comparable to a 16 mm home movie system which has about a 450 line resolution.

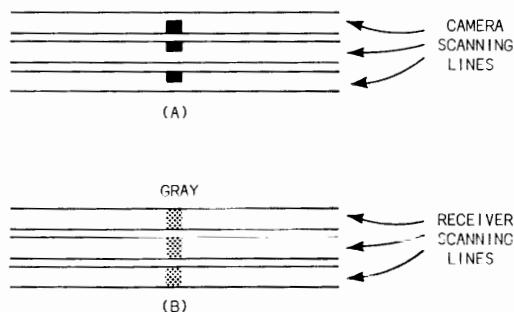


Fig. 2-1. Vertical resolution limitation.

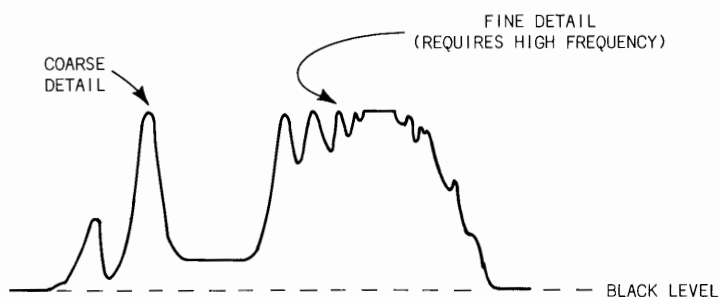


Fig. 2-2. Typical waveform of one horizontal scanning line.

horizontal
resolution

The horizontal resolution ideally should be approximately the same as the vertical resolution. Since the picture is one-third wider than the height, the horizontal resolution should be four-thirds times 350 (the vertical resolution) or 466 lines. If the cathode-ray-tube spot geometry is neglected for the moment, the resolution on the horizontal axis is limited by the maximum number of times the CRT-spot tracing out the scanning line can be turned on and off by a control voltage. As illustrated in Fig. 2-2, the control voltage rate-of-change, and therefore the horizontal resolution, is limited by the bandwidth of the system. As mentioned earlier, adjacent black and white lines are counted separately; so if the positive half of a sinewave produces a white element, and the negative half of the sinewave produces a black element, 233 cycles-per-line will be required to produce the maximum of 466 lines.

horizontal
resolution
limitation

Now, if the CRT scanning spot itself is considered, the spot shape will appear slightly elliptical, because the spot is moving along the horizontal axis faster than the light can decay from the phosphor. The result is a slight reduction in the horizontal resolution capability to about 450 lines or 225 cycles per scanning line. The required bandwidth then becomes 225 cycles times the time necessary to scan one line. Since 525 lines must be scanned in one-thirtieth of a second, one line must be scanned in

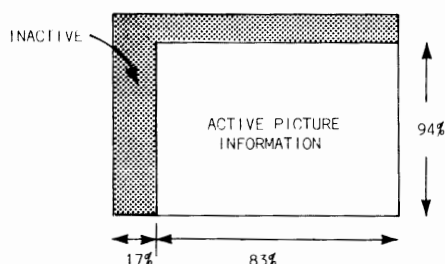


Fig. 2-3. Percentage of time used to transmit active picture information. Inactive portions are used for synchronizing information.

63.5 μ s. As illustrated in Fig. 2-3, 17% of the horizontal scanning time is used for nonactive picture information leaving a net of about 52 μ s of scanning time left. If 225 cycles are to be produced in 52 μ s, the maximum sinewave frequency will have to be 4.26 MHz.

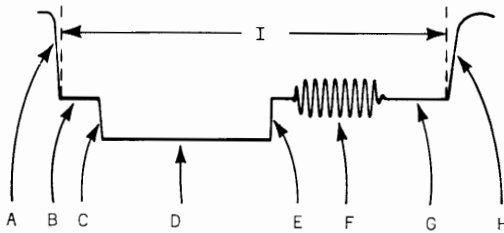
The maximum number of elements possible in each television picture frame then is 450 horizontal lines times 350 vertical lines, or about 157,000 elements -- comparable to the 200,000 picture elements per frame of a 16-mm home movie system.

As a final note about picture detail, some industrial closed circuit television systems require greater picture resolution than the conventional television system is capable of reproducing. Doubling the bandwidth of the system will increase the overall picture resolution 50% since only the horizontal resolution has been increased. The bandwidth will have to be increased from 4.25 MHz to at least 16 MHz to double both the horizontal *and* the vertical resolution and hence the overall picture resolution. (Vertical axis resolution is doubled by increasing the number of scanning lines.)

Meeting just three basic requirements (electronically scanning a scene in one-thirtieth of a second with a reasonable amount of detail) has resulted in a complex television waveform. This complex waveform is made up of neither pulses nor sinewaves, but a combination of pulses and sinewaves conveying seven different types

composite
video

of information to satisfactorily reconstruct a televised picture. Each of the seven components are intended to eventually perform specific functions in the receiver to reconstruct a quality television picture. The seven components, when combined into one continuous electrical wave, form what is called the television composite video waveform, or simply "composite video." The television waveform oscilloscope must extract, process, and display each one of the seven different pieces of information.



- A. Video voltages at right side of picture.
- B. Front porch $1.59\mu\text{s}$ or 2.5% of horizontal line interval (H).
- C. Leading edge of sync.
- D. Sync tip $4.76\mu\text{s}$ or 7.5% H.
- E. Trailing edge of sync.
- F. Color sync burst 8-11 cycles of 3.579545MHz (present only during colorcast).
- G. Back porch $4.76\mu\text{s}$ or 7.5% H.
- H. Video voltages at left side of picture.
- I. Horizontal blanking interval $11.11\mu\text{s}$ or 17.5% H.

Fig. 3-1. Horizontal line scanning interval.

3

TELEVISION WAVEFORM
COMPONENTS

The seven components contained in the television waveform are:

1. Horizontal line synchronizing pulses.
2. Color sync (burst).
3. Set-up (black) level.
4. Picture elements.
5. Color hue (tint).
6. Color saturation (vividness).
7. Field synchronizing pulses.

Since the television picture is reconstructed in time sequence by line scanning, most of the individual components and their purpose can be examined by discussing one complete scanning line.

horizontal
line sync

The first component of the composite video consists of three segments. (See Fig. 3-1.) The first segment, the front porch, blanks the receiver just before the start of horizontal retrace in the camera, and also isolates the sync pulse from the active picture information of the previous line. The next segment is the horizontal line sync pulse whose leading edge synchronizes horizontal line sweeps of the receiver to the camera. Following the trailing edge of the sync pulse is the back porch, which provides picture blanking until the start of the active line of video information, isolating sync and beam flyback. This isolation gives the sweep circuitry time to "settle down" before presenting another line of picture information. The front porch, sync pulse, and back porch comprise the horizontal sync and blanking interval, the first component of the composite video waveform.

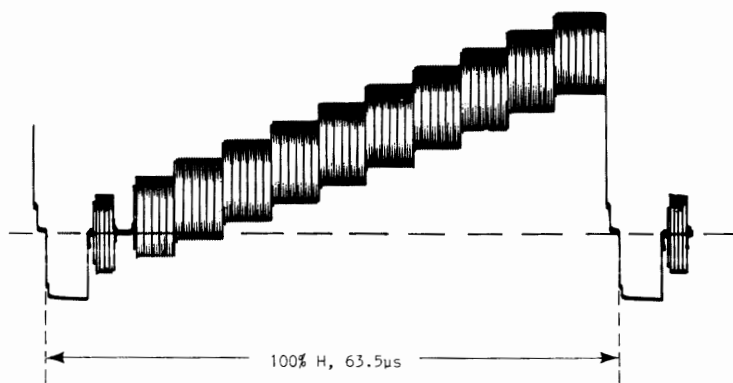


Fig. 3-2. One horizontal scanning line (100% H) illustrating the reference set-up level.

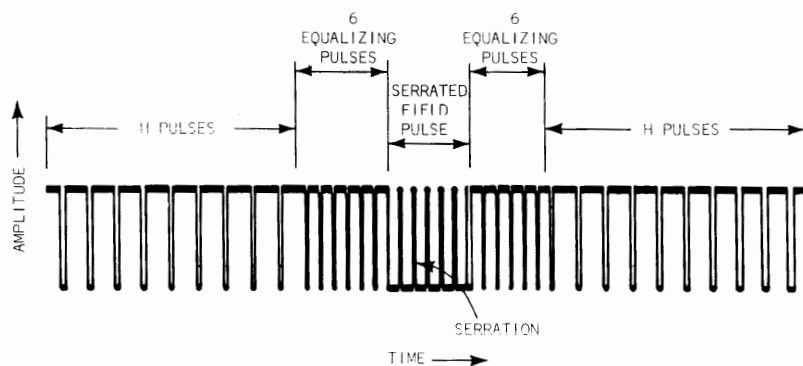


Fig. 3-3. Field synchronizing pulses.

color burst	The second component is one of the optional components -- the color burst (Fig. 3-1). During the transmission of a color program, color sync in the form of a burst of sinewaves is added on the "back porch" of the sync pulse to frequency-and-phase lock the picture color information.
set-up (black) level	The third component is the set-up level often referred to as the black level (Fig. 3-2). RETMA standard RS-189 (amended 12/66) specifies the set-up level as 7.5% above the back-porch blanking level. The average picture brightness level discussed in the DC restorer section is referenced to the set-up level.
picture detail	The fourth component is the black and white picture detail information. The gradations of light including the <i>average</i> light of the entire scene are included with the picture detail information (Fig. 3-2).
hue color saturation	The fifth and sixth components contain hue and color saturation information. Both components are optional but when transmitted are added to the existing black and white picture detail portion of the waveform.
field sync pulse equalizing pulses	The last component is the field synchronizing information comprising eighteen pulses (Fig. 3-3). Six equalizing pulses precede the serrated field synchronizing pulse referred to simply as the "field sync pulse" and six more equalizing pulses follow. The purpose of the initial equalizing pulses is to assure proper field interlacing and maintain horizontal line sync. The usefulness of the last group of equalizing pulses is not too obvious in light of present-day circuitry; however, according to the literature, at the time the video standards were adopted, field scanning generators were triggered rather than synchronized. This important difference required more accurate shaping of the integrated field sync pulse while at the same time assuring continuity of horizontal scanning repetition rate.
serrated field sync pulse	The serrated field sync pulse consists of six pulses with a duty factor of 84% (compared to 4% duty factor for the normal line sync pulse) to charge the field integrating network. The field sync pulse is serrated to simultaneously maintain horizontal synchronization. (A more complete discussion of the field sync information will be deferred to the discussion of synchronizing processing circuitry.)

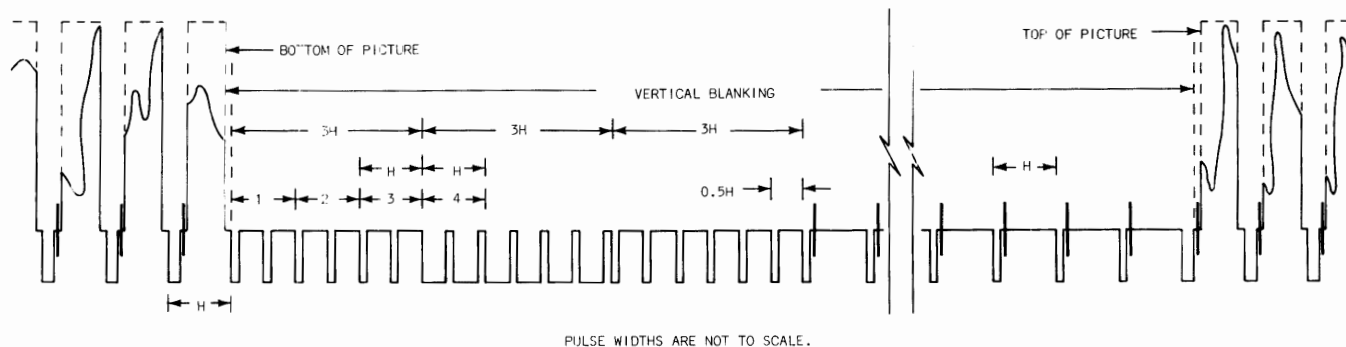


Fig. 3-4. Expanded illustration of composite video during the field (vertical) blanking interval.

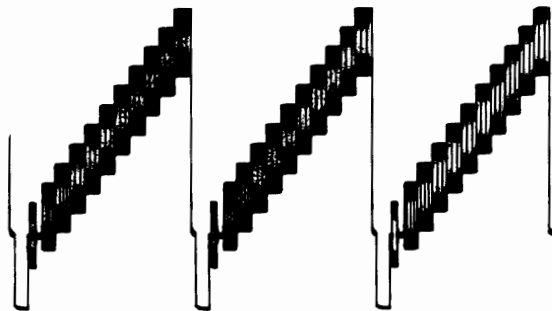


Fig. 3-5. Expanded illustration of composite video horizontal scanning lines.

composite
video

These seven components, when combined to form a continuous waveform, form the composite-video waveform illustrated in Fig. 3-4 and Fig. 3-5.

waveform
distortions

Once the composite video is properly formed care must be taken to insure that waveform distortions do not occur -- either in the transmission or the reception and processing of the video, -- if a quality picture is to be reconstructed. The waveform monitor oscilloscope is intended to "monitor" *and* measure:

1. The existence of composite-video waveform distortions.
2. The amount of waveform distortions.

No attempt will be made to discuss all the different types of composite video waveform distortions that can exist. However, these distortions fall into three general categories:

1. Waveshape distortions.
2. Amplitude errors.
3. Crosstalk between two or more of the components in the composite video waveform --- particularly between black and white detail information and chroma information.

Accumulated signal distortions of many varieties cause picture quality degradation, which must be quickly identified and diagnosed. The broadcasting industry requires an oscilloscope that will quickly and accurately determine the nature and amount of picture quality degradation. Over the last ten years, signal distortion limits have been established in terms of the subjective impairment of the displayed picture. Based on these established distortion limits, suitable test signals have been developed by which quantitative measurements of picture quality can be made. To a very large extent, these test signals determine the requirements of the television waveform circuits that will be discussed in the subsequent sections.

Display and measurement of any of the seven waveform components by an oscilloscope can be done only by first selectively separating the individual components within the oscilloscope. The oscilloscope then utilizes the separated information in two ways:

1. Direct display of the waveform (such as the separated color difference signals or the

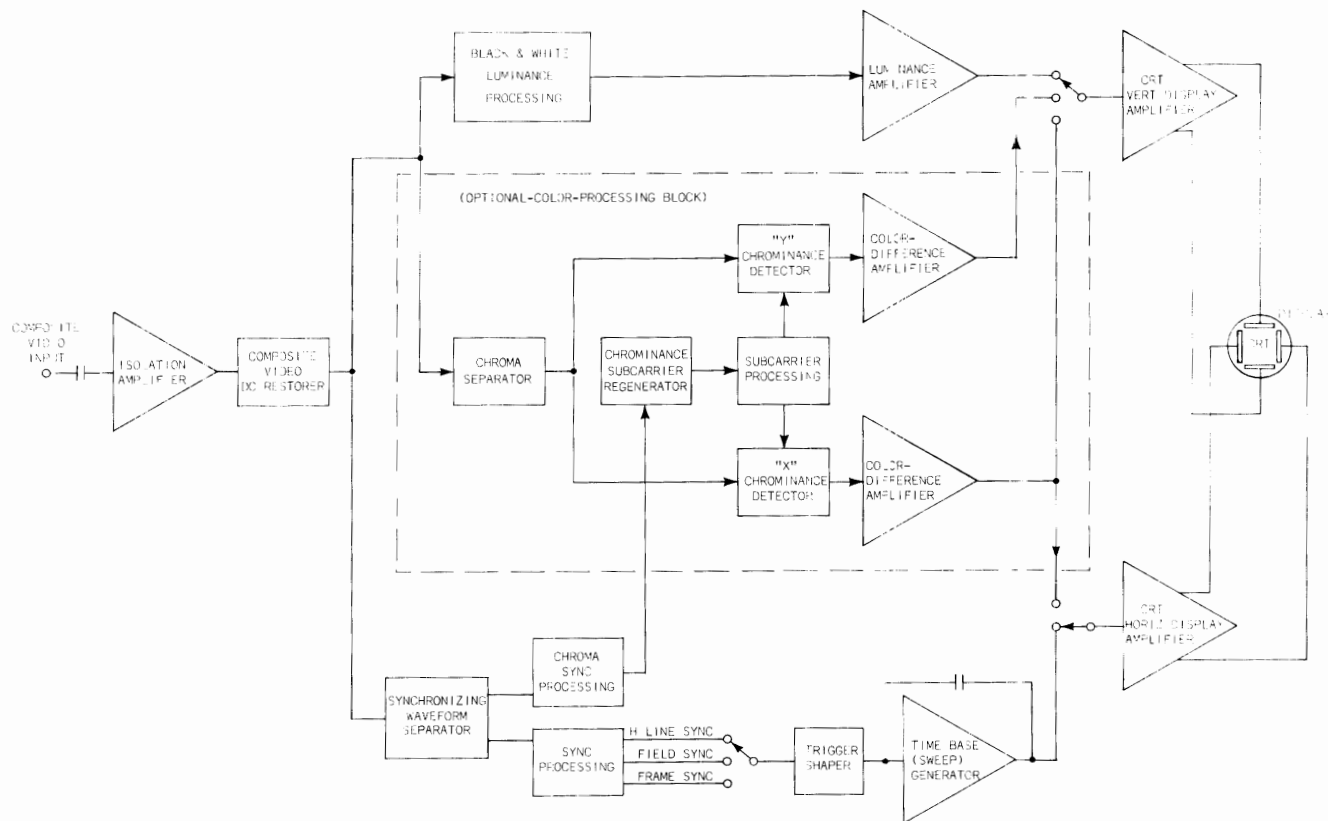


Fig. 3-6. Simplified block of oscilloscope signal processing.

luminance picture detail information) on the oscilloscope CRT.

2. Indirect display (such as the synchronizing information obtained from the composite video to time-reference the waveform display on the CRT.)

oscilloscope system requirements Requirements of an oscilloscope to display the composite video (or selected portions of the composite video) on the face of the cathode-ray tube in terms of time and voltage can be more easily visualized by observing the simplified signal processing block diagram of Fig. 3-6. The block diagram of Fig. 3-6 is intended only to illustrate the signal-processing relationships and *not* the actual signal processing techniques which will be covered in subsequent chapters. The chrominance processing blocks inside the dashed line are optional and not found in oscilloscopes intended primarily to measure a conventional black-and-white composite video waveform.

A complete discussion of each signal processing block appears in the following chapters. A summary of the detailed waveform processing system of the various television waveform oscilloscopes will be studied in the final chapter.

4

DC LEVELS IN VIDEO SYSTEMS

If the composite video signal is to eventually produce various shades of gray at the picture tube, some portion of the waveform must logically represent the absence of light and other portions of the waveform must represent shades of gray. The video signal amplitude must then be "unidirectional" -- one DC level always representing black and a second level DC always representing maximum white. However, DC levels cannot be transmitted in an AC system.

Any unchanging test pattern observed either on a picture monitor or waveform monitor produces a relatively constant video waveform and DC levels are not needed. Under these conditions, the DC levels can be manually inserted at the picture monitor by setting the contrast and brightness controls for a pleasing picture; the waveform on the monitor oscilloscope can be vertically positioned to a reference line drawn on the face of the cathode-ray tube.

average
picture
level

The situation changes when live program material is being observed, because the average picture level (background brightness of the picture) constantly varies. In the absence of DC reference levels in both the camera and picture monitor, the scene displayed on the picture monitor will have poor black tones resulting in low quality pictures that are hazy and possibly streaked. Horizontal picture tearing can also exist from partial loss of sync when the average picture level changes. What is the average picture level (often referred to as APL) and why is it important in television?

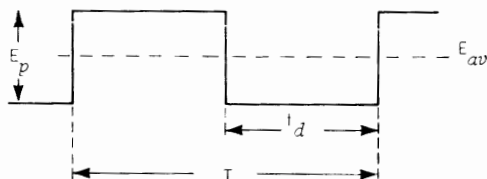


Fig. 4-1. Average DC level.

Reviewing briefly, any waveform has an average DC level which can be mathematically stated as

$$E_{av} = E_p \times \left(\frac{t_d}{T} \right) \text{ where } E_p \text{ is the peak voltage of the}$$

pulse, t_d is the duration of the pulse, and T is the period or repetition rate of the pulse (Fig. 4-1). The average picture level then is defined as the DC component of the television signal which contains information concerning the *average* light in the televised picture.

Before the average picture level can be accurately reproduced on a picture monitor, that level must be referenced to the black level or another portion of the composite video waveform which remains constant with respect to the black level. The black level *DC-insertion* is initially accomplished in the camera by clamping the black portions of the picture signal, which occur at the end of each line and frame, to some fixed DC voltage level. During the active scanning time all values of light are then represented by various signal levels extending unidirectionally from the clamped DC black level. A reference now exists to accurately reproduce the average picture brightness. As an example, Fig. 4-2 shows two successive sync pulses where E_{av} is 0.92 V with respect to ground. Any video added above the 1-V level will move the average level more positive, giving information about the average brightness of that particular scanning line independent of the detail information that may be contained in that same scanning line. However, as seen in Fig. 4-3, if the signal is passed through a capacitor the average DC becomes zero and erroneous brightness information will result. The DC reference can be reinserted with a simple peak-detector called a DC restorer, which holds the waveform extremities to a *fixed* DC value.

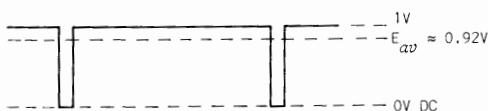


Fig. 4-2. Successive line sync pulses average voltage level.

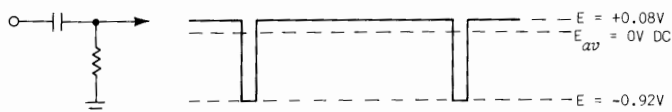


Fig. 4-3. Average DC level becomes zero if the signal is passed through a capacitor.

DC restorer
function
in picture
monitors

DC restorers perform a variety of functions when processing composite video for display on either a picture monitor or waveform monitor oscilloscope. The two principal functions of a DC restorer in a picture monitor are to:

1. Clamp the composite video to a fixed DC level and remove any 60-Hz hum and low-frequency distortions that may be present.
2. Clamp composite video for effective sync processing.

DC restorer
function
in waveform
monitors

The two principal functions of a DC restorer in a waveform monitor oscilloscope differ slightly from restorers used in a picture monitor. The two principal functions are:

1. Maintain average DC levels in the vertical system *without* destroying error signals and hum that may exist on the waveform.
2. Clamp sync tips to a fixed DC level in such a manner that sync processing and finally jitter-free sweep triggering can take place.

These two DC restorer applications in waveform monitor oscilloscopes impose more rigid requirements on the circuit design than picture monitor applications. As a result, the DC restorer circuits are more varied in configuration and operation. At this point then, the DC restorers may further be classified into two types:

slow
restorers

1. Slow restorers have long time constants requiring many frames to establish the DC axis. Generally, slow restorers are desirable in the oscilloscope vertical systems for observation and measurement of low frequency abnormalities such as 60 Hz hum without adding distortion to the waveform.

fast
restorers

2. Fast restorers have much shorter charge-time constants, requiring from one horizontal line sample to one field to establish the DC axis. Generally, fast restorers are used to solidly clamp sync tips prior to sync processing.

5

DC RESTORER CIRCUITS

sync tip
clamping

Fig. 5-1 illustrates a simple DC restorer. The necessary sync tip clamping is produced in the grid circuit of V2 to move the average DC level of the composite video from zero volts to some more negative value. The sum of the correcting DC signal and the composite video waveform appears at the output plate circuit (inverted). The result is a fixed DC level at the sync tips which remains constant independent of signal amplitude variations so that all the sync tips line up evenly.

The circuit operates as follows: Quiescently, V2 is operating at zero grid bias. When negative-going composite video is applied to the control grid through the coupling capacitor C1, the grid is forced slightly positive into a grid current condition by the positive-going sync pulse tip. This current charges the capacitor to a *new average value*. R1 is very large so that during the active picture scanning time, when the grid is driven negative, the capacitor will only partially discharge through R1.

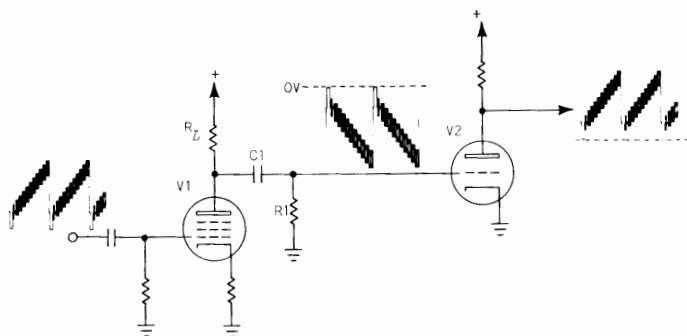


Fig. 5-1. Simplified DC restorer.

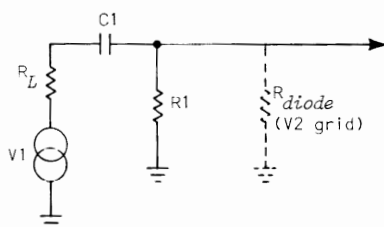


Fig. 5-2. Simplified DC restorer equivalent circuit.

simple
DC restorer

Fig. 5-2 shows the equivalent circuit of the simple DC restorer during the sync tip interval when the diode (V2) forms part of the circuit. When the composite video is initially applied to the DC restorer, the sync pulse tips will have a small tilt on the top because the charging time constant ($C1-R_{diode}$) is relatively short, but still longer than the pulse duration. When C1 becomes charged after a number of pulses, less current flows through the diode causing its forward resistance to increase. The charging *time constant* continues to increase until eventually it is much longer than the pulse duration.

Finally at equilibrium the charge through $C1-R_{diode}$ equals the discharge path through $C1-R1$. The discharge time is predictable, but since a conducting diode's forward resistance follows a $3/2$ power law the charge time is more difficult to predict.

After nearly complete restoration has taken place, over a period of many fields, the sync tips are still slightly more positive than zero volts.

simple DC
restorer
limitations

When serrated field sync pulses occur, the charging time is suddenly increased from an 8% duty factor to an 84% duty factor, permitting C1 to charge closer to 0 VDC. At the output then, the level of field sync pulses will be depressed as shown in Fig. 5-3, because C1 has had an opportunity to charge closer to the correct 0 VDC level. The DC level difference between the line sync pulse tips and the field sync tips is a measure of the detection error of the simple DC restorer.

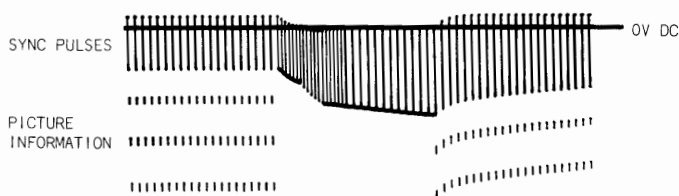


Fig. 5-3. Field synchronizing pulse DC restoration error.

The detection error can be minimized any one of several different ways:

1. Make C_1 larger, increasing the RC time constant. The *difference* in charging time allowed by the wider field sync pulse and the narrower line sync pulse then becomes a smaller percentage of the RC time constant, making the apparent detection error less. For example, if the time constant is 100 μ s, the field sync pulse will permit C_1 to charge 30% of the time constant, while the line sync pulse will allow C_1 to charge 4% of time constant. This 26% difference causes an appreciable change in DC levels between the field sync pulse tips and the line sync pulse tips.

If C_1 is now increased until the time constant is 1.0 ms, the field sync pulse will allow C_1 to charge 3% of the time constant while the line sync pulse will permit C_1 to charge 0.4% of the time constant. The difference of 2.6% creates a smaller detection error than is apparent in the first example.

2. Select a diode (or tube) whose forward resistance is as low as possible. The lower diode resistance will reduce the voltage drop across the diode, minimizing the detection error.
3. Increase the amplitude of the applied signal so the DC level error is a smaller percentage of the total applied signal amplitude.

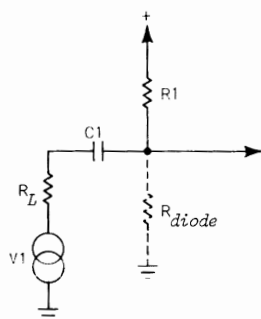


Fig. 5-4. Equivalent circuit of positive-bias DC restorer.

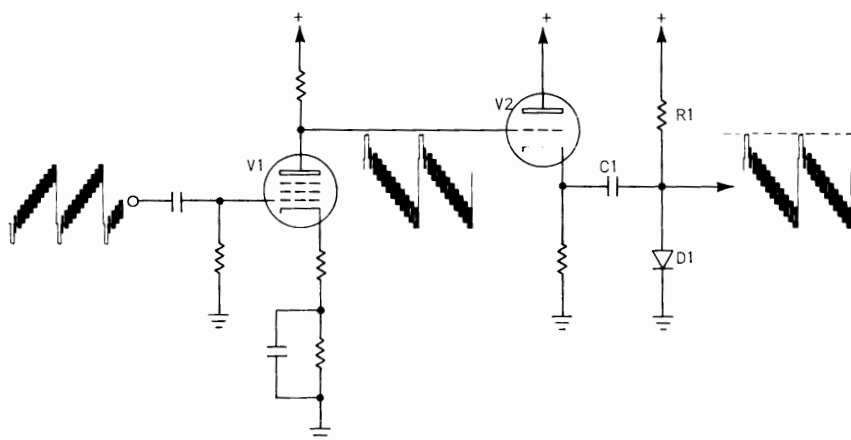


Fig. 5-5. Positive bias DC restorer.

circuit
limitation

In simple DC restorers, if the circuit has a short time constant, detection errors within one field change due to average picture level changes. If the circuit has a long time constant, the variation in detection errors is averaged out at the expense of partial loss of the DC component. For these two reasons the simple circuit just illustrated is not normally used in waveform-monitor oscilloscopes. However, the concept of operation is important to evaluate the performance of the following circuits.

positive
bias DC
restorer

The positive bias circuit illustrated in Fig. 5-4 differs only slightly from the previous circuit arrangement in that R_1 is now returned to a positive voltage. However, the operation is very much different.

The total DC voltage across R_1 is now the sum of the average signal DC level and the applied bias. If the bias voltage is several times larger than the average signal DC, the total voltage variation from a white scene to the field blanking level will be only a small percentage of that in the previous circuit. However, the field sync level error will still be present. R_1 , of course, will have to be larger to keep the current to a minimum.

Fig. 5-5 illustrates a typical positive bias DC restorer circuit. The composite video signal is amplified twenty times by pentode V_1 to minimize the field sync level error. High-frequency peaking has been added in the cathode circuit to preserve the leading edge of the sync pulses, to allow triggering a shaping multi with a minimum of time jitter. A cathode follower is added to reduce the driving impedance to the DC restorer. A semiconductor diode has been used to keep the forward voltage drop of the diode to a minimum. The diode is also selected for minimum leakage current so that C_1 is not excessively discharged between sync pulses.

Since the DC restorer output signal will be used for triggering, an effort is made to insure that sync tips will remain at a constant DC level independent of average picture level or 60 cycle hum. C_1 is made relatively small ($0.015 \mu F$) so the charging time will be short.

keyed
clamp DC
restorer

A further improvement on the positive bias DC restorer is the keyed clamp circuit illustrated in equivalent form in Fig. 5-6. By using a keying signal DC restoration can be greatly improved in special applications. Several advantages over the simple restorer are:

1. Impulse noise -- particularly in the direction of sync tips -- can cause the diode in a simple restorer to conduct, putting an unwanted erroneous charge on the capacitor. In the keyed clamp, noise occurring between the keying pulses has no effect on the restorer circuit.
2. The charge time is limited only by R_L , the driving source. Therefore, the clamped level will always remain the same eliminating hum and low-frequency components.
3. Any flat, recurrent portion of the video waveform can be selected for clamping -- either back porch or sync tip.
4. When the key is opened C is essentially infinite in size to the video signal, since the only discharge path is the leakage resistance of the tube.
5. The input signal can be either a positive-going or negative-going, low-amplitude waveform.

diode
switch

Fig. 5-7 illustrates a keyed clamp DC restorer similar to the ideal circuit of Fig. 5-6. The composite video is again applied to C1 from a low impedance cathode follower. The "switch" consists of two diodes simultaneously driven into conduction by push-pull pulses of equal amplitude coupled through C3 and C4. The charging time nonlinearity due to the changing diode forward resistance is not present in the double keyed restorer, because the current through the two diodes during pulse time is relatively constant -- the current required by C1 is only a small percentage of the total current flowing through both diodes. Most of the keying pulse amplitude is coupled through C3 and C4, because the time constant of C3-R1 and C4-R2 is long compared to the duration of the pulse. When the push-pull pulses forward bias the diodes D1 and D2, C1 starts to charge to the same potential as

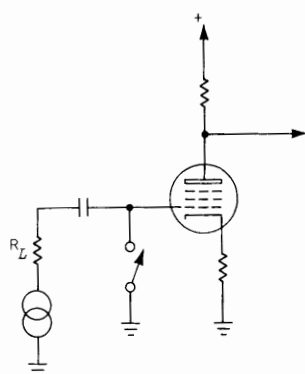


Fig. 5-6. Equivalent form of keyed clamp DC restorer.

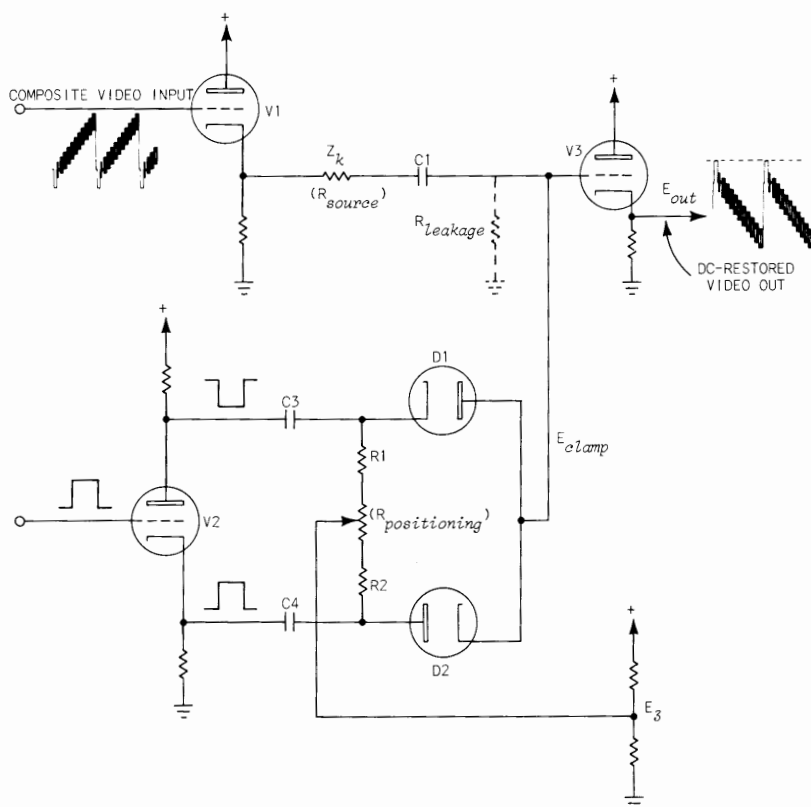


Fig. 5-7. Keyed clamp DC restorer.

E_3 . The charging time is determined by the driving source impedance of V1 and the driving impedance of the pulse source. In this case, if C1 is 0.01 μ F and the total driving impedance of V1 and V2 is about 5 K Ω , the time required to charge C1 to the correct level is about 50 μ s. Since the pulse duration is only 1 μ s, a *minimum* of 50 samples are needed to charge C1 to the correct level. However, C1 will discharge slightly through $R_{leakage}$ between samples so the minimum number of samples will be greater than 50.

The voltage available to charge C1 is the difference between E_3 and E_{clamp} .

$$E_{clamp} = \frac{E_3 + e_p \left(\frac{R_1 - R_2}{R_1 + R_2 + 2 R_{diode} D} \right)}{1 + \frac{R_p + R_{source} D}{R_{leakage}}}$$

where:

R_p = the parallel resistance of R1 and R2 in series with the diodes.

$$D = \frac{\text{pulse time off}}{\text{pulse time on}}$$

e_p = differential pulse amplitude.

R_{diode} = the forward resistance of the diodes in series with the pulse source impedance.

$R_{leakage}$ = the parallel leakage of V1, D1 and D2.

clamping-
circuit
equation

The clamping-circuit equation above helps to illustrate many factors in the circuit configuration. When C_1 is charged to the correct level the circuit reaches equilibrium. At equilibrium E_{clamp} will always equal E_3 if:

1. $R_{leakage}$ is high -- a potential problem if leakage or gas current in V_1 is present.
2. R_{source} is low.
3. Pulse off/on ratio is as low as operating conditions will permit.
4. Pulse amplitude is low, making e_p (the differential pulse amplitude) as low as possible.

This particular circuit has design trade-offs of all four conditions. R_1 and R_2 are fairly large (2.2 M Ω) for several reasons:

- (a) Minimum pulse loading.
- (b) The pulse amplitude must be made large to accommodate the increased dynamic range required by the positioning control. The dynamic range is the sum of the maximum composite-video-waveform amplitude and the total positioning voltage. The fraction in the numerator indicates that larger pulse amplitude will result in greater drift sensitivity unless R_1 and R_2 are made fairly large.

low-
frequency
limitation

The principal limitation of this keyed restorer, when used in the vertical system of an oscilloscope, is the removal of 60-Hz hum and other low-frequency errors. The low-frequency limitation requires a DC restorer ON/OFF switch so the restorer circuit can be disabled when the presence and amount of low-frequency signal must be known.

ideal
keyed
fast-clamp
DC restorer

Fig. 5-8 illustrates a keyed fast-clamp DC restorer very closely approaching the ideal circuit of Fig. 5-6. Through the use of semiconductors, most of the design trade-offs of the previous circuit are not necessary. As in the previous two circuits, the composite video waveform is applied to C1 from a low-impedance source -- in this case an emitter follower. Quiescently, the discharge path for C1 is the combined reverse leakage of D1-D2 and the FET gate. D1 and D2 are low-leakage diodes which, combined with the very low reverse leakage of the FET gate, provides a long discharge time constant between keying pulses. D3 and D4 are both 3-volt zener diodes, reverse-biasing D1 and D2 enough to prevent composite video from forward biasing D1 and D2. When a pulse is applied to the primary of the transformer D1 and D2 conduct. The charge path for C1 is now the source impedance of the emitter-follower (emitter resistance) and the forward resistance of D1 and D2 -- providing a charge time constant of about 1 μ s.

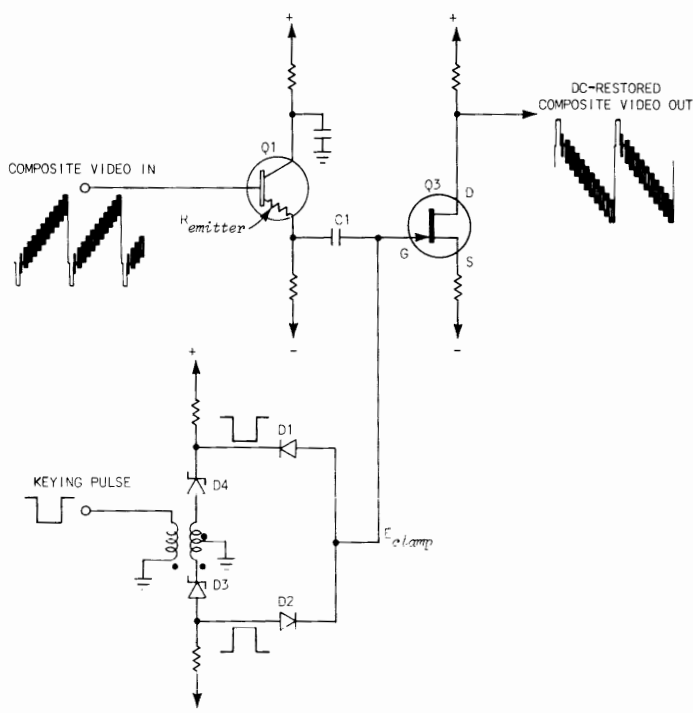


Fig. 5-8. Transistorized keyed clamp circuit--fast restorer.

Except for the use of semiconductors, the circuit configuration of Fig. 5-8 appears to be nearly identical to the circuit shown in Fig. 5-7. Looking at the clamping circuit equation and examining the four main criteria that determine whether E_{clamp} will always equal E_z , several differences between the circuits shown are apparent. Some major differences in operation efficiency and stability may not be quite as apparent.

leakage
resistance

The leakage resistance (discharge path for C1 between samples) is *very* high because:

- (a) The gate-leakage current of the FET is at least an order of magnitude smaller than the control-grid current of a typical tube.
- (b) If the reverse leakage resistance of D1 and D2 are identical, any leakage current through the diodes will cancel and the net current change at the junction of D1 and D2 will be zero. Under these conditions, C1 can be almost any value.
- (c) The total source resistance is low -- consisting of the series forward resistance of D1-D4 in parallel with the series forward resistance of D2-D3 plus the emitter resistance of Q1, for a total of less than 1 k Ω .
- (d) The pulse on/off ratio is about 31 compared to a pulse on/off ratio of 63 for the circuit in Fig. 5-7.
- (e) The keying pulse peak-to-peak amplitude is low, making e_p , the differential pulse amplitude, a low value. Since DC positioning is not done in this circuit, the dynamic range is much smaller, permitting a smaller pulse to be used.

circuit
limitations

The two principal limitations of this circuit are:

1. The differential pulse amplitude e_p . If D3 and D4 do not have the same zener voltage, the error *difference* between the zener voltages will appear at E_{clamp} when D1 and D2 are forward biased. However, the error will be constant.

2. The leakage resistance the principal discharge path for C1 is the leakage resistance of the diodes. The diodes must have radically different leakage characteristics for the discharge to be objectionable, since C1 is 0.001 μ F.

The circuit illustrated in Fig. 5-8 will solidly clamp the composite video without distortion, but any very low-frequency information will be removed. However, CRT displays of composite video at a horizontal line rate will have virtually no vertical jitter. Since the restorer is very fast, signal source switching will not cause annoying jumps on the CRT display.

FEEDBACK DC RESTORER CIRCUITS

four-
diode
keyed
clamp

The deluxe form of the two diode clamp circuits is the four-diode keyed clamp illustrated in Fig. 6-1. This circuit is mainly used for *measuring* rather than setting reference levels.

The circuit is inherently balanced since no center tap is used. Both the AC *and* the DC center-tap point is supplied by the second pair of diodes. The circuit gives rise to a great deal of circuit application versatility -- either as a normal DC restorer or a phase detector. When a pulse is applied to the primary of the transformer, a signal path is formed between point A and B.

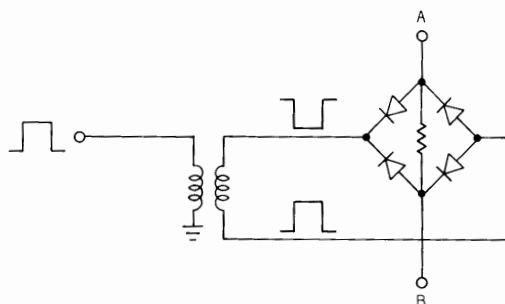


Fig. 6-1. Basic four-diode keyed gate.

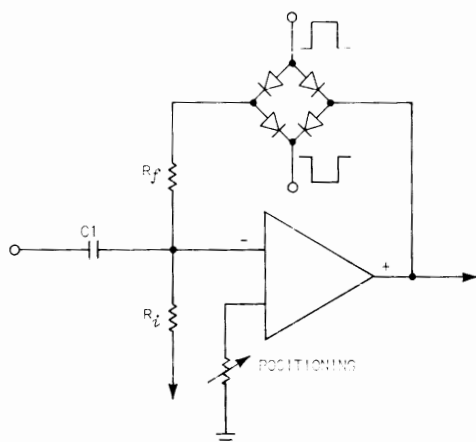


Fig. 6-2. Keyed feedback DC restorer.

keyed-
feedback
network

Fig. 6-2 shows one of the unique applications of the four-diode keyed clamp as a keyed-feedback network. The keyed-feedback network, when applied around an amplifier system, not only establishes a selected portion of the composite video waveform to a reference point, but also provides DC stability to the amplifier system -- eliminating the effects of amplifier drift.

The circuit (Fig. 6-3) has three basic parts:

1. A comparator to measure the difference between a fixed or variable reference voltage and the amplifier output voltage.
2. A diode gate to periodically close the feedback loop.
3. A memory circuit to "remember" the reference point between samples -- when the diode gate is open.

The circuit of Fig. 6-3 can be simplified still further for a preliminary illustration of the circuit principle.

gate
timing

loop gain

Fig. 6-2 best illustrates the concept of the DC-feedback restorer. The pulses that close the feedback loop are arranged to occur when either the sync tip or back porch of the composite video waveform appears at the input. When the gate is closed, the loop gain of the amplifier system is essentially the ratio of R_f/R_i

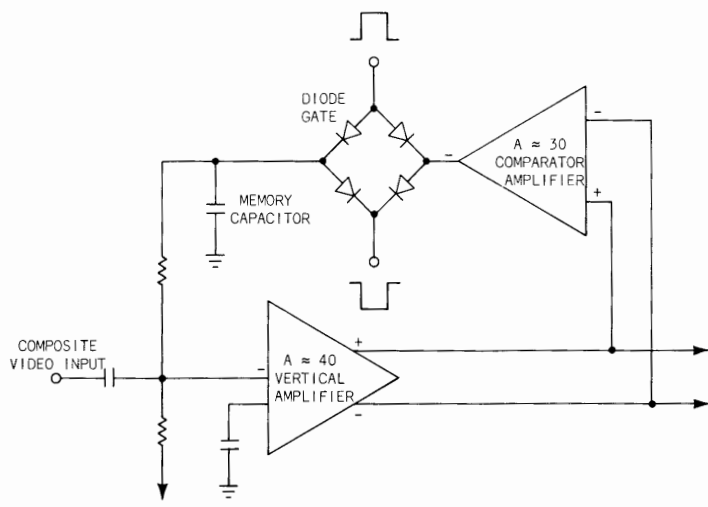


Fig. 6-3. Block diagram of an actual keyed DC feedback restorer.

-- in this case a closed-loop gain of about 1.6. A closed-loop gain of 1.6 is too good to be true, so closer examination of the circuit is necessary. Fig. 6-4 illustrates the equivalent circuit during feedback time. When the gate is closed for feedback purposes C_1 is in parallel with R_i , so the loop gain of 1.6 exists only at, or very near DC. The result is a fast-operating feedback DC restorer with a low-pass filter in the feedback loop, having the equivalent effect on the overall amplifier system of a "slow" restorer. Putting it another way, insufficient feedback current is available to charge C_1 to the correct value in one sample or less.

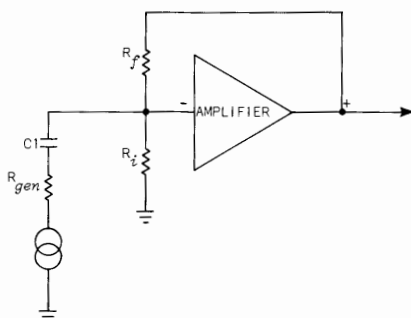


Fig. 6-4. Simplified block of feedback DC restorer - when diode gate is closed.

effective
time
constant

The second consideration is the low frequency 3-dB point "looking" into C_1 . The low frequency 3-dB point is, of course, determined by the RC network consisting of C_1 , R_i , and R_f . R_i and R_f *cannot* be considered in parallel, because the voltage on one end of R_f is moving in the opposite direction to the input signal. If R_i and R_f were considered in parallel, the time constant formed with C_1 would appear to be 0.15 seconds corresponding to a 3-dB point of 1 Hz.

However, R_f is part of the negative feedback loop, so the equivalent R_{f1} is approximately equal to $\frac{R_f}{A}$, where A is the open loop gain. The equivalent R_f in this case will be about 10 k Ω . Since the equivalent R_f is much smaller than R_i the time constant is essentially R_{f1} or about 1.0 ms, corresponding to a 3-dB point of about 150 Hz instead of 1 Hz.

Two important conclusions can be reached at this point when considering low-frequency abnormalities existing on the applied composite video waveform:

1. Under normal conditions, the field sync pulses will not be appreciably distorted but less hum will be displayed on the CRT than may actually exist.
2. The actual low-frequency 3-dB point will be affected by the loop gain of the amplifier system. When the GAIN control is turned to minimum the reduced system gain will *increase* the low-frequency response.

LF
response

The third consideration is the dynamic voltage range of the feedback system. The dynamic range of the feedback amplifier is limited because the theoretical dynamic range is much larger than practical. For example, a 1.0-V change at the input would result in a 400-V change at the output (other end of R_f) if the dynamic range was unlimited and the system open-loop gain was 400.

Therefore, the practical dynamic range of the feedback amplifier is limited because:

1. A diode gate exists in the feedback circuit. The output voltage applied to the sampling gate cannot exceed the gating-pulse voltage amplitude.

2. The absolute voltage level of the sampled portion of the applied composite video is approximately at the same DC level each time the feed back gate is closed.
3. A large supply voltage for unlimited swing is not available or needed.

eliminate
drift

Fig. 6-5 illustrates the circuit details of the DC-feedback-restorer feedback-amplifier system. Since the DC-feedback restorer has two functions -- providing amplifier-system DC stability and establishing a selected portion of the composite video waveform to a reference point -- the circuit operation will be described in terms of the two separate conditions. The first condition is eliminating the effects of amplifier drift. Assume that pulses are applied to the diode bridge whether or not the composite video signal is applied to the amplifier input. Under ideal quiescent conditions, $+E_o$ should equal $-E_o$. If amplifier drift causes $+E_o$ to not equal $-E_o$, the difference voltage applied to

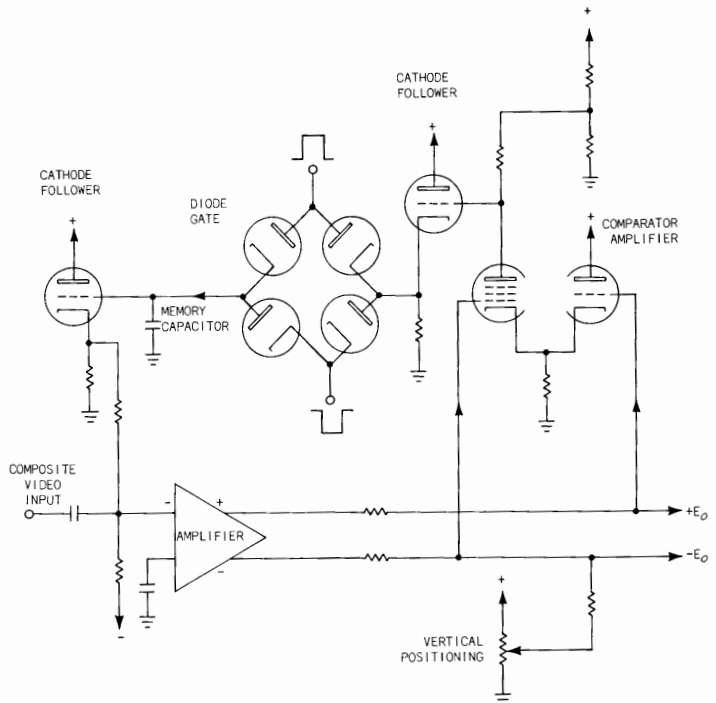


Fig. 6-5. DC feedback restorer.

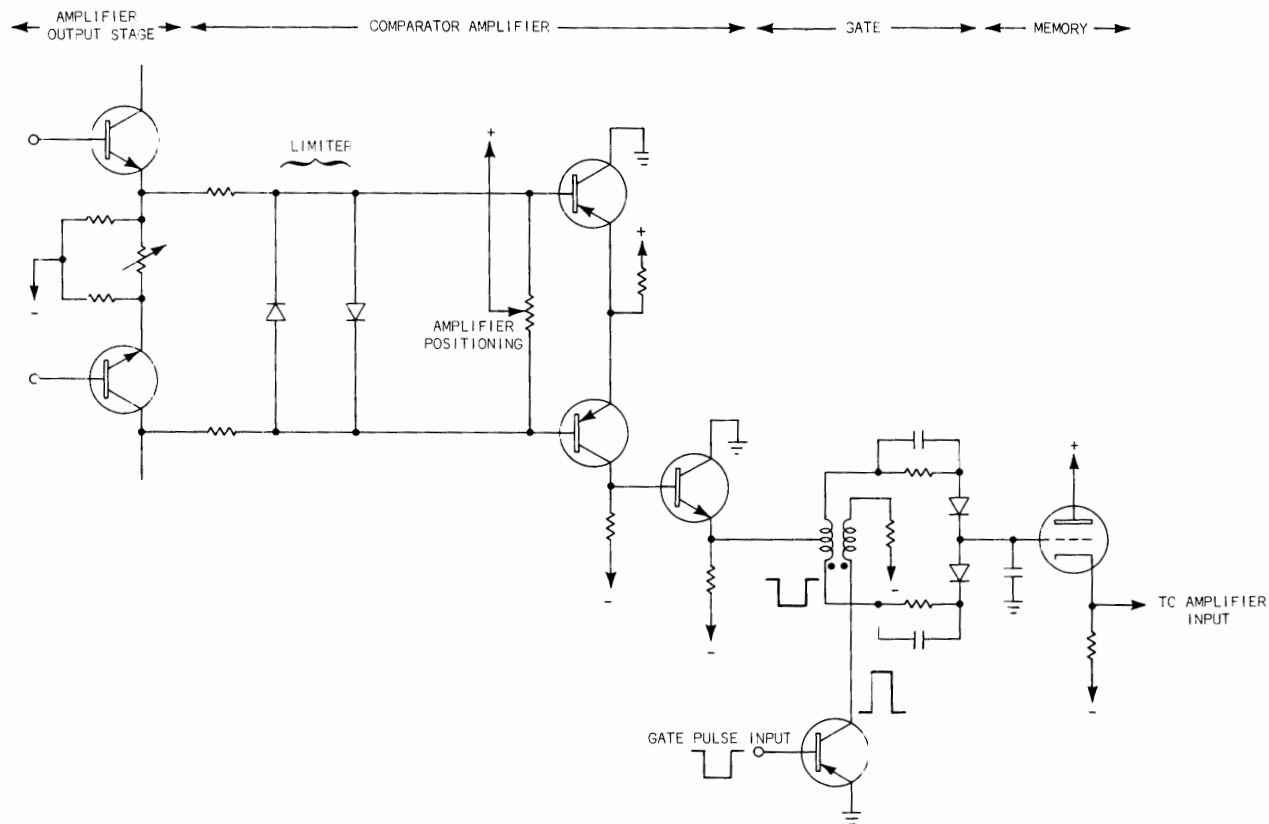


Fig. 6-6. DC-feedback restorer with dynamic-range limiting.

the differential amplifier (V1) will be amplified. The next pulse to close the diode gate will allow the amplified error voltage to charge the memory capacitor C2. Since the current available through the large value R_f is limited, C1 cannot be charged immediately to the correct voltage. However, the amplifier drift is normally a much longer time constant than the time required to charge C1 through R_f so the amplifier is virtually drift free.

reference level
sustained error signal

When a composite video waveform is applied to the amplifier input through C1, the second condition now applies in addition to the first condition just described. The feedback diode gate is closed during a selected portion of the composite video waveform -- in this case the back porch -- and the absolute voltage level of the back porch is compared to the reference voltage at the center arm of the position control. Any difference voltage between the back porch DC level and the positioning control DC level is amplified and applied to the memory. Since C1 and R_f form a low-pass filter to the feedback error signal, the error must exist *continuously* for at least 1.0 millisecond -- the equivalent time constant of the low-pass filter -- before the error will be completely corrected. The time constant of the low-pass filter prevents complete removal of 60-Hz hum that may exist on the composite video waveform.

dynamic range limiting
low-Z drive

The circuitry shown in Fig. 6-6 is very similar to the circuit just described, with a few minor changes. Emphasis has been placed on minimizing the DC effects of color burst that may be present on the back porch of the sync pulse and reducing the effects of overdriving the comparator amplifier by off-screen signals. Since the dynamic range of the comparator amplifier is similar, two of the diodes are replaced with a center-tapped floating transformer winding. The composite video is applied push-pull from the amplifier output stage to the comparator amplifier. Back-to-back diodes in the base circuit of the comparator amplifier limit the dynamic range of the comparator of the feedback system to prevent overdriving the feedback system from unusually large signals. The output of the comparator is from an emitter follower to provide a low-impedance current source to drive the floating transformer winding. Since the feedback restorer is used in a vertical amplifier system, some form of positioning the CRT

display
positioning

display is needed. The DC restorer is a gated negative feedback loop, so positioning can be easily accomplished by introducing an error signal into the comparator amplifier.

Notice in Fig. 6-6 that the necessary positioning is done by applying a differential current to the two inputs (bases) of the comparator amplifier.

With the limited dynamic range of the comparator amplifier, the feedback signal at the center tap of the transformer winding will not forward-bias the sampling diodes. A pulse applied to the primary of the transformer will produce large enough push-pull pulses on the secondary to forward bias the diodes; the pulse current, plus the DC current from the comparator emitter follower combined, then charge the memory capacitor.

7

WAVEFORM SYNCHRONIZATION

sync pulse
purpose

Historically, the use of scan-synchronizing pulses dates back to the initial developments of the electronic television system carried out in the early 1930's. The primary function of the synchronizing pulses is to assure simultaneous vertical and horizontal spot location in both the camera and the receiver. The sync pulses do NOT control the spot's rate-of-change, but only the spot's occurrence (repetition rate).

composite
sync

Nearly one-third of the total peak-to-peak voltage amplitude range of the composite video waveform is used exclusively for synchronizing information. If the remaining two-thirds of the composite video are removed, only a sequence of pulses remain and the waveform is called "composite sync." The composite sync waveform contains three basic pieces of *time* information:

line sync

1. Line sync pulses -- Line sync pulses occur at 15.750-kHz repetition rate to synchronize the horizontal scanning line termination in both the camera and the receiver.

field sync

2. Field sync pulses -- Field sync pulses occur at the beginning of each scanning field at a 60-Hz rate to synchronize the camera and receiver field scanning. The field sync pulse is serrated to simultaneously provide line sync information during the occurrence of the field sync pulse.

frame sync

3. Frame sync pulse -- The frame sync pulse (usually called a field identification or recognition pulse) is a recent innovation necessitated by the use of in-service field blanking interval test signals. The frame sync pulse is derived from the time relationship between the serrated field sync pulse and the subsequent horizontal pulses. The use of frame sync pulses is confined exclusively to waveform monitors.

One basic difference between picture monitors and waveform monitors dictates different circuit requirements in the processing of composite sync. The basic difference is that picture monitors are intended primarily for visual *observation* while waveform monitors are intended primarily for visual *measurements* -- either voltage measurements or time measurements.

picture
monitor

Since the picture monitor is intended for observation only; the cathode-ray spot is time-referenced on two axes, the horizontal axis and the vertical axis, at fixed time rates to trace out a complete picture. The picture information in the form of voltage waveform, is used to Z-axis-modulate the intensity of the spot. The sync pulses, when processed in the picture monitor, terminate the normal scanning motion of the CRT spot.

waveform
monitor
technology

In contrast, the CRT spot in the waveform monitor is usually time-referenced on only one axis, the horizontal; the picture information in the form of a voltage waveform is used to modulate the vertical axis. The horizontal time rate can be made to display all the waveform of one complete frame or a small increment of any one selected line. The wide selection of time rates usually found in a waveform monitor requires that the waveform be displayed on the CRT virtually free of time-jitter. The sync pulses when processed in a waveform monitor can either terminate or initiate the normal horizontal spot motion, but in the process of recovering the line and field sync information care must be taken to:

1. Preserve the leading edge of the sync pulses for jitter-free triggering.
2. Minimize the effects of noise that may be present in the composite video.
3. Prevent crosstalk between field and line sync -- a problem more serious in picture monitors.

The foregoing criteria determine to a large extent the circuits required in waveform monitors.

Recovering specific timing information from composite video is normally accomplished in two steps:

timing
information

1. The first step is amplitude discrimination. The video is amplitude-separated from sync -- commonly referred to as sync separation.
2. The second step is frequency discrimination. Generally, the common method of simultaneously separating the line and field sync is to recover the *leading edge* of the pulses for line sync and recover pulses of a specific *duration* for field sync.

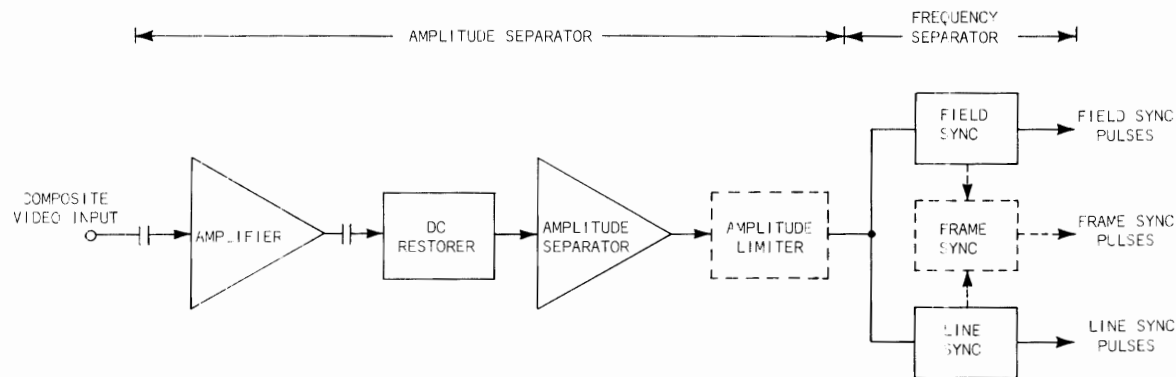


Fig. 8-1. Block diagram of sync separator system.
(Dotted blocks are optional performance items.)

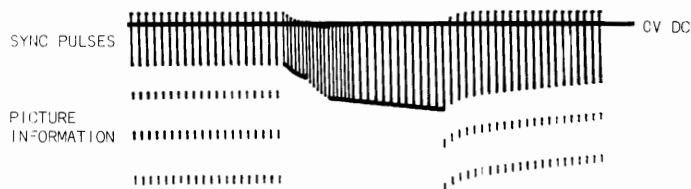


Fig. 8-2. Field synchronizing pulse DC restoration error.

0
0

SYNCHRONIZING PULSE SEPARATION

A complete sync separator system block diagram is shown in Fig. 8-1, but rather than discussing each sync system completely all the amplitude separators in their various configurations will be discussed first and compared for relative performance requirements since they all perform the same function. Then the frequency separators will be discussed according to relative complexity. However, the illustrations will detail the combined amplitude and frequency separators as a system to avoid confusion.

DC restorer
prerequisite

For effective and complete amplitude separation of the sync pulses from video, the composite video waveform must first be DC-restored, otherwise changes in the average picture level may cause clipping of the video to take place *above* the blanking level. Usually, a simple DC restorer (peak detector) is used. Since the simple DC restorer will distort the vertical field sync pulse, the composite video waveform is amplified first to reduce the field sync pulse distortion to a lower percentage of the total composite video amplitude. (See Fig. 8-2.)

Please note the waveforms used to illustrate the sync processing techniques employ composite video and composite sync waveforms at both the line rate and the field rate. Line rate waveforms are used to illustrate a circuit's effect on the individual sync pulses; field-rate sync waveforms are used to illustrate a circuit's effect on the equalizing and serrated field sync pulses.

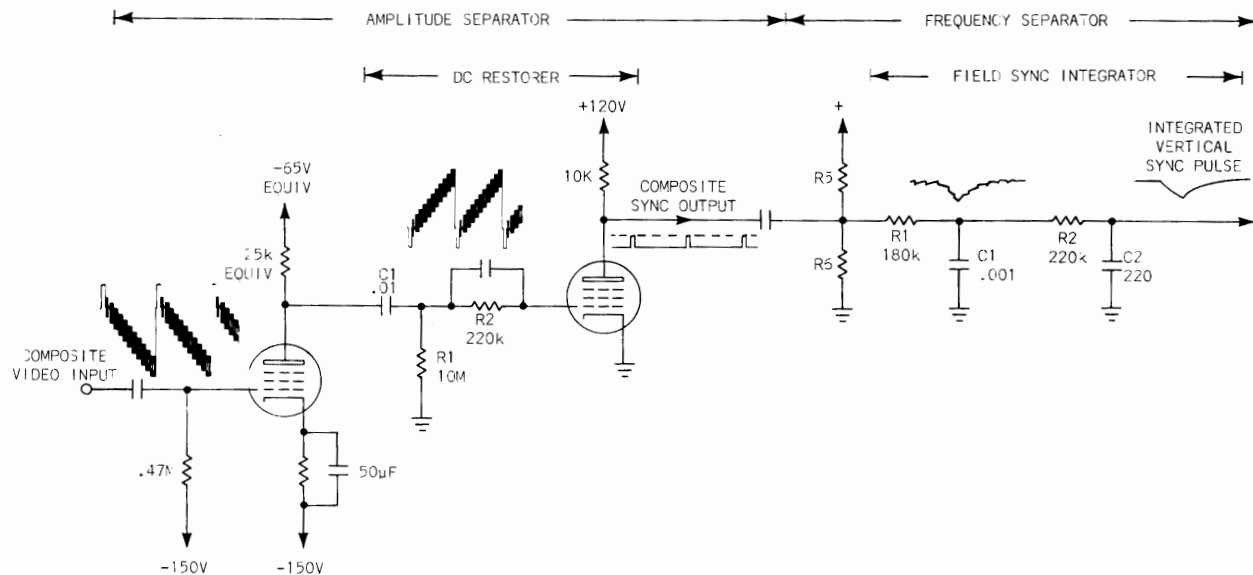


Fig. 8-3. Circuit diagram of complete sync separator system.

Fig. 8-3 illustrates a simple amplitude separator. The composite video is applied to the input grid of V1 to amplify the composite video about 80X and invert the composite video so DC restoration can take place in the grid circuit of V2. R2 is intended to limit the grid current that occurs during sync peaks -- at the expense of a slight loss of DC restoration. The peak-to-peak amplitude of the composite sync video at the grid of V2 is large enough to drive the plate current from saturation to cutoff, effectively removing the active video information.

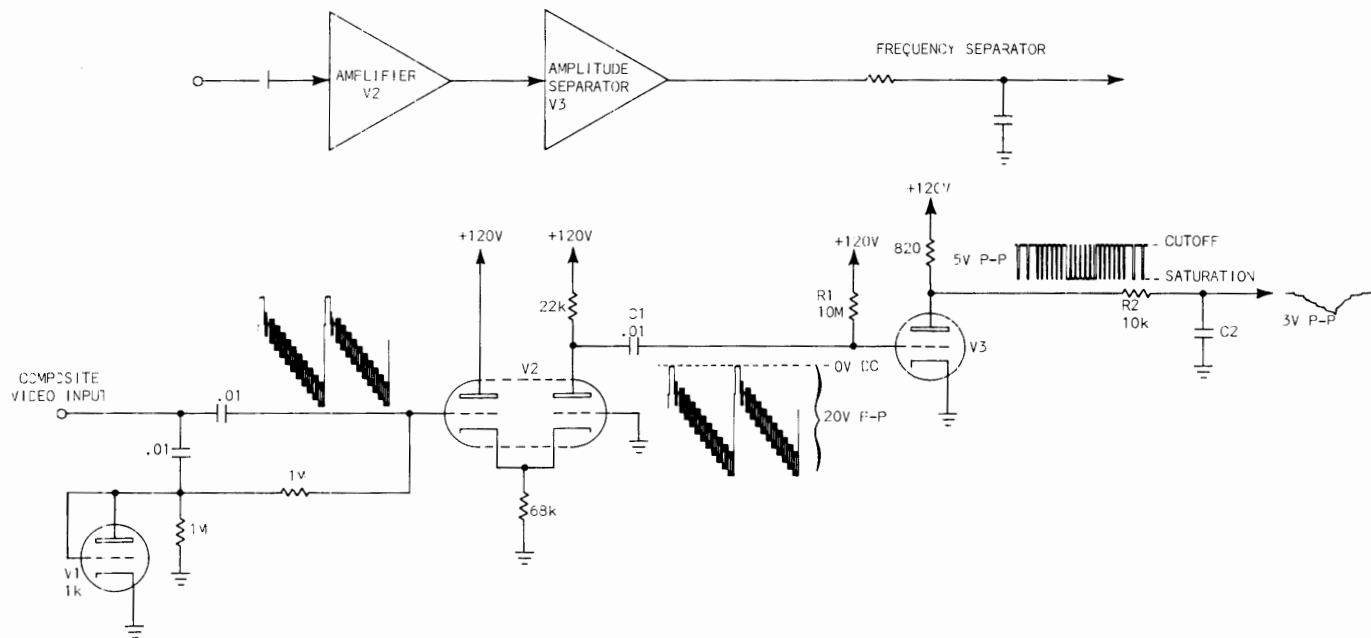


Fig. 8-4. Simple sync separator system.

sync
separators

The circuit in Fig. 8-4 is identical in operation, but slightly different in configuration from the previous circuit. The composite video is applied to the input of V2 cathode-coupled amplifier. Since a wide range of composite video amplitudes are applied to the input grid of V2, large peaks (sync pulses) could cut off the grounded-grid amplifier. Therefore, a bias diode V1 is added to keep the bias of V2 slightly negative -- insuring that sync is always amplified regardless of the total peak-to-peak amplitude of the applied composite video. The composite video is amplified by V2 to about 20 V peak-to-peak (when 4 cm of video are displayed on the CRT).

The grid circuit of V3 is a positive bias DC restorer that clamps the positive-going sync tips at or near ground. A circuit utilizing a relatively high- μ triode is arranged to have a fairly steep plate load-line so a few volts of signal at the control grid will drive the tube from cutoff to saturation. As a result, only a small portion of the sync pulse appearing at the control grid is amplified. The rest of the sync pulse and all the video hold V3 in cutoff, providing complete amplitude separation of the video. (See Fig. 8-5.)

time
jitter

Whenever sync pulses are used for triggering, noise rejection is highly desirable to reduce triggering time jitter. The problem of jitter can be better appreciated by considering the fastest sweep range of a waveform monitor oscilloscope magnified times 25. Under these magnified conditions the sawtooth rate-of-rise is approaching the rate-of-rise of the

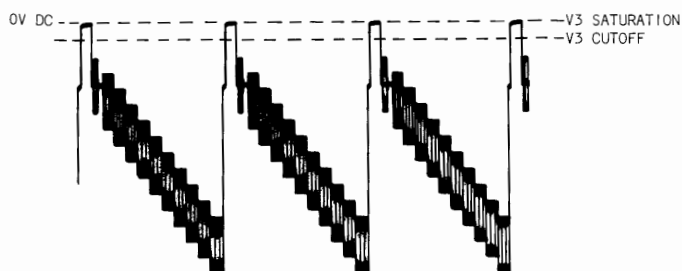


Fig. 8-5. Composite video signal applied to V3 control grid.

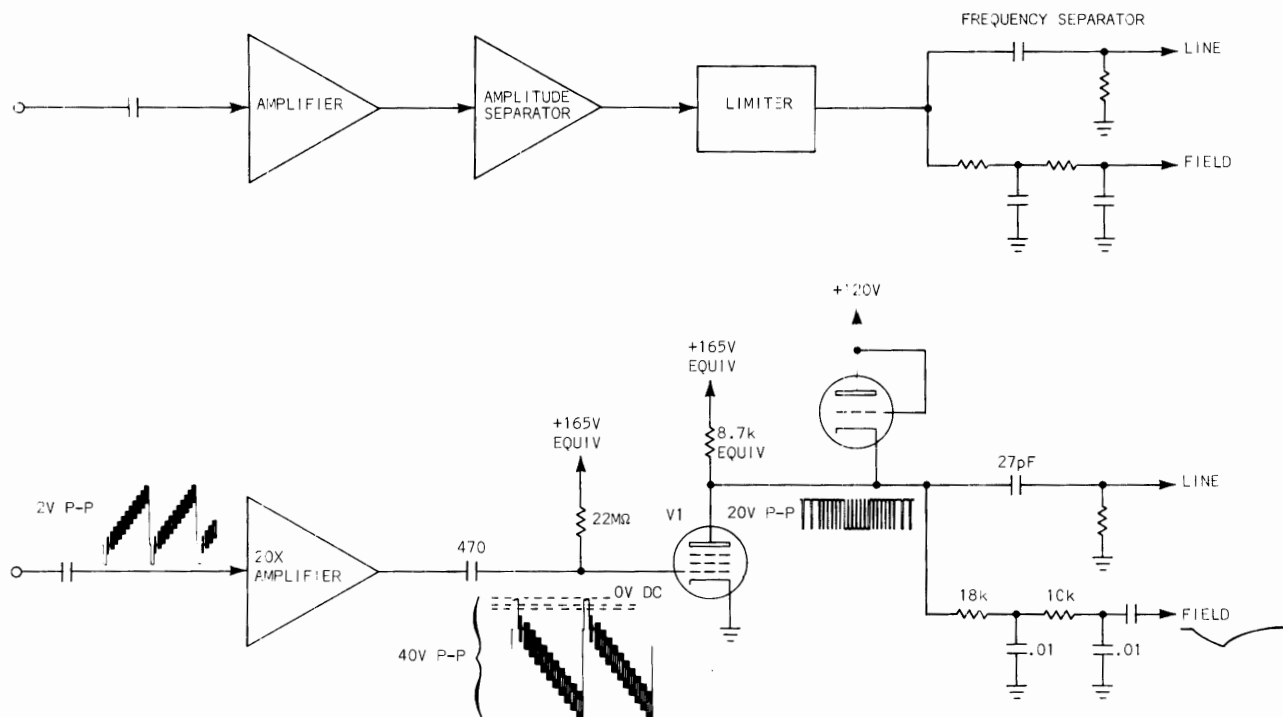


Fig. 8-6. Block and circuit diagram of amplitude-limiting sync separator.

sync pulse leading edge used to trigger the sweep multi. Any time jitter on the sync pulse leading edge will be easily seen on the CRT display under these conditions.

noise

One method used to reduce the time jitter caused from noise on the sync pulse is to "section" the sync pulses. The block diagram and circuit illustrated in Fig. 8-6 shows how the sectioning is done.

First the composite video signal is amplified to about 50 V peak-to-peak before being applied to the positive-bias DC restorer in the grid circuit of V1. Quiescently, the plate of V1 is clamped to +120 V by the diode-connected triode. When negative-going composite video is applied to the control grid of V1, the plate current starts to decrease, but the plate voltage remains constant because of the clamping action of the diode. After the grid voltage goes negative about 2 V, the plate current is reduced further; reverse-biasing the diode. The plate voltage then starts to rise. Finally, when the grid voltage is about -4 V, the tube is cut off and held off until the next sync pulse.

jitter
reduction

Any noise appearing on the sync pulse tips is removed by the clamping action of both the control grid circuit of V1 and the plate-clamping diode. The time jitter on the leading (and trailing) edge of the sync pulses caused by noise is reduced by the high gain of the input amplifier (the amplification makes the effective rate-of-rise of the sync pulse leading edge shorter), since only a small portion of the amplified sync pulse is actually used. The signal appearing at the plate of V1 is then a sectioned portion of the input signal as shown in Fig. 8-7.

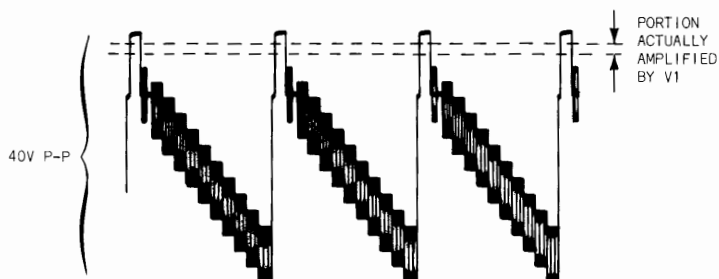


Fig. 8-7. Waveform signal applied to V1 control grid showing portion of sync actually amplified by V1.

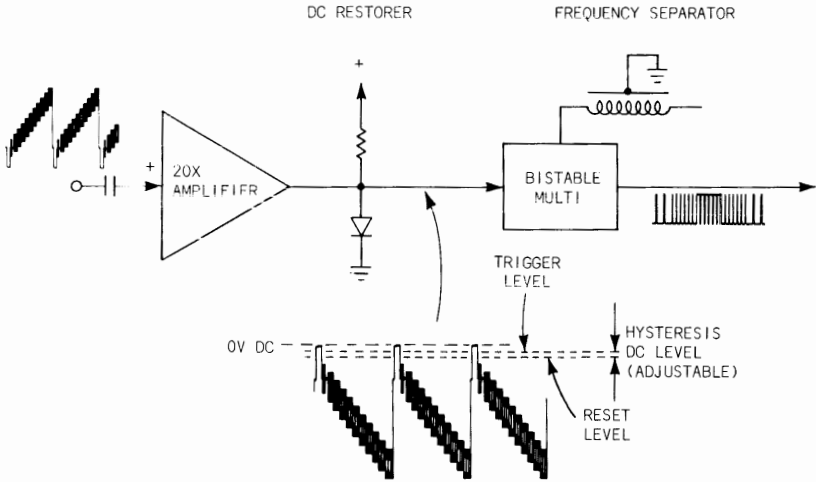


Fig. 8-8. Sync regenerator system.

Another method used to reduce the time jitter of sync pulses is illustrated in Fig. 8-8. Rather than strip the video from composite sync, the DC-restored composite video is DC coupled to a bistable multivibrator, often referred to in this application as a sync pulse regenerator.

sync pulse
regenerator

The principal advantages of using a sync pulse regenerator instead of a sync amplitude separator are:

advantages of
sync pulse
regenerator

1. The time display jitter can be reduced by making the risetime of the regenerated pulses shorter than the incoming sync pulse risetime.
2. The risetime of the regenerated pulses is independent of the signal source risetime.
3. All the regenerated pulses including the field sync pulses are of uniform amplitude.
4. Noise occurring on sync tips can be amplitude-discriminated by adjusting the trigger DC level adjustment to trip (trigger) the multivibrator below the noise level. (See Fig. 8-8.)

circuit
discussion

The composite video is amplified as in the previous circuit and the tips of the sync pulses are clamped to ground with a positive-bias DC restorer. The composite video waveform is then applied directly to the multi. The bistable multivibrator, triggered by the leading edge and reset by the trailing edge of sync pulses, acts as a pulse shaper circuit similar in operation to the bistable pulse shaper of a conventional oscilloscope trigger system. An internal trigger level control adjusts the trip reset DC level to lie in the sync pulse region.

The trigger level control does not need to be a front panel control as in conventional oscilloscopes, because the sync tips are clamped to a fixed DC level by the DC restorer.

To reduce the DC restoration error that normally exists on the serrated field sync pulses (without having to use a keyed-restorer keying pulse source and/or a high-gain amplifier) the DC restoration can be done in the amplifier itself by using a non-linear negative feedback loop.

nonlinear
negative
feedback
clamp

The nonlinear negative feedback loop is arranged to reduce the overall gain close to zero when any negative-going portion of the composite video waveform is applied to the amplifier. The input waveform is capacitively coupled to the amplifier, but the input capacitor is *not* charged to the average DC voltage of the applied signal but instead is charged to the average voltage *plus* the feedback signal, clamping the negative extremity of the signal (sync tips) to a fixed DC level. The simplified diagram is illustrated in Fig. 8-9. The nonlinear feedback loop consists of D1; D2 is the emitter-base junction of a transistor. If the peak amplitude of the applied signal goes too positive, D2 is reverse biased. In effect, D1 and D2 form a dynamic "window" at the input of the amplifier approximately 0.5 V wide. Since the negative peak of any applied voltage waveform is clamped to a preset DC voltage level, the remainder of the waveform within the window will be amplified by the ratio of R_f/Z_k -- about 33 in this case. If the peak-to-peak amplitude of the input signal exceeds 0.5 V, only the portion within the 0.5 V window will be amplified.

dynamic
window

With the concept of operation in mind, look at the simplified circuit illustrated in Fig. 8-10. The composite video waveform is applied to a cathode-follower through C1. The cathode-follower provides current gain to drive the base of the transistor. Quiescently, the negative feedback through R1 will start to charge C1 to the same DC voltage level as the transistor collector until the diode starts to conduct. So initially the DC voltage difference between the input control grid and the output collector is essentially the voltage drop across D1. When an input signal is applied to C1, D1 will conduct even more when the input signal starts to go negative, charging C1 more positive than the average DC level of the signal -- clamping the negative sync tip portion of the waveform to the amplifier quiescent DC level. Since the cathode of the CF is slightly more positive than the grid, any positive-going transition greater than about 0.5 V will completely cut off the transistor.

Composite video, 1.5 V peak-to-peak, is normally applied to C1. The sync tips are negative-going and so are clamped to the quiescent DC level. Since the normal input composite video amplitude is 1.5 V peak-to-peak, the sync pulse amplitude is about 0.5 V peak -- corresponding closely to the window of the amplifier. Thus the output waveform is virtually free of video.

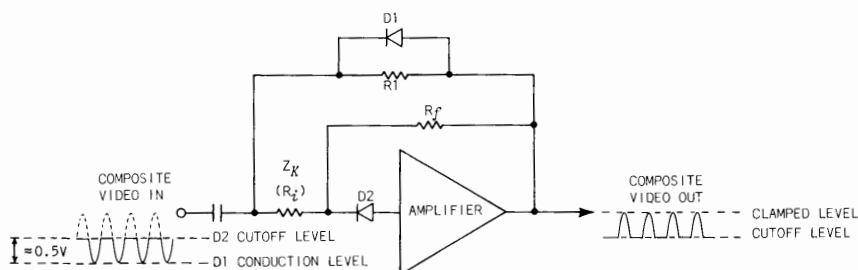


Fig. 8-9. Simplified diagram of a nonlinear negative feedback sync-tip clamp circuit. For purposes of illustration, waveforms shown are sinewaves rather than the composite video.

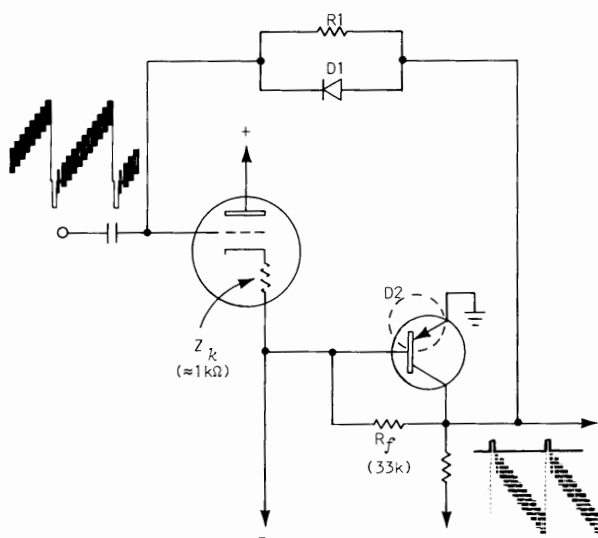


Fig. 8-10. Simplified circuit of the negative feedback clamp.

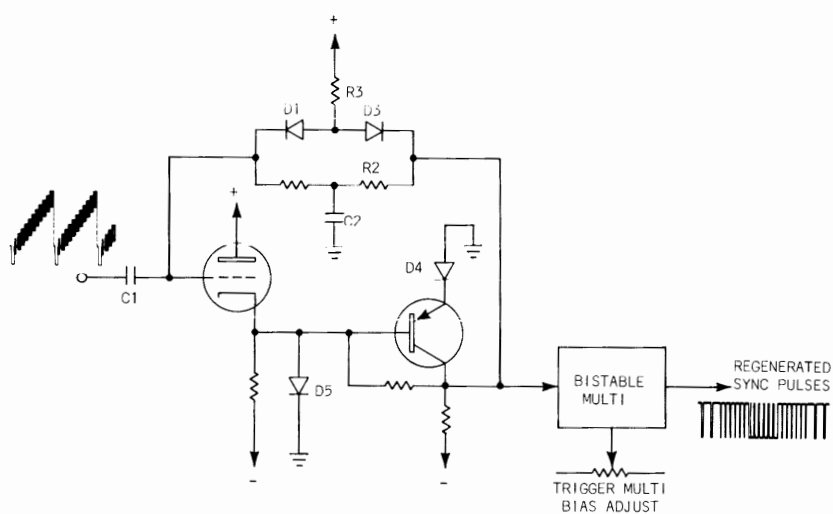


Fig. 8-11. Complete circuit diagram of the negative feedback clamp.

In the actual circuit of Fig. 8-11 several additional components are added to satisfy different operating conditions; the basic operation of the circuit is still the same. These different operating conditions are:

1. Signal overload protection. D3, D4 and D5 are added for very large signal overload protection. Without D3, a large negative-going transition would cause an unusually large positive charge on C1, reverse-biasing the transistor until C1 could slowly discharge through the resistor in parallel with D1. With D3 added, a large negative-going transition will reverse-bias D3 -- limiting the charge on C1. R3 provides biasing for D3. D4 and D5 protect the emitter-base junction from reverse breakdown when large positive-going transitions occur.
2. Low signal amplitudes. The RC network R2 and C2 maintains an average bias on the input control grid when low-level amplitude signals are applied to the amplifier. Without the RC network the clamped level will not be constant -- particularly during the field serrated sync pulses -- when the applied signal is less than 0.5 V peak-to-peak. Normally when D1 conducts the gain is reduced close to zero. With small amplitude signals, the feedback current through D1 is small enough that the forward resistance of the diode increases, changing the voltage drop across the diode, which in turn changes the clamp voltage appearing at the output.

The composite sync waveform is then DC-coupled to a bistable multi, which is triggered and reset by the leading and trailing edge of the sync pulses. In order to stabilize the operation of the bistable multi in the presence of temperature variations, the multi is not permitted to operate in an off-to-saturated mode. Since transistor bistable multivibrators normally have a wider separation between the trip and reset levels, a resistor added between the emitters of the multi will reduce the voltage level difference between the trip and reset levels to a more practical difference -- about 1 V in this circuit. An adjustable DC level control labeled Trig. Multi Bias sets the trip-reset DC level to trigger the multi on the desired portion of the

narrow-
difference
trip and
reset levels

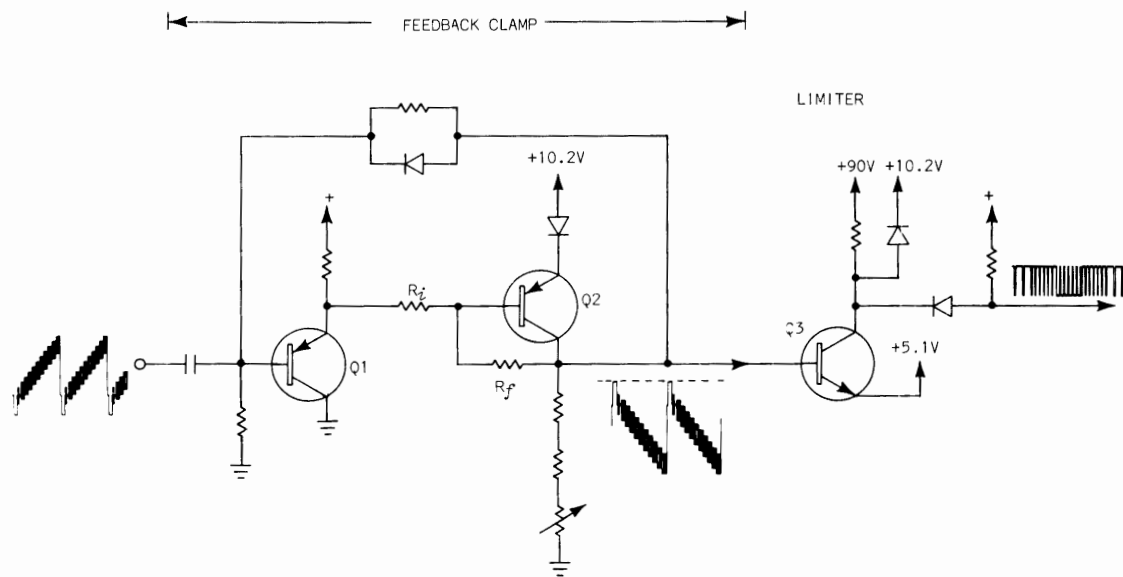


Fig. 8-12. Transistorized feedback clamp circuit.

clamped composite sync waveform. The multi output waveform then consists of regenerated sync pulses completely free of video.

The circuit of Fig. 8-12 is similar in operation to the circuit of Figs. 8-10 and 8-11. The principal differences are:

1. The vacuum-tube cathode follower has been replaced with a transistor emitter follower. Since the emitter impedance of a transistor is much lower than the cathode impedance of a vacuum tube, a current-limiting series resistor R_i has been added between the emitter follower Q1 and the base of the amplifier Q2.
2. The DC clamp level (quiescent level of Q2 collector) has been made variable.
3. Instead of using a bistable multi as a sync regenerator, the composite sync is applied to a limiter Q3 which completely removes the video.
4. The voltage gain of the input "window" is slightly higher than the previous circuit -- essentially the ratio of R_f/R_i . The gain is about 40.

After the sync information has been removed from composite video, the pulses must then be further processed to recover specific timing information -- the line and field timing. All the sync pulses have the same amplitude and the same repetition rate (the pulses occurring during field blanking have twice the line scanning repetition rate for field interlace reasons), but differ in duration or duty factor.

The concepts of this chapter have dealt with the various techniques of sync pulse separation from the composite video waveform. These techniques are determined in large measure by the specific triggering requirements of the monitor oscilloscope. The trigger requirements are in turn dictated by the intended measurement versatility of a specific oscilloscope model.

Reviewing briefly, the two basic objectives in trigger circuit design are:

1. Minimize effects of noise.
2. Minimize triggering jitter.

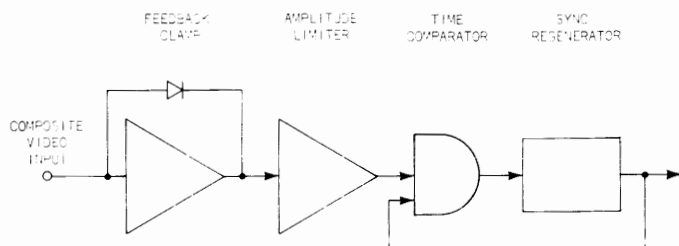


Fig. 8-13. Block diagram of noise and jitter-free sync processing system.

performance
compromises

In the previous circuits one of the two objectives is compromised in favor of the other. The degree of compromise varies according to circuit application. For example, the effects of noise can be reduced if triggering time-jitter can be tolerated. Conversely, minimum triggering jitter can be achieved by tolerating the existence of noise.

The objective of the following circuit is to minimize the effects of impulse transients and random noise that may be present in the composite video sync pulses while still achieving minimum triggering time-jitter. In addition, the circuit must have few adjustments. The result is a sync separation system employing virtually all the sync separation techniques discussed in this chapter -- complete sync pulse separation from composite video as well as sync pulse regeneration. The system block diagram is illustrated in Fig. 8-13. The sync system has three basic parts:

1. Video sync pulse separator
2. Sync pulse regenerator
3. Time comparator to establish a time relationship between the video sync pulses and the regenerated sync pulses.

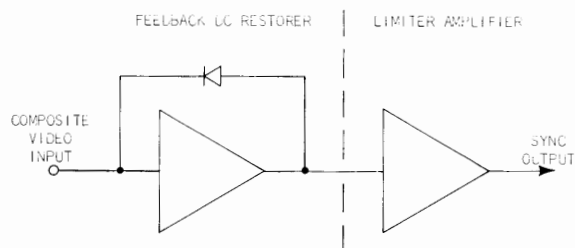


Fig. 8-14. Block diagram of sync separator and limiter.

Fig. 8-14 illustrates the functional block and circuit diagram of the sync pulse separation circuit. The circuit consists of a nonlinear feedback DC restorer and a limiter amplifier.

The DC restorer is functionally identical to the circuit illustrated in Fig. 8-11 but completely transistorized.

feedback
clamp
DC restorer

The DC restorer amplifier clamps the sync pulse tips to a fixed DC level; the limited dynamic range of the DC restorer amplifier is just slightly greater than the anticipated amplitude of the sync portion of the composite video applied to the input, so only the sync pulses (and possibly a small amount of video) will be amplified within the dynamic range of the DC restorer amplifier. The sync pulses are amplified by a factor of 18; the video is clipped and therefore does not appear at the output.

The separated sync pulses are then applied to the limiter amplifier. The purpose is to amplify only the center portion of the sync pulses while not distorting the leading and trailing edges of the original sync pulses. By using a comparator-type amplifier, a limited dynamic range can be utilized without driving the amplifier into saturation.

The circuit details of the sync separator system are illustrated in Fig. 8-15. The waveforms on the circuit diagram are superimposed on dotted replicas of the input waveform to show the portion of the original sync pulse appearing at the output of the circuit.

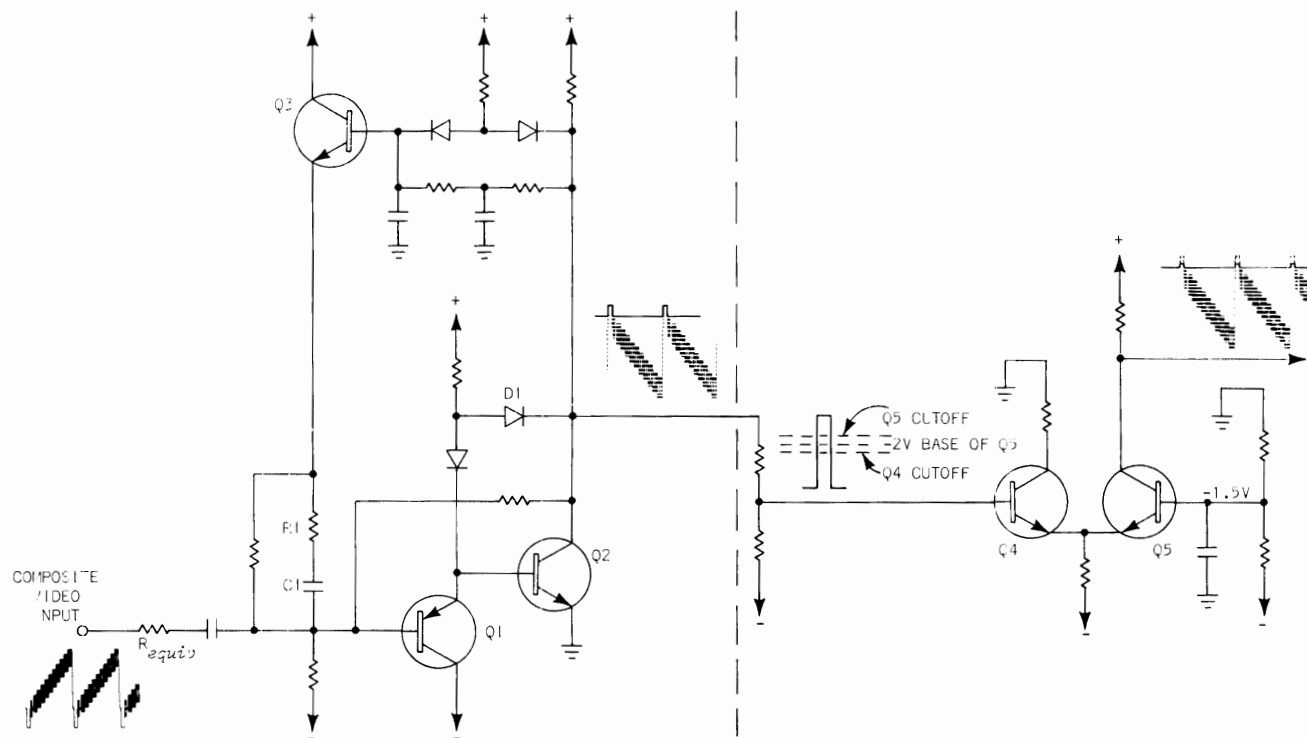


Fig. 8-15. Circuit details of sync separator and limiter amplifier.

The feedback DC restorer concept of operation has been described in the previous circuit but additional details include:

1. D1 prevents the collector-base junction from forward-biasing (saturating).
2. D2 prevents noise impulses that are more negative than sync tips from providing an erroneous feedback error signal.
3. R1-C1 reduces the feedback effectiveness at low frequencies to prevent 60-Hz modulation of composite video from affecting the clamp level.
4. The two networks (R4-C4 and R5-C5) prevent changes in the DC clamp level due to average picture level changes (APL).

The sync pulse amplitude at the output of the DC restorer amplifier is approximately 5 volts peak; the dynamic range of the comparator amplifier is only one volt peak as shown in Fig. 8-15. Therefore, only the center portion of the video sync pulses will be amplified and appear at the output.

requirements
of sync
system

The purpose of the video stripping and sync "sectioning" action of the limiter amplifier is to reduce the possibility of random noise from affecting subsequent processing circuitry. Since random noise, when present, is usually located on the *tips* of sync pulses, much of the noise can be selectively removed by amplitude-limiting techniques.

The effects of large impulses and transients occurring in the sync pulse region of the composite video are reduced with the next portion of the sync pulse system -- the sync pulse regenerator.

The sync regenerator is intended to independently regenerate noise-free sync pulses with a fixed risetime and pulse width, yet maintain a strict time relationship with the composite video sync pulses for jitter-free triggering.

sync pulse
regenerator

Fig. 8-16 illustrates the functional block of the free-running sync pulse regenerator. The sync regenerator is actually a sweep generator system, consisting of a monostable multivibrator and a sawtooth generator. The sweep generator system is similar to the type used in conventional oscilloscopes with two exceptions:

1. Only one time range is available -- corresponding to the television line rate.
2. The sweep generator has *two* feedback loops. The conventional feedback loop controls the total sweep length or duration. The second feedback loop is a gated loop which is used to precisely control the sawtooth rate-of-rise much like the sweep time adjustment control of a conventional oscilloscope.

The normal operation of the sweep generator will be considered first before describing the special gated feedback loop.

The monostable multi is used to turn the sawtooth generator on and off. The sawtooth generator by means of a feedback loop, is used primarily to control the repetition rate of the monostable multivibrator.

The output waveform of the multivibrator is a "gate pulse" corresponding to the +GATE output found on the front panel of most conventional oscilloscopes. The waveform is used not only to turn the sawtooth generator on and off, but also serves as the regenerated sync pulse. The pulse is labeled a +H pulse and used in the rest of the oscilloscope.

The multi circuit is arranged so that an initiating trigger signal is not needed. The sweep gate multi, in its stable state, allows the sawtooth generator to operate. The sawtooth generator output waveform is fed back to the multi to change the multi to its unstable state -- terminating the operation of the sawtooth generator. The monostable multi time constant is arranged to initiate a new cycle shortly after the sawtooth waveform has terminated the previous cycle. Therefore, the sweep system is a recycling or freerunning sweep system.

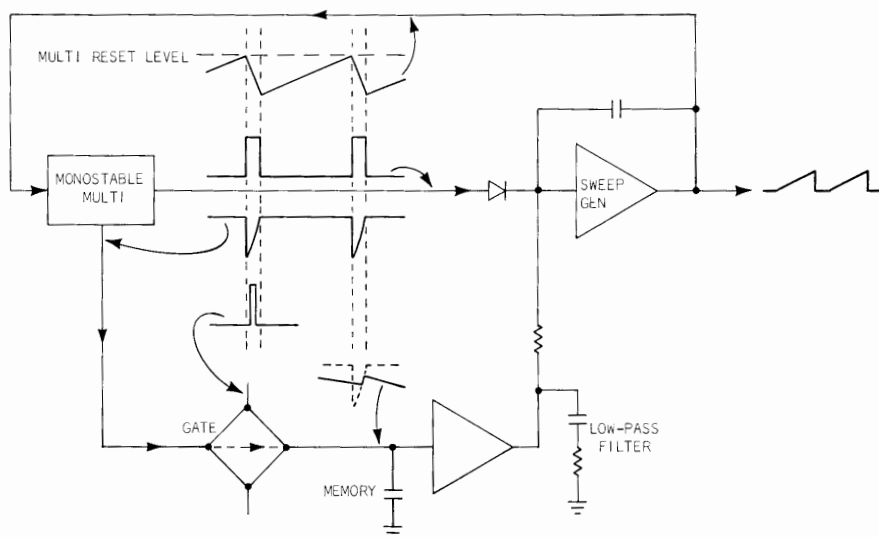


Fig. 8-16. Functional block diagram of sync regenerator system.

The sweep gating multi time-constant determines the time *interval* between the termination of one sweep and the initiation of a subsequent sweep. This interval is fixed.

The rate of rise of any sawtooth generator voltage is determined by the available charging current of a capacitor. That charging current can be controlled several ways:

1. The conventional method of controlling the charging current is by means of a resistor. The resistor is usually made variable to allow manual control of the charging current and usually functions as a Sweep Time Control in a conventional oscilloscope.
2. Another method of varying the charging current of the capacitor is to supply the current from an amplifier. If the amplifier is driven by a feedback loop or other control circuit, the charging current can then be varied automatically.

The sawtooth generator's total rate-of-rise determines the repetition rate or period of the sawtooth cycle. By using the automatic charging current control, the repetition rate can either be held within very close limits or made to coincide with the repetition rate of external repetitive pulses such as the sync pulses. The latter purpose is used in the sync regenerator.

This second feedback loop utilizes the multi output to precisely control the repetition rate of the sawtooth generator.

The feedback loop consists of:

1. A sampling gate to close the loop only during the time the sawtooth generator is *not* running.
2. A memory circuit to "remember" the sample when the loop is open.
3. An amplifier.
4. A low-pass filter to prevent the circuit from overcorrecting. (The corrective action of the feedback loop can be either positive or negative.)

The sampling gate is a "switch" which is closed for a small time interval -- completing the second feedback loop. The switch is closed by the composite video sync pulses, if the sawtooth generator is *not* running. Therefore, time coincidence between the video sync pulse and the multi gate pulse must occur for feedback action to take place.

In order to derive time-error information from the sampled gate waveform the gate is differentiated, forming what is called a "gate ramp." Then when the ramp voltage is sampled, the voltage sample will be proportional to the *time* the sample is made as shown in Fig. 8-16.

Since only one small sample is made during each cycle of sweep operation, a memory circuit is used to maintain or "remember" the sampled voltage level during the rest of the sweep cycle.

The voltage sample is then applied to a comparator amplifier. The comparator amplifier serves as a reference; any difference between the voltage sample and the reference voltage is amplified as an error signal and used to control the sawtooth generator charging current.

The circuit diagram of the sync regenerator system, illustrated in Fig. 8-17, shows the essential elements of circuit operation.

Q1 and Q2 form the monostable multivibrator. Q2 is initially conducting. The base-circuit resistive divider allows Q1 to be forward biased by the sawtooth waveform at the desired amplitude -- 20-V peak in this case. When Q1 is forward biased, the negative transition at the collector of Q1 reverse-biases Q2.

The differentiating network R1-C1 has a time constant arranged to allow the base of Q2 to rise back to a forward-biased condition after 6 μ s. Since R1 is returned to a positive voltage, D1 serves to clamp the positive-going voltage decay at ground potential.

The gate waveform at the collector of Q2 serves two purposes:

1. The wide negative-going portion of the gate waveform reverse biases D4 -- allowing the sawtooth generator to operate.

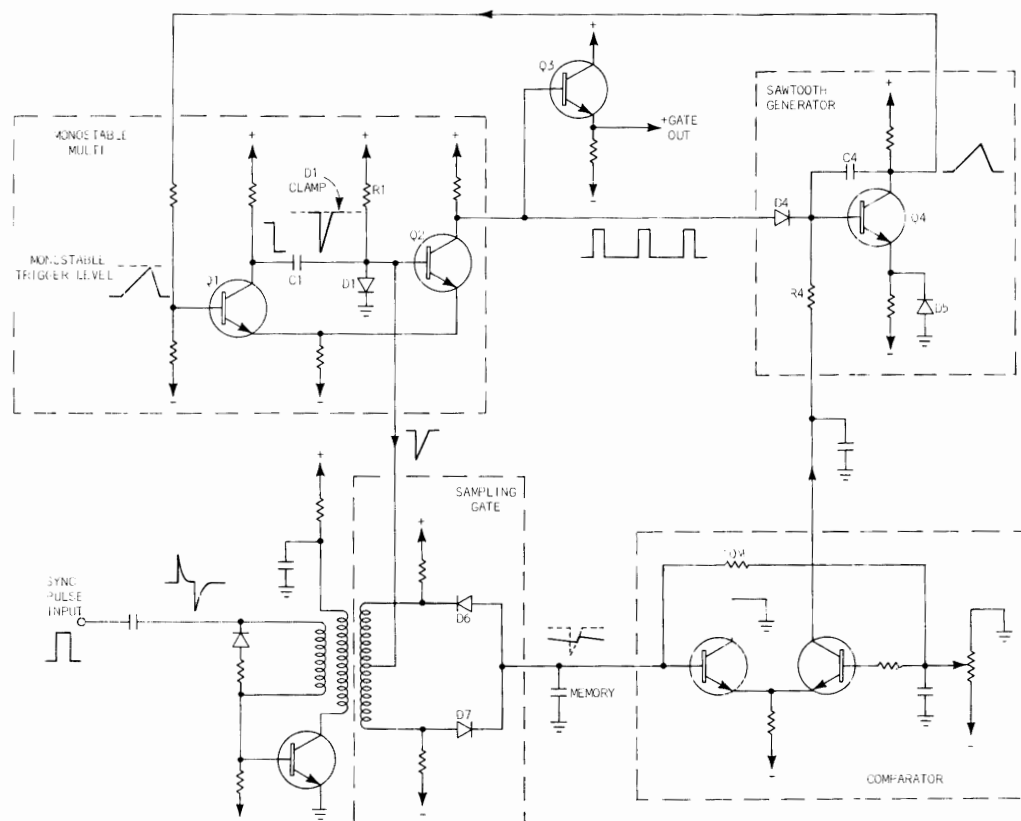


Fig. 8-17. Functional circuit diagram of the sync regenerator.

2. The positive-going portion of the gate waveform serves as the regenerated sync pulse. Q3, an emitter follower, is used to couple the regenerated sync pulses to external circuitry.

The sawtooth generator consists of transistor Q4, timing capacitor C4, clamp diode D4 and charge-current limiting resistor R4.

Instead of grounding the emitter of the sawtooth generator amplifier, the emitter is returned to a negative potential, limiting the operating current. A diode, D5, is used to clamp the emitter near ground.

Since the sawtooth generator operates as a high-gain voltage amplifier, limiting the emitter current eliminates the necessity of selecting transistors for special gain characteristics.

The charging current for the sawtooth generator timing capacitor is supplied from the comparator amplifier consisting of Q6 and Q7. The DC voltage on the base of Q7 determines the collector current and therefore the rate of charge of the timing capacitor.

The DC voltage is made adjustable so the rate of charge can be varied slightly from its nominal rate. The adjustment has two effects:

1. When sync pulses are present, sampling feedback takes place causing the adjustment to determine the time relationship between the leading edge of the regenerated sync pulse and the leading edge of the composite video sync pulse. The adjustment does *not* affect the repetition rate.
2. In the absence of sampled feedback, the adjustment does slightly affect the repetition rate.

The voltage at the base of Q6 is ideally slightly more negative than the base voltage of Q7.

The 10-M Ω resistor between the two bases of Q6 and Q7 is intended to make the base voltage of Q6 more positive in the absence of sampling. The charging current will then be reduced, lowering the sawtooth rate of rise. The repetition rate then becomes about 67 μ s instead of 63.5 μ s. The repetition-rate

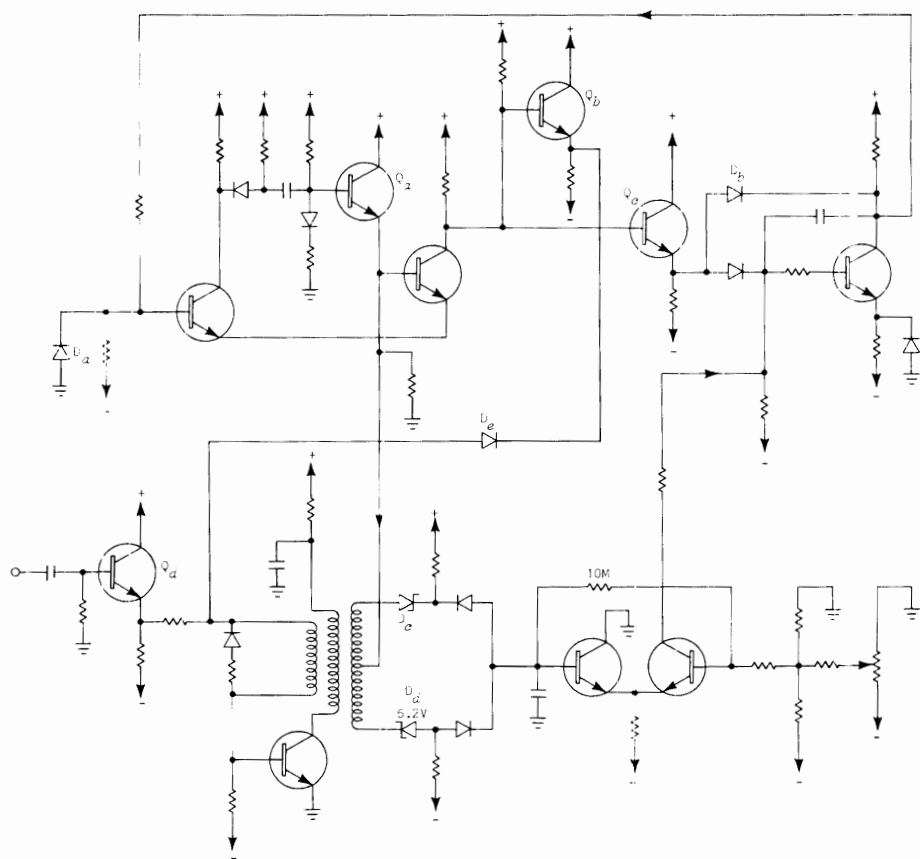


Fig. 8-18. Operational circuit diagram of sync regenerator.

difference between composite sync and regenerated sync will assure eventual time coincidence of the two pulses so sampling *can* take place.

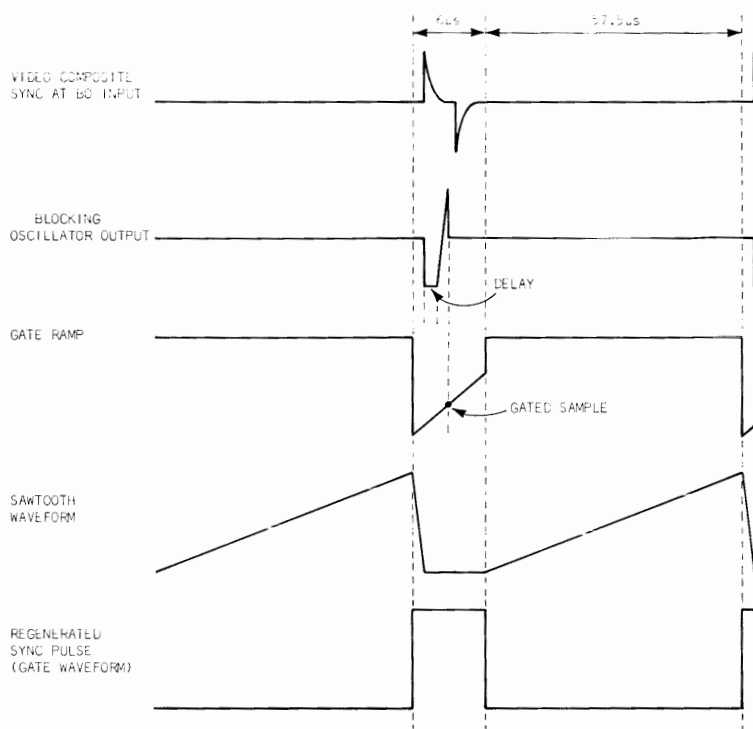
The sampling gate operates the same as a symmetrical four-diode gate except that two of the diodes are replaced by push-pull windings of a transformer. The push-pull transitions do not appear at the output but do forward bias D6 and D7, permitting the common-mode signal -- the differentiated gate waveform -- to appear at the output.

The transformer is part of a blocking oscillator circuit. The blocking oscillator serves two purposes:

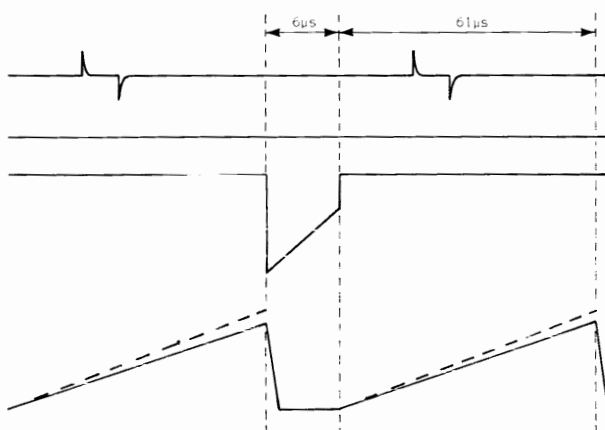
1. Provides a narrow (0.5 μ s) push-pull pulse to operate the gate.
2. Provides the necessary time delay from the leading edge of the sync pulse which is used to trigger the blocking oscillator.

In the complete circuit diagram of Fig. 8-18, several details are added for practical considerations.

1. Four emitter followers, Q_a through Q_d , are added to provide current drive to the various elements without interfering with the basic operation of the system.
2. Diode D_a , clamps the baseline of the fed-back sawtooth near ground to insure a fixed DC level.
3. Diode D_b , clamps the terminated sawtooth voltage to a fixed DC level to insure a constant starting voltage level.
4. When the sampling gate is not operating, D_c and D_d maintain a reverse bias of 6.2 V on the gate diodes, insuring that no portion of the gate ramp will accidentally forward bias the gate diodes.
5. Diode D_e , maintains the input of the blocking oscillator at low impedance, preventing composite video sync pulses from triggering the blocking oscillator while the sawtooth generator is operating. During the sawtooth reset interval D_e is reverse biased by the regenerated sync pulse. If a composite sync pulse occurs during that interval, the sampling gate will be closed, allowing feedback to occur.



(A) SYNC LOCK HAS OCCURRED



(B) BEFORE SYNC LOCK OCCURS

Fig. 8-19. Waveform diagram of time sequence of events in the subcarrier regenerator.

The time sequence of events is shown in the ladder diagram of Fig. 8-19A and 8-19B.

In Fig. 8-19A, video sync pulses are occurring time coincident with the regenerated sync pulses. Gated feedback is then occurring to maintain time coincidence.

In Fig. 8-19B, video sync pulses are *not* occurring time-coincident with the regenerated sync pulses. Therefore D_e is forward biased, preventing video sync from triggering the sampling gate. Since gated feedback is not occurring, the time duration is increased to 67 μ s until coincidence does occur.

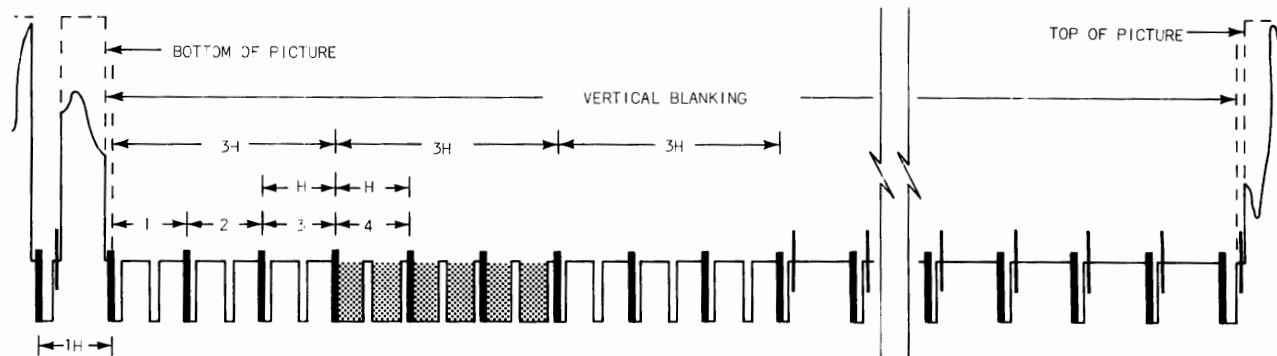
The stability and freedom from time-displacement jitter of the regenerated sync pulses is determined by:

1. The sampling time duration which is ideally as short as possible.
2. The equivalent feedback loop gain.

The sampling time duration is 0.5 μ s to provide enough time to charge the memory capacitor. The feedback loop gain is very high; therefore, a low-pass filter in the collector of the comparator amplifier is added to reduce the AC gain to less than unity. Otherwise, instability and oscillation will occur. In many positive feedback loops, the low-pass filter establishes what is called damping factor or simply K factor.

The sync regenerator system may appear to perform as an AFC system. Conventional AFC systems, however, generally maintain an *average* rather than a strict time relationship to incoming composite video sync pulses.

The sync regenerator system is best described as a gated, time-coincidence AFC system -- AFC action occurs *only* when the regenerated sync pulse and the incoming composite video sync occur simultaneously. In the conventional system, AFC action occurs *until* the two pulses occur simultaneously.



PULSE WIDTHS ARE NOT TO SCALE.

Fig. 9-1. Vertical Blanking portion of the composite video waveform. The heavy lines show the leading edge of the pulses containing line-timing information during Field #1 (odd field). The shaded area is the serrated field sync pulse.

9

SYNCHRONIZING PULSE
PROCESSING

The line timing information is contained in the *leading edge* of all the pulses -- line pulses, alternate equalizing pulses and alternate serrated field sync pulses. (See Fig. 9-1.) The field timing information is contained in the greater energy content or *pulse duration* of the serrated field sync pulses.

The two basic methods of recovering the line and field timing reference are:

line sync
pulse

information in
leading edge

1. Line Sync pulses: Because the line timing information is contained in the leading edge of the pulses, a differentiating circuit whose output is proportional only to the *rate of change* (leading edge) of the pulse and *not* the pulse duration, is commonly used to recover the line timing pulses. These recovered pulses can then be used to trigger the start or termination of the CRT scanning beam or to simply synchronize the repetition rate of the scanning beam.

During equalizing and field pulse portions of the input wave, *two* leading-edge components occur during *each* line-scanning period. The extra leading-edge components occur when the horizontal scanning generator is already active and are, therefore, ignored by the scanning generator. The line pulse repetition rate is

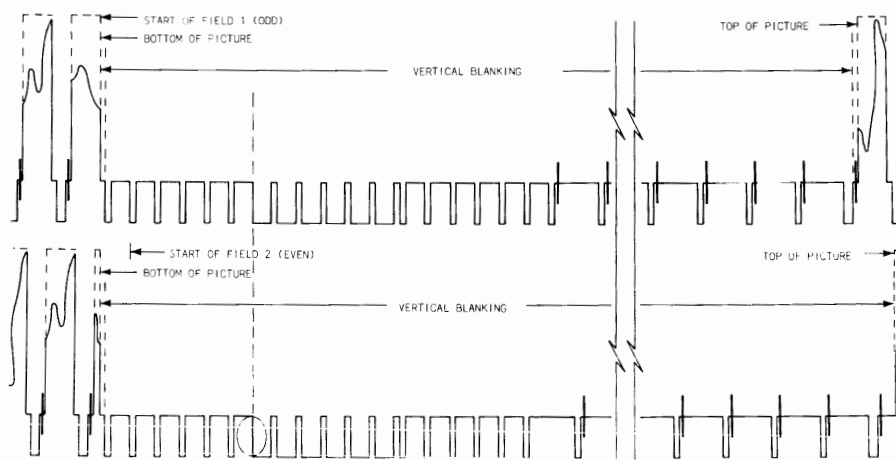
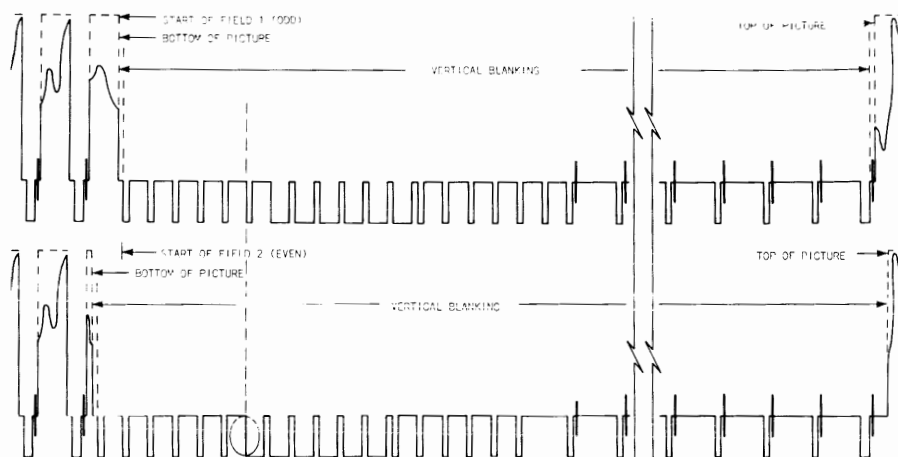


Fig. 9-2. Field One and Field Two.
(Pulse widths are not to scale.)

an odd multiple (525 lines instead of 524 lines) of the serrated field sync pulse, so the extra pulses will be used for line timing on the *next* (interlaced) field. (See Fig. 9-2A and 9-2B.)

field sync
pulse

2. Field Sync pulses: Field sync timing is contained in the pulse duration of the six serrated field sync pulses. The field sync can be recovered from the composite sync waveform by applying the waveform to a circuit whose output voltage is proportional to the *time duration* of the pulses.

An RC integrating network is the most commonly used circuit to recover the field sync information. The time constant of the RC integrator is arranged to be at least equal to the expected pulse duration. The rate of discharge of the integrator is equal to the rate of rise. Therefore, the time constant is usually made longer than the expected pulse duration to minimize the discharge during the serrations.

avoiding
spikes

Referring back to Fig. 8-3, the integrating network consisting of R1-C1 can be readily identified. The time constant is $180\ \mu\text{s}$, considerably longer than the $27\text{-}\mu\text{s}$ duration of each serrated sync pulse, but very close to the $160\ \mu\text{s}$ combined duration of the six serrated sync pulses. The output voltage amplitude then could be expected to be about 63% -- one time constant -- of the amplitude of the applied signal. But the serrations discharge the capacitor 16% of the time, so the output voltage will be less than half the amplitude of the applied composite sync. R2-C2 forms a second integrator with a time constant of $48\ \mu\text{s}$ -- almost ten times longer than the width of the horizontal sync pulses. The purpose of the second integrator is to smooth the integrated field sync pulse by removing the charge and discharging spikes caused by the line sync pulses and the serrations. The integrated field sync pulse can then be applied to a sweep multi to initiate a sweep or unblanking generator at a field rate.

A similar field pulse integrator can be seen in Fig. 8-4 consisting of R2-C2. The time constant in this case is $100\ \mu\text{s}$ -- slightly shorter than the $162\text{-}\mu\text{s}$ duration of the six serrated field sync pulses, so that the output voltage will be slightly greater than 63% of the applied voltage amplitude.

cascaded
integrator

Fig. 8-6 illustrates a cascaded field sync pulse integrator. The time constant of the first section equals the total time duration of the field sync pulses, so the output voltage amplitude will be about 63% of the input voltage minus the 16% that the serrated pulses are discharging the network to less than half the applied voltage. The second integrating network will further reduce the amplitude, but will also remove the discharging serrations leaving a smooth integrated pulse.

Field sync pulses can be formed with very economical and simple pulse-integrating circuits. The pulses formed with these integrating circuits provide adequate triggering or synchronizing information to operate a sweep scanning generator at a 60-Hz or 16.6-ms repetition rate. However, even with carefully selected time constants slight time variations between the recovered field sync pulses can occur because of a slight initial charge on the integrating network from noise, transients, and other voltage impairments. Stating the problem another way, the repetition rate of the recovered field sync pulses at the output of the integrator network is not precisely the same as the originally-generated field sync pulse repetition rate. In practice, the repetition rate error is small (less than 1 μ s in 16.6 ms) and not normally noticeable when the sweep scanning generator is used to displace a spot from the left side of the CRT to the right side in 16.6 ms (one field) or 33 ms (one frame). But if that same sweep is magnified to displace the CRT spot from left to right in 3 μ s instead of 16.6 ms, a 1 μ s repetition rate error in the pulse used to trigger the sweep causes very noticeable display jitter on the CRT. The sawtooth *in a practical circuit* is generally not "magnified" directly, but instead a second sweep generator with a rate-of-rise on the order of 0.3 μ s/div is triggered by the slower sweep generator and called "line selection."

One method used to eliminate the possible time jitter encountered with integrator networks is to replace the circuit. A simple open-circuited transmission line will provide an output voltage amplitude that is related to the pulse time duration and yet maintain a sharp leading edge suitable for jitter free triggering. Recall that a positive-(or negative) going transition driven into an open-circuited transmission line will travel down the line and be

pulse
reflections

reflected back toward the source. If the losses of the line are small, the amplitude of the reflected transition will be the same as the initial transition and furthermore, the reflected transition will also be positive-going. If the initial pulse is still present when the reflected pulse arrives back at the source, the amplitude will be doubled, permitting, in this application, selectivity of the wider serrated field sync pulses.

Consider the simplified circuit of Fig. 9-3. The available current is limited by R_{gen} ; a pulse from the generator will develop a voltage across R_L . When the delay line with an impedance Z_L made equal to R_L is connected as in Fig. 9-4, the voltage amplitude appearing at the output will be one-half the previous voltage (because of the limited current) until the pulse traveling down the delay line is reflected back

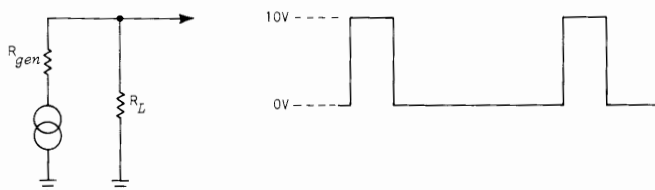


Fig. 9-3. Simplified diagram of a pulse source driving a resistive load.

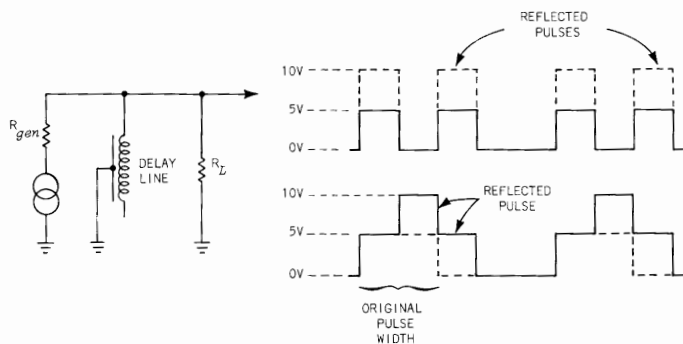


Fig. 9-4. Simplified diagram of a pulse source driving an unterminated delay line.

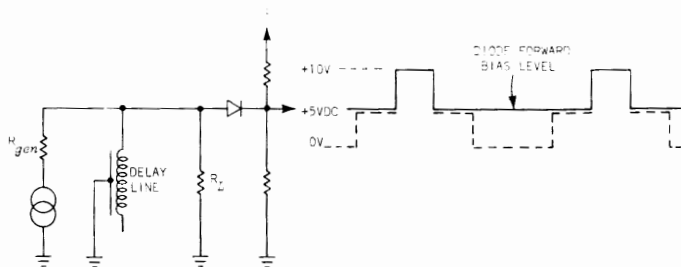


Fig. 9-5. Reflected-pulse amplitude discriminator.

diode pulse
discrimina-
tion

to the source. If the original pulse is still present, the reflected pulse will add to the original pulse amplitude. By adding a diode biased for amplitude discrimination, as shown in Fig. 9-5, the leading edge from the wider field pulses can be selectively removed. If the transmission line is made long enough so the reflected horizontal sync pulse occurs *after* the original pulse has been terminated, two time-displaced pulses will be observed as shown in Fig. 9-4. When a longer duration pulse is applied to the delay line, the reflected pulse arrives back at the source *before* the original pulse is terminated, doubling the amplitude of the voltage step. Since D1 will only conduct when the pulse amplitude exceeds the diode bias only the long-duration field sync pulses appear at the output. The pulses can be differentiated; the first differentiated pulse to occur can then be used to trigger a sweep generator -- with virtually no time jitter.

The actual circuit is illustrated in Fig. 9-6. The pulse generator consists of a bistable multi, triggered and reset by the leading and trailing edges of the horizontal sync pulses. When the selector switch is turned to the LINE position, the delay line is disconnected and regenerated sync pulses (composite sync) 15 V in amplitude appear at the multi output.

The voltage on the cathode of D1 is set by R3 and R4 so that only the upper portion of the sync pulses cause D1 to conduct. Only about half of the voltage amplitude (about 7 V P-P) appears at the differentiator. When the selector switch is turned to FIELD, a 3.3- μ s delay line with an impedance equal to R_L is added in parallel to R_L . Since the plate current of V1 is limited by R_k , the peak-to-peak voltage amplitude of the sync pulses at the plate of

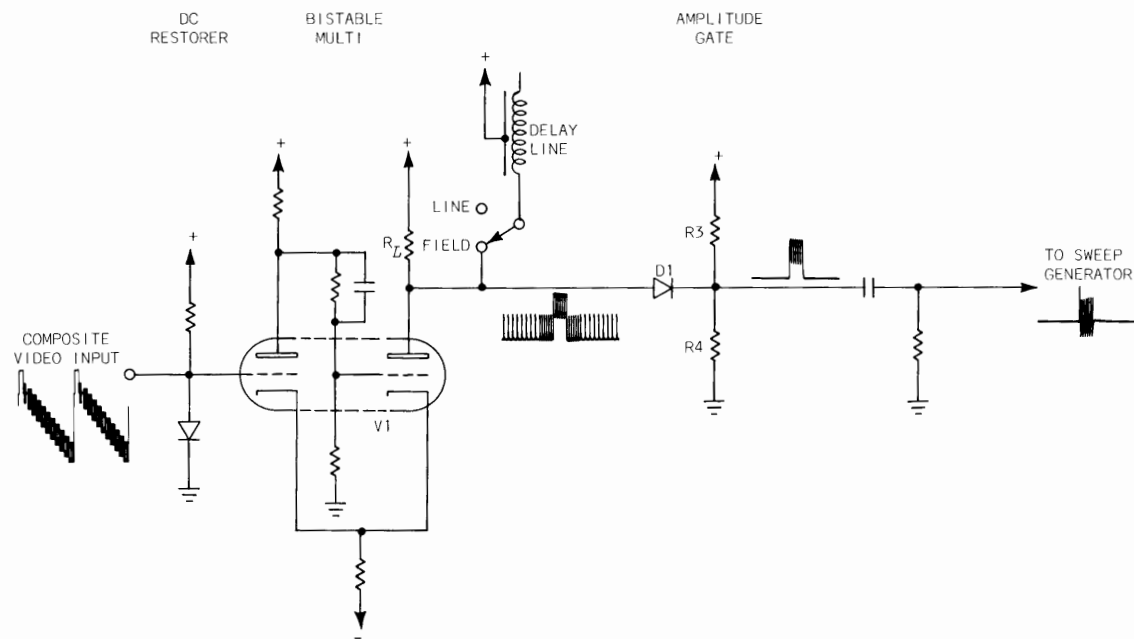


Fig. 9-6. Pulse-duration discrimination circuit.

V1 will be reduced to one-half when the delay line is connected in the circuit. $6.6 \mu\text{s}$ later the pulse driven down the delay line is reflected back and appears at the plate of V1. If the original pulse was a standard $4.5\text{-}\mu\text{s}$ line sync pulse, the reflected pulse will occur time-displace $6.6 \mu\text{s}$ from the leading edge of the original pulse at the same amplitude and D1 will not conduct. When a $27\text{-}\mu\text{s}$ field sync pulse is driven down the delay line, the reflected pulse still appears back at the plate of V1 $6.6 \mu\text{s}$ later. Since the original pulse still is present at the plate, the reflected pulse amplitude is added, doubling the pulse amplitude and forward-biasing D1. The pulse is then differentiated, applied to the sweep generator and a sawtooth is generated -- $6.6 \mu\text{s}$ after the first field sync pulse actually occurs. The other five field sync pulses will also appear at the differentiating network, but are not used.

Using the reflected pulse technique instead of a conventional integrator is more versatile in that the equalizing pulses are not needed to minimize the initial charge on the integrating network. Since waveform monitor oscilloscopes are also used in foreign markets where equalizing pulses are not always a component part of the composite video, providing line selection free of jitter is somewhat more practical.

An alternate method, eliminating the need of a transmission line, is to differentiate rather than integrate the serrated field sync pulses and use the *serrations* for field timing instead of the wide field sync pulses themselves. As illustrated in Fig. 9-7, the composite sync is applied to an RC differentiating network with a time constant arranged to completely differentiate *only* the wider field sync pulses. By differentiating only the wider field sync pulses, the difference in the average DC level that normally exists between field sync pulses and line pulses is eliminated. The result is line sync pulses positive-going from ground and field sync pulse serrations negative-going from ground.

pulse
discrim-
ination
and
processing

Further pulse shaping takes place by connecting the differentiating network to the base of a transistor as illustrated in Fig. 9-8. A by-passed voltage divider sets the emitter to -4 V so the transistor is initially reverse-biased until the negative-going 7-V field serration pulses drive the transistor into saturation.

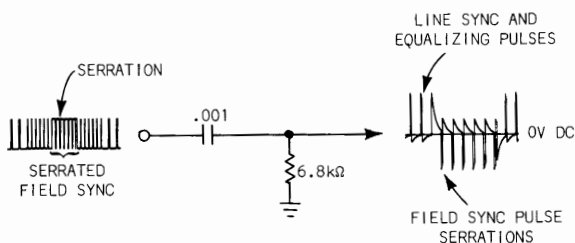


Fig. 9-7. Differentiation network utilized for field sync discrimination.

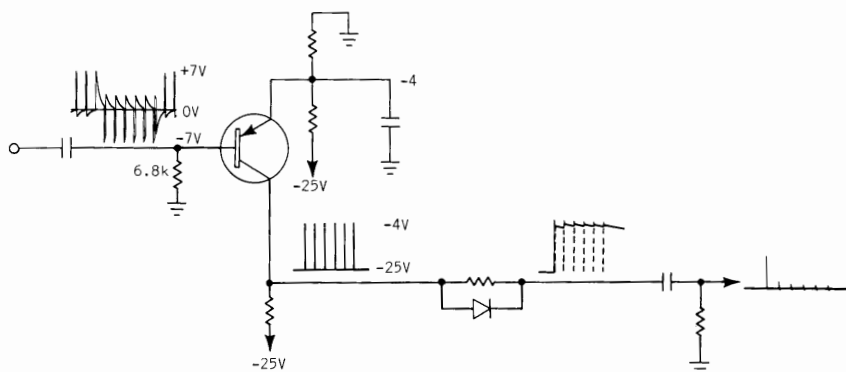


Fig. 9-8. Differentiation network applied to an amplitude-discriminating amplifier.

Finally, the serration pulses are differentiated (before, the *field* pulses were differentiated) to provide only positive-going output spikes. A diode added in series with the differentiating network provides two additional features:

1. Only the positive-going transitions are coupled through the capacitor, because the diode is reverse-biased during the negative-going transition.
2. Since the negative transition can only discharge through the bias resistor wired in parallel with the diode, each succeeding pulse is reduced in amplitude.

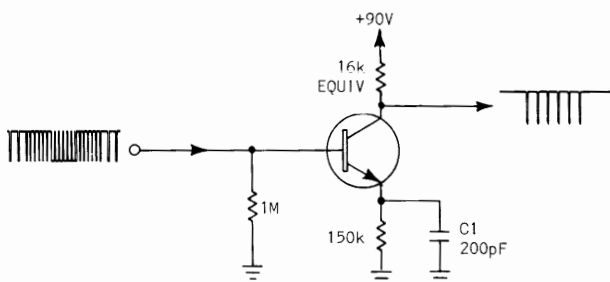


Fig. 9-9. Differentiation concept used in the form of a high-pass amplifier.

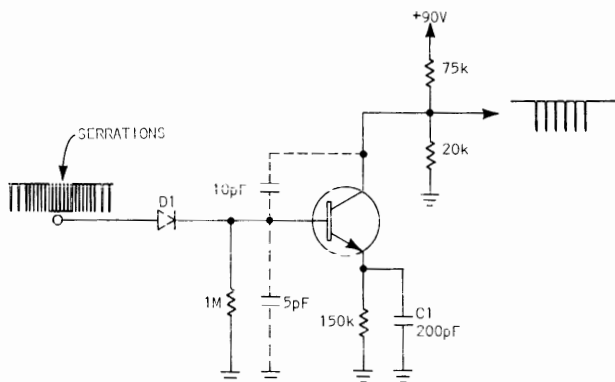


Fig. 9-10. Actual circuit of high-pass amplifier differentiating and amplifying the field sync pulse serrations.

The final single output spike is time-coincident with the trailing edge of the first of the six field sync pulses. The leading edge of the output spike will be precisely the same repetition rate as the serrated field sync pulses of the applied composite sync.

Another circuit utilizing the differentiation concept to process the field sync pulse serrations is illustrated in Fig. 9-9.

The circuit is essentially a nonlinear amplifier where:

1. The gain at high frequencies greatly exceeds the gain at DC and low frequencies. The average DC voltage of the field sync pulses is much different than the average DC voltage of the line sync pulses. However, when composite sync is applied to the amplifier, the line and field sync pulse serrations will form a common average DC level, because the DC gain of the amplifier is very low.
2. The amplifier is quiescently operating very near cutoff, because of the limited current through the large emitter resistor. The common average DC level of the pulses forms the operating bias of the amplifier, which is very near cutoff. The negative-going line sync pulses will cut the transistor completely off and the positive-going serration pulses will turn the transistor on -- producing the desired field serration pulses at the output.

The output field serration pulses would be rectangular like the input pulses except that C1, the emitter by-pass capacitor, charges when the transistor is turned on. C1 will charge at a rate determined by the emitter impedance -- 1.5 μ s, biasing the transistor back to the quiescent level. The output waveform then, instead of being rectangular in shape will be a negative-going spike with a 1.5- μ s decay, corresponding to the charge time of C1. In the actual circuit of Fig. 9-10, about 15 pF of circuit stray capacitance must also be charged and discharged so a series diode D1 and biasing resistor are added in the base circuit to make the stray capacity discharge time much longer than the time required to charge C1. Without D1, the output pulse decay would have a double-slope decay.

extending
discharge
time

The final timing pulse to be derived from the composite sync waveform is the frame sync pulse.

Normally the picture information occurring on any numbered line of one field is essentially the same as the picture information occurring on the corresponding line of the subsequent interlaced field. However, if the picture information on any numbered line of one field were different than the information contained on the corresponding line of the interlaced field, separation of the two fields would be necessary. The need for some method of selecting the field containing the desired information and rejecting the alternate field would be apparent.

VIT

During a live picture transmission, measurements of electrical parameters affecting picture quality are limited to amplitude measurements only. To overcome this limitation, four scanning lines -- two lines during each field blanking interval -- have been assigned for the use of suitable test signals, to permit continuous monitoring of electrical parameters that directly affect picture quality. The use of *different* test signals on *each* field is possible only if a means of identifying, selecting, and displaying the desired field is available. The frame sync pulse, a pulse derived from a specific time relationship between field sync pulses and line sync pulses, is used to identify one of the two fields. The frame sync information is *not* included in the composite sync waveform as a discrete pulse.

Since the use and concept of the term "frame" is not universal, the frame pulse will be referred to as a field-identification pulse. (The concept of the term "frame" in some foreign television systems is directly borrowed from cinematography and is synonymous with the concept of "field" in the United States.) Before field identification can take place the characteristics of each field must be properly defined.

Field ONE, also called the "odd" field, is identified by the following characteristics (see Fig. 9-11):

field ONE
identification

1. The start of field ONE is preceded by a *whole scanning line interval* between the first equalizing pulse and the last line sync pulse of the preceding field.

2. The leading edge of the first serrated field sync pulse is time-coincident with the leading edge of line sync pulses.
3. The first active scanning line normally starts at the top left corner of the raster.
4. The last active line of the field normally terminates at the bottom *center* of the raster.
5. The scanning lines are numbered in consecutive order beginning with the *first* equalizing pulse.
6. The first line sync pulse of the field occurs $1/2$ scanning line ($31.75 \mu s$) after the last equalizing pulses.

Field TWO, also called the "even" field is identified by: (See Fig. 9-12.)

field TWO
identification

1. The start of field TWO is preceded by a $1/2$ scanning line interval between the first equalizing pulse and the last line sync pulse of the preceding field.
2. The leading edge of the first serrated field sync pulse occurs midway between the leading edge of successive line (equalizing) sync pulse.
3. The first active scanning line of the field normally starts at the top center of the raster.
4. The last active line of the field normally end at the bottom right corner of the raster.
5. The scanning lines are numbered in consecutive order beginning with the *second* equalizing pulse.
6. The first line sync pulse of the field occurs a whole scanning line after the last equalizing pulse.

Six principal differences exist between field ONE and field TWO, but one, the time occurrence of the first horizontal line sync pulse of the field, is the easiest to identify electronically.

On field ONE, the first horizontal line sync pulse of the field occurs exactly six scanning lines after the leading edge of the first serrated field sync pulse, whereas on field TWO the first line sync pulse occurs six and one-half lines after the leading edge of the first serrated field pulse.

A gate circuit can be arranged to open exactly six scanning lines after the leading edge of the first serrated field pulse to allow any pulse present at that time to go through the gate. Identification of field one is then possible because a line sync pulse does occur six scanning lines after the start of field sync. Absence of an output pulse would, of course, indicate field two.

forming
field-
identi-
fication
pulse

One method used to form the field identification pulse is illustrated in Fig. 9-13. Basically, two coincident pulses -- one, a generated negative-going delayed pulse and the other, the desired line sync pulse, also negative-going -- are added together. The combined amplitude of the two pulses forward-biases a diode gate, producing the desired identification pulse. The delayed pulse occurs every field, six lines (381 μ s) after the start of the field sync pulse. The desired sync pulse occurs [every alternate field (field ONE)] six lines after the start of the field sync pulse.

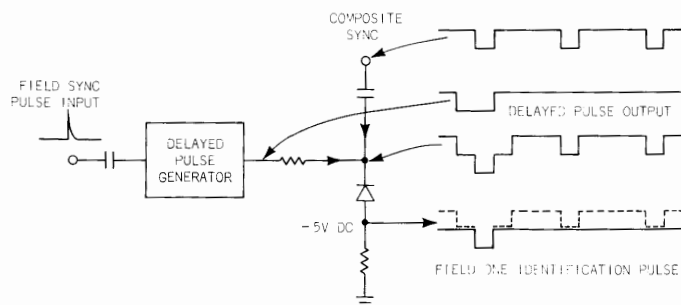


Fig. 9-13. Block concept of a field identification pulse-forming circuit.

The simplified circuit is illustrated in Fig. 9-14. Quiescently, R5 maintains Q2 in conduction and D9 clamps the output DC level close to ground.

A field sync pulse, applied through C_{in} , triggers the monostable multi, turning Q2 off, but since D9 is already conducting, the positive-going transition does not appear at the output. The time constant R5-C5 is adjusted to reset the multi back to the stable condition six lines (381 μ s) after the multi is triggered. When the multi does reset, Q2 turns on again -- coupling the negative-going transition to the output and biasing D11 into conduction. The next composite sync pulse to occur will be coupled through D11. A negative-going pulse is then developed across R12 which can be used for field-ONE identification. The duration of the negative-going transition is determined by the time constant R9-C9/0.11. (C9 is only permitted to charge 11% before being clamped by D9.)

Fig. 9-15 illustrates the details of the circuit which insure the operation of the basic concept under various practical operating conditions. The various considerations are:

- | | |
|----------------------------|---|
| field sync pulse amplitude | 1. Field sync pulse amplitude The positive-going field trigger pulse applied to C_{in} is amplitude-limited by D1. R1 and R2 set the quiescent bias to about -2 VDC so the pulse amplitude will never exceed 2.5 V peak-to-peak. The limited pulse amplitude prevents the possibility of Q1 ever saturating and affecting the reset time. |
| jitter | 2. Reset jitter Composite sync is applied to the base of Q2 through C7 to reset the monostable multi. C7 is a small value so only the leading and trailing edges of the sync pulse are applied to the base of Q2. The trailing edge of the last equalizing pulse, a positive-going transition, occurs just before the multi would normally reset. The trailing edge actually resets the multi, assuring a precise reset time independent of transistor and component variations. |
| diode gate bias | 3. Diode gate bias To establish the bias on the gate diode D11, just below the negative sync tips, C10 is added in series with R12. |

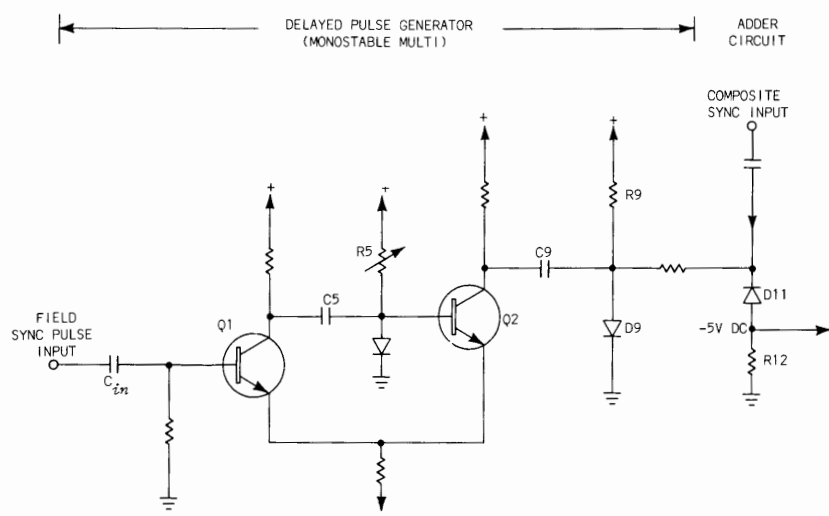


Fig. 9-14. Simplified field pulse circuit.

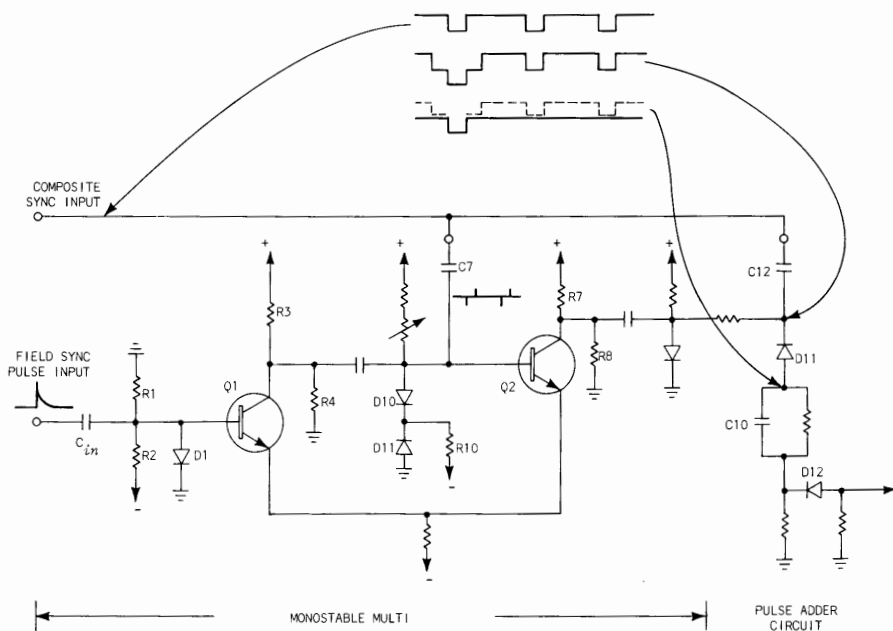


Fig. 9-15. Detailed illustration of a field identification-pulse-forming circuit.

Initially, negative-going sync pulses applied through C12 will forward bias D11 and start to charge C10. Eventually the charge on C10 will equal the peak negative amplitude of the composite sync pulses, which are 5 V peak in this case. (The resistor across C10 is only to maintain bias on D11.) After C10 is charged, D11 will not conduct until the combined amplitude of the composite sync pulse and the multi output pulse exceeds 5 V peak. D12 prevents negative transitions from subsequent circuitry from affecting the average charge on C10.

time
coincidence
pulse
identifi-
cation

An alternative method of obtaining a field identification pulse also employs two time-coincident pulses but without the need for a delayed pulse generator.

The system consists of two elements:

1. An integrator amplifier to integrate the composite-video field sync pulses.
2. A time-coincidence gate operated by the integrated field sync pulses.

The two signals required at the time-coincidence gate are:

1. The integrated field sync pulse to operate the gate.
2. H-rate sync pulses.

The usefulness of this circuit is based on the availability of sync pulses occurring *only* at a line scanning rate, therefore, every alternate equalizing pulse is absent.

Notice in Fig. 9-2, on field ONE, the first equalizing pulse occurring after the serrated field sync pulse is at the scanning-line rate while on field TWO, the *second* equalizing pulse after the serrated field sync pulse occurs at the scanning-line rate.

The time-coincidence circuit utilizes the first line-rate pulse of field ONE.

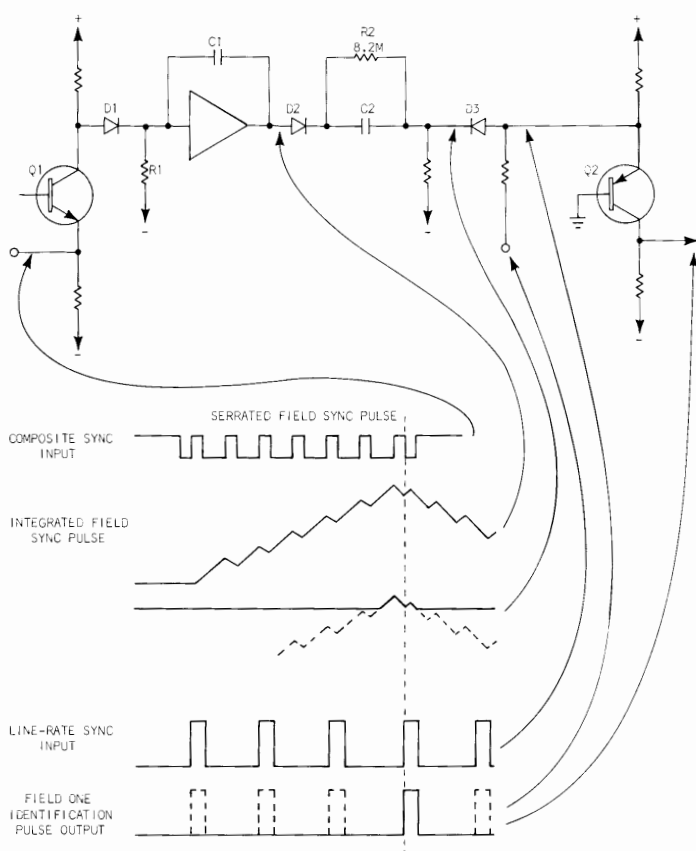


Fig. 9-16. Field identification pulse circuit utilizing a field sync pulse integrator.

Fig. 9-16 illustrates the basic operation of the circuit. Composite video sync is applied to the integrating amplifier through an input amplifier. The narrow line sync pulses and the serrated field sync pulse reverse-bias D1, allowing the integrator capacitor to charge.

The comparatively narrow line sync pulses do not reverse-bias D1 long enough for C1 to charge to any extent. The wider serrated field pulses do, however, allow the capacitor to charge, resulting in a serrated sawtooth waveform at point B.

D2, R2, and C2 perform the function of a peak detector after one or two integrated field sync pulses charge C2. The resistance of R2 is made very large to prevent any appreciable discharge of C2 between integrated field sync pulses. R2 is necessary to provide bias current for D2 which will then conduct only at the peak of the sawtooth waveform.

When the peak of the sawtooth does forward-bias D2, the signal is coupled through C2 to reverse-bias D3. D3, when conducting, functions as an inhibiting diode preventing line sync pulses applied to point D from appearing at point E. Any line sync pulse occurring within a half a scanning line after the termination of the integrated field sync pulse (sawtooth peak) will appear at point E and forward bias the amplifier Q2.

BASIC COLOR CONCEPTS

Color television is a system comprising two completely separate fields of science:

1. The practical application of colorimetry to televised images, and
2. The methods by which the color characteristics of the image are converted and packaged into the standard composite video waveform (Engineering).

Evaluation in depth of the circuits used to process and display color information requires a review into the application of colorimetry to electronically-reproduced images as well as a review of the electronic system and its resultant composite waveform.

The science of Colorimetry is the measurement and specification of color in numerical terms. The purpose of the following discussion is to define the qualities of color and why color television is dependent on Colorimetry (assigning numerical quantities to color).

Much of the fundamental work that forms the modern technical treatment of color in light was done by Sir Isaac Newton, who observed that if the radiation from the sun did not contain more than one wavelength, "there would be but one Colour in the whole World." He concluded that the visual sensation of color was not because light rays were colored, but because of selective-reflection or transmission of the incident radiation of light from an object. Newton further found that white light could be produced by adding discrete mixtures of separate colors either simultaneously, or sequentially at a high enough rate for the eye itself to perform the additive process.

While Newton's experiments were concerned mainly with the mixture of different color light sources, he failed to make a distinction between the mixture of pigments (on which modern color printing is based) and the mixture of light sources. The lack of distinction between the two processes led to a considerable amount of confusion. About 100 years ago James Clerk Maxwell demonstrated that full color pictures could be reproduced by making three slides, each photographed through a different colored glass. The colored glasses that Maxwell chose were red, green, and blue.

additive
color
process

The process of producing a color reproduction with colored light sources is now known as the additive color process. The additive color process is the only practical method of producing a color television picture where the colored lights are replaced with phosphors in a cathode ray tube.

subtractive
process

The second method of reproducing color was first demonstrated in 1722 using an entirely different principle, empirically improved over the years into what is now the basis of the modern printing process. Basically, pigments selected to absorb certain colors of light and reflect other colors are placed in layers on a surface and light is reflected from the surface. When viewing a color photograph, proper color rendition to the eye is dependent on ambient light -- not colored light but white light. The photograph is actually three separate pictures laid down on the paper in layers. One layer reflects the red and green light and absorbs the blue; this layer is called the minus-blue or yellow layer. The second layer reflects blue-red and absorbs green; this layer is known as minus-green or magenta. The third layer reflects green and blue and absorbs red; this layer is known as the minus-red or cyan layer. Since each layer absorbs part of the ambient light, the system is known as the subtractive method of color production.

The primary color pigments used in the subtractive process are the complimentary colors of the additive process.

Since color television employs the additive color process exclusively, the subtractive process had been discussed at this point only to form a comparison. The important point to remember is that the additive

process of color mixing employs colored *light sources* while the subtractive process is dependent on selective filtering of ambient *white light*.

Before colored light can be numerically measured and specified, the variable qualities of light must be defined. The three qualities necessary to specify a color exactly are:

charac-
teristics
of colored
light

1. Brightness
2. Hue
3. Saturation

brightness,
luminance

Brightness is the total amount of light energy perceived by the eye. Brightness is interpreted relatively by the eye as being very dim to very bright. For example, the color of the CRT display from an oscilloscope is fixed by the chemical composition of the phosphor. As the intensity control is increased, only the brightness of the display is changed.

hue,
wavelength

Hue (also called color or tint) is generally the most noticeable quality of light perceived by the eye. Hue is a characteristic of visible light energy that identifies the wavelength of radiant light and is usually interpreted by the normal eye as red, blue, yellow, etc.

saturation,
vividness

Saturation can be described as the "vividness" of a color in terms of pale, pastel, deep, etc. Saturation is a measure of how much the particular color appears to differ from grey or white. White light is the sum of three or more suitable primary colors of *equal light energy*. Therefore, saturation is defined as the degree of a predominant hue observed by reflection, by transmission or directly from a light source.

These three qualities of light can be measured in terms of three quantities:

1. Luminance
2. Dominant wavelength
3. Color purity.

The television system transmits electrical signals proportional to the qualities luminance, hue, and saturation of a colored object. Reference to the

electrical signals is generally made in terms of luminance, hue, and saturation. Before the additive color process could be commercially applied to television on a systematic basis (without having to empirically tailor each receiver to produce pleasing color), these fundamental questions had to be answered:

primary
colors

1. How many primary colors are needed?
2. What primary colors should be used?
3. What proportions of each primary color will produce a new desired color?

Newton concluded from his experiments with the prism that there were seven primary colors, but later experimenters found that four of Newton's primary colors could be reproduced with mixtures of the remaining three primary colors -- red, green, and violet lights. These experimenters found the choice of the primary colors to be loosely determined by two factors:

1. Each primary color must be different from the other two primaries.
2. Any combination of two primary colors must not equal the color of the third primary.

These two factors permit numerous choices of primary colors, however an additional factor, the range of color to be reproduced from the primary colors, narrows the choice of primary colors considerably.

The third question, the proportions of primary colors to produce a new color combination, has led to the development of colorimetry.

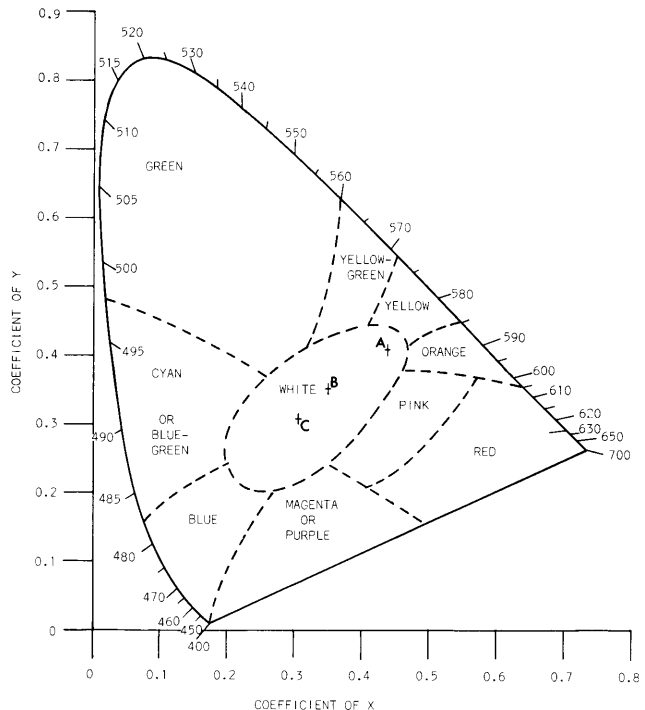
Maxwell suggested that any color could be specified in terms of three primary colors mixed to produce that color. Grassman, a contemporary of Maxwell, suggested that a color be specified in terms of its percentage contribution to "white" light. Both suggestions were adopted by the International Commission on Illumination and used to form four major conclusions:

1. Choice of three primary colors with standard values.
2. Reference white is specified and standardized as "equal energy white," that is, equal amounts of each primary. Three standards of white are specified: A, B, and C. Standard Source C is commonly used to specify CRT phosphors.

3. A color match at one brightness level will be maintained over a wide range of brightness levels. Normally, all colors have different brightness levels; for example, yellow light is brighter than blue light. By normalizing the brightness levels of all colors, the brightness can be omitted when specifying a color. With the brightness component of colored light removed, the values of color can be drawn on a two-dimensional diagram (as devised by the commission).
4. Color matches using three primary colors obey the laws of addition and subtraction.

ICI
chromaticity
chart

As a result of these and other conclusions, the standard ICI chromaticity chart (Fig. 10-1) was developed. For the same reason that the study of geography is facilitated by the use of maps, the

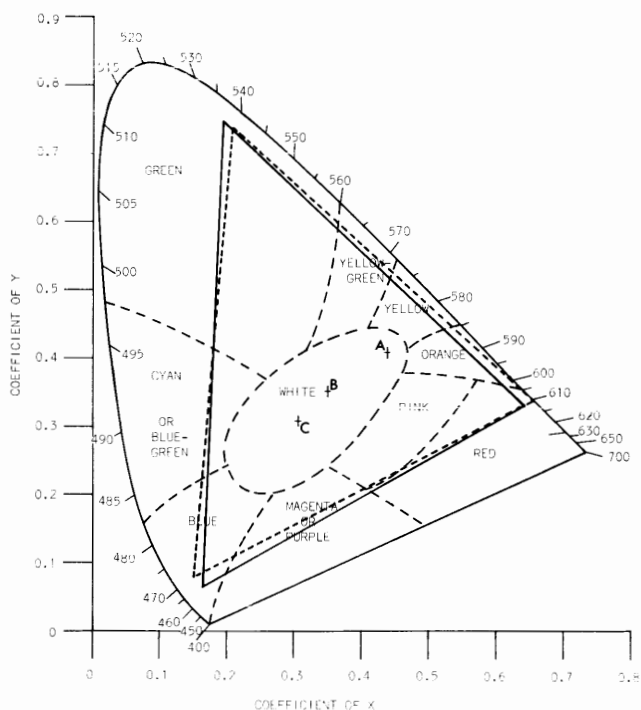


The bounded area includes all real colors perceptible to the average human eye. The colors located on the periphery are spectrally pure with the exception of the purples which are combinations of blue and red. Numbers along the periphery are wavelengths in nanometers of the spectral colors. Standard ICI illuminant "C" is the standard representation of average daylight (Source C). The transition from one color to the next is gradual rather than sharply defined as shown.

Fig. 10-1. Standard Chromaticity chart.

numerical quantities of color can be represented graphically on a chromaticity diagram. The chart permits three-color specification to represent any of the 10,000 or more colors and brightness levels which may be distinguished by the eye.

The color and luminance characteristics of the filters in the television camera and the phosphors in the receiver picture tube are based on the numerical quantities of the chromaticity diagram shown in Fig. 10-2.



The solid triangle encloses the area of colors capable of being reproduced by a P-22 tri-color phosphor. The dotted triangle has been established by the N.T.S.C. (National Television Systems Committee) as a reference.

Fig. 10-2. Typical phosphor characteristics.

11

COLOR TELEVISION SYSTEM REQUIREMENTS

The television system is not capable of transmitting color as color, but only as electrical signals proportional to certain *characteristics* of color. The transformation is made as follows:

At the camera, each color in the scene is optically reduced to three primary colors through the use of filters as shown in Fig. 11-1. The amounts of light passing through each filter into the three camera tubes produce picture signals representative of the amounts of the three primary colors present in colored objects being scanned by the camera. The three primary color signals could then be transmitted directly over three separate channels if system compatibility (transmission and reception on existing black and white equipment) were not required. As it turns out, only the terminal points of the television system -- the camera and the picture tube -- analyze and reconstruct color light in the form of red, green, and blue primary colors.

Dichroic mirrors which reflect most of the incident light at specific wavelengths and pass most of the incident light at other wavelengths, are used as color-selective light splitters.

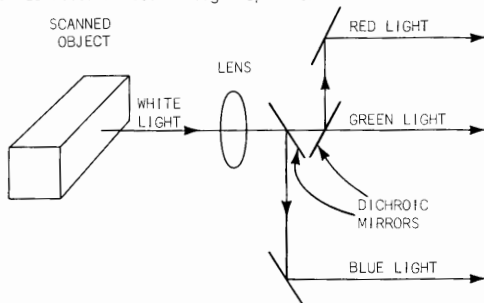
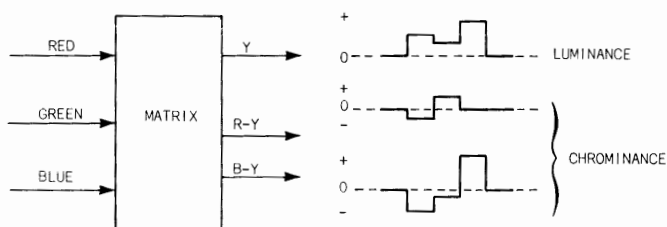


Fig. 11-1. Translation process block diagram of colored light into three electrical signals.

bandwidth economy	To conserve channel bandwidth and permit color transmission over conventional television systems, the red, green, and blue primary color video signals are translated into three related quantities; luminance, hue, and saturation; which are then processed and "packaged" as a composite video signal consisting of two separate signals, luminance and chrominance. Luminance is the brightness information essentially identical to conventional black and white systems and transmitted as white or shades of gray. The chrominance signal is the combined hue and saturation representing the colorimetric <i>difference</i> between a color and a reference white light of the <i>same luminance</i> as the particular color; in other words, the luminance characteristic of the light is electronically removed, leaving only the signals proportional to hue and saturation.
luminance signal	
chrominance signal	
signal-combining matrix	The transformation from red, green, blue video signals to luminance, hue, and saturation is carried out by means of a resistive matrix -- algebraically combining percentages of each primary color signal to form the luminance (Y signal) and the two color-difference signals (R-Y, B-Y) as shown in Fig. 11-2.



The luminance signal consists of fixed percentages of each primary light.

The two color-difference signals form the CHROMINANCE signal. Chrominance is defined as the colorimetric *difference* between any color and a reference white of equal luminance (in this case, standard "C").

Hue is determined by the *relative amounts* of the two difference signals; Saturation is determined by the *absolute amounts* of the two difference signals.

Fig. 11-2. Transformation of red, green and blue video signals to new related quantities called luminance, hue, and saturation.

The red, green, and blue camera output signal percentages contributing to the luminance and chrominance signals respectively are:

color-signal
constituents Luminance,
 Y: $0.30 \text{ red} + 0.59 \text{ green} + 0.11 \text{ blue} = 100\%$
 brightness. While the absolute values of
 each signal output voltage vary as the
 scene is being scanned, the *ratios*
 contributing to brightness remain the same.

The chrominance signal consists of the two color-difference signals,

R-Y: $0.70 \text{ red} - 0.59 \text{ green} - 0.11 \text{ blue}$ (when
transmitting saturated red).

and,

B-Y: $-0.30 \text{ red} - 0.59 \text{ green} + 0.886 \text{ blue}$ (when
transmitting saturated blue).

The color-difference signals produce correct colored light output from the phosphors of the CRT only when recombined with the luminance signal.

For example:

$$\begin{aligned} \text{White} &= 0.30R + 0.59G + 0.11B \\ \text{Add difference signal, R-Y} &= \frac{0.70R - 0.59G - 0.11B}{1.00R + 0.00G + 0.00B} \end{aligned}$$

100% of the brightness is now produced by the red phosphor while at the same time the green and blue phosphors have been turned off -- saturated red is produced. Since the R-Y, B-Y signals are difference signals, all the values of the color-difference signals reduce to zero when white light is present.

The luminance signal "Y" is essentially identical to the standard black and white video signal and, therefore, does not require further processing.

After the color-difference signals are derived, several transmission possibilities exist:

system
options

1. Transmit the two difference signals on two additional channels.
2. Transmit the two difference signals on two low-level subcarriers within the luminance channel with a loss of luminance detail.
3. Combine the two difference signals onto *one* subcarrier within the luminance channel.

The best compromise (adequate color information and minimum degradation of the black and white "luminance" signal) can be achieved by combining the two color-difference signals onto one subcarrier within the luminance channel bandwidth by using a special modulation process.

Transmission of two signal components on one subcarrier is made possible by a process called quadrature-amplitude modulation. This process is illustrated in the simplified block diagram of Fig. 11-3. The two independent signals (R-Y, B-Y) amplitude-modulate carriers of identical frequency, but different phase; the two modulated outputs are then added together to feed a common transmission channel, forming a sum of the two modulated sinewaves. The resultant sinewave varies in both amplitude and phase when compared to the original reference sinewave.

modulating
in
quadrature

For transmission of two independent signals, the two carriers may be separated by any phase angle except 0° or 180° , but in practice the phase separation between the two carriers is established at 90° so the original signals modulating the two carriers can be recovered at the receiving end without cross-talk. The modulation system is, therefore, called quadrature amplitude modulation, providing a means for using the two sidebands surrounding a single

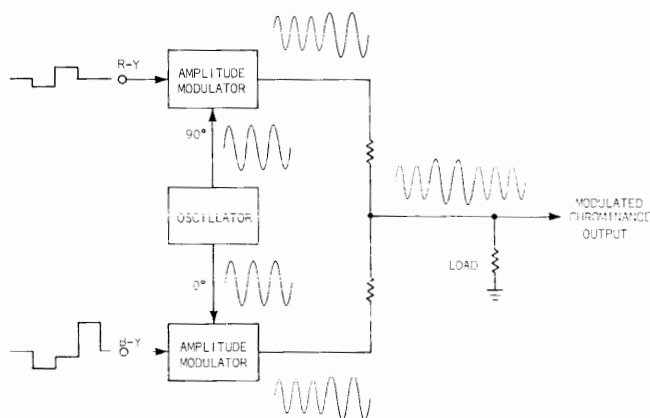


Fig. 11-3. Simplified block diagram of the color-difference signal modulator system.

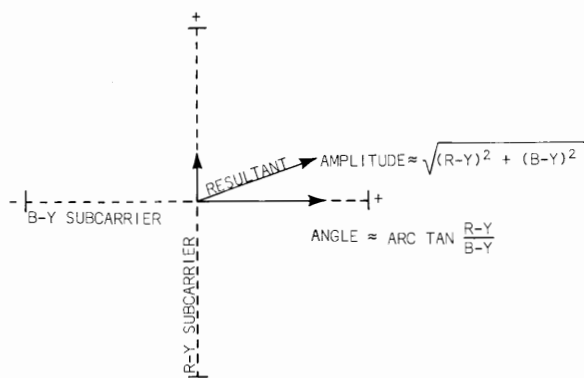
carrier frequency for the transmission of two variables. In contrast, a normal amplitude-modulated carrier contains two sidebands -- one above the carrier (sum) and one below the carrier (difference) with identical information on both sidebands.

The requirements for quadrature amplitude modulation are:

1. Two pieces of information modulate carriers of the same frequency.
2. Phase separation of 90° for the two carriers.
3. The carrier is omitted from the transmission.

The resultant modulated signal can be visualized two different ways:

1. Two separate amplitude modulated subcarriers identical in frequency but different in phase.
2. A single resultant subcarrier sinewave with amplitude modulation proportional to $\sqrt{(R-Y)^2 + (B-Y)^2}$ and phase modulation (when compared to the reference frequency) proportional to $\arctan \frac{R-Y}{B-Y}$. (Fig. 11-4.)



Dashed vector lines show the maximum subcarrier modulation. Solid vector line illustrates the instantaneous amplitudes of the two subcarriers and the resultant amplitude when they are added together.

Fig. 11-4. Vector diagram of the two chrominance subcarriers.

suppressed
carrier

Earlier, mention was made that the carrier must be omitted from the transmission. Carrier suppression in any transmission system can have several advantages as well as several disadvantages. In the case of the compatible color television system, the advantages outweigh the disadvantages. The advantages of suppressed carrier are:

1. Conservation of signal energy.
2. Reduction of luminance interference -- the possibility of spurious effects on the luminance (black and white) information is reduced because the subcarrier amplitude goes to zero when the camera scans a white or gray object. (An FCC requirement.)

The disadvantages of a suppressed carrier are:

generate
local
carrier

1. Since the carrier normally is used as a reference frequency to recover the modulated information contained in the sidebands a local carrier must be generated, correctly phased, and added to the sideband information before any intelligence can be recovered.
2. A sample of the carrier either in the form of a harmonically related sinewave, or a small sample of the carrier itself is required to control the locally-generated carrier. In the compatible television system a harmonically related sinewave is not practical because additional bandwidth is required. Since *time* is available during horizontal blanking when the subcarrier sidebands are never transmitted, a sample of the carrier frequency with a predetermined phase is added to synchronize the locally-generated carrier.

synchronize
local carrier
with original

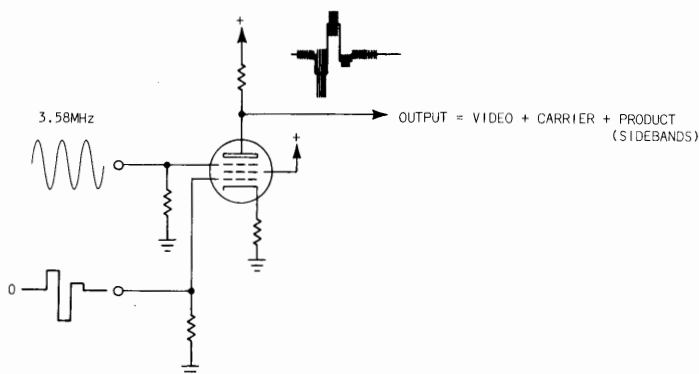
12

COLOR MODULATION AND DEMODULATION

The modulation and demodulation process of the additional chrominance information is the basis of the compatible color transmission system. Since modulation and demodulation are very similar, a basic discussion of the modulation process will serve to better describe the principles of the demodulation process and their application in the circuits of the color vectorscope.

simple
modulator

Fundamentally, the simple modulator is nothing more than an amplifier with *two* independent controls over the plate current, as illustrated in Fig. 12-1. Assume for the moment that each grid has *equal* control over the plate current. If either signal is held constant (but still permitting plate current to flow), the other control grid can affect the change in plate current entirely. On the other hand, if both grids move positive the same amount the plate current is proportional to the *product* of the two input signals, but *in addition* to the product output the amplifier also acts as an ordinary amplifier to the individual input signals, inverting them at the output. The



For purposes of illustration, the control grid and the suppressor grid are assumed to have equal control over plate current.

Fig. 12-1. Simplified amplitude modulator.

output signal voltage is then the product of two input signals *plus* the sum of the two input signals (if both input signals move in the same direction), or the difference of the two input signals (if the two input signals move in opposite directions).

In the circuit just described, either or both the input signals could be filtered at the output except that the frequencies (the carrier and the desired sidebands) are not widely separated. The alternate solution is to balance both input signals to cancel at the output, leaving only the product of the two input signals. For this reason, the circuit shown in Fig. 12-1 is not suitable for modulation of chroma information.

The circuit of Fig. 12-2 illustrates the effect of balancing the video to cancel at the output, leaving only the carrier and the sidebands generated in V1.

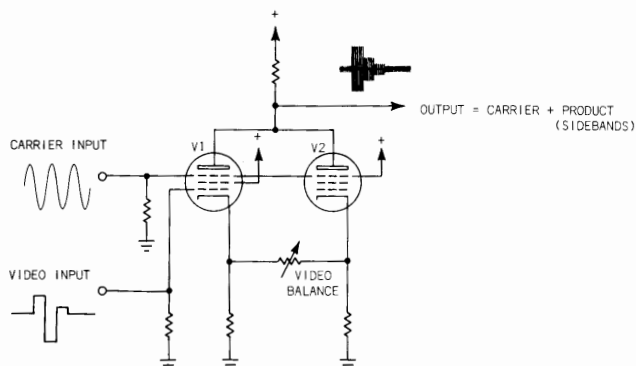
removing
carrier

Similarly, the carrier can be cancelled at the output, leaving only the video and carrier sidebands as illustrated in Fig. 12-3.

If both carrier balance and video balance are combined in a circuit similar to the one illustrated in Fig. 12-4, the result will be a true product output with both input signals cancelling at the output. The circuit is then termed a "doubly balanced" product modulator. Since modulation is so much like demodulation, any one of the three circuits previously mentioned might also be used as a demodulator. When the modulator is used as a demodulator one of the inputs, instead of being video, is a pair of sidebands. The difference frequencies will be the original video information and the sum frequencies will be the sidebands around the second harmonic of the carrier which must either be cancelled at the output or removed with a filter. Since a low-pass filter provides the most practical means of removing not only the second harmonic, but also the carrier frequency as well, demodulators as a rule are not balanced. When any one of the three circuits are used as demodulators, the output waveform will appear as illustrated for the modulator with two exceptions:

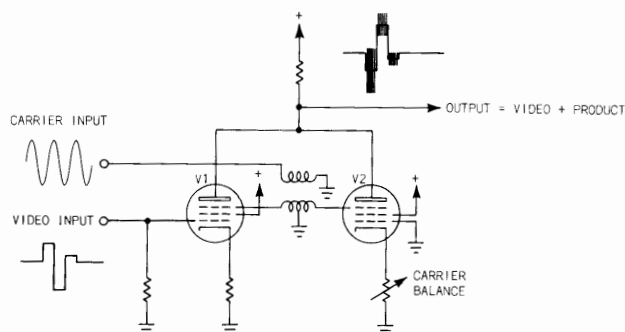
low-pass
filter

1. Second harmonics of the carrier will be present.
2. The amplitude will be half the amplitude of the original video after the second harmonics are removed.



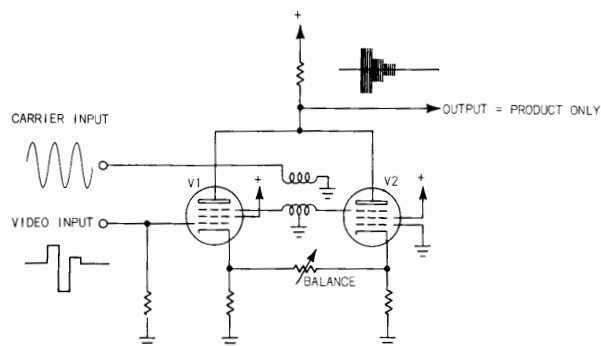
Only the video on the plate of V1 is inverted which when combined with the noninverted video on the plate of V2 is cancelled from the output signal, leaving only the carrier and the video/carrier product (sidebands).

Fig. 12-2. Video-balanced modulator.



The carrier is applied push-pull to the suppressor grids of V1 and V2 which cancel at the output.

Fig. 12-3. Carrier-balanced modulator.



Both input signals are cancelled at the output leaving only the product (sidebands).

Fig. 12-4. Doubly-balanced modulator.

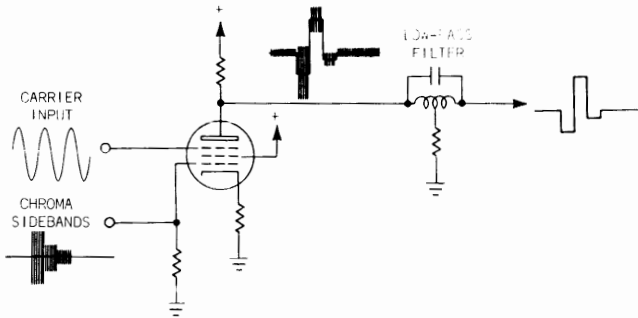


Fig. 12-5. Simplified demodulator-also called a synchronous detector.

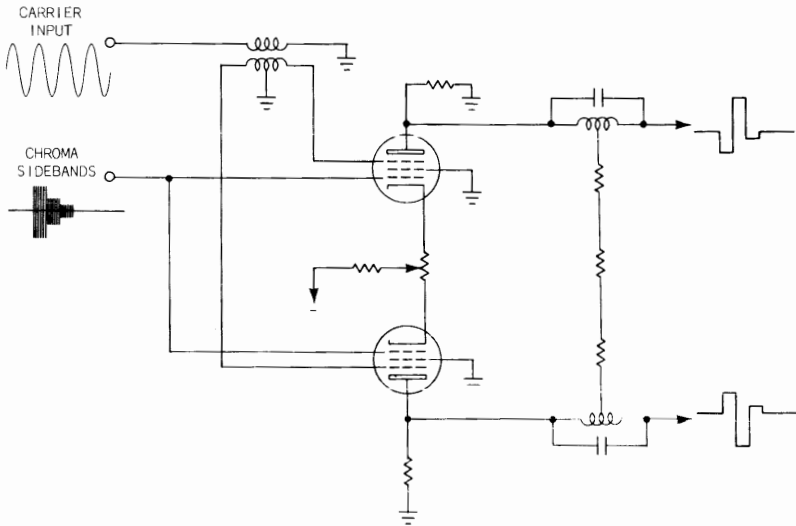


Fig. 12-6. Push-pull synchronous detector.

synchronized
detector

Fig. 12-5 illustrates a simplified diagram of a synchronous detector. The term "synchronous" implies that a locally-generated carrier, synchronized to the original carrier, is applied to the detector along with the incoming chroma sidebands. Sideband information containing the original video cannot be recovered until the suppressed-carrier signal is first multiplied by a synchronized carrier. Note that the waveform at the plate of the detector is very similar to the waveform produced by the simple modulator of Fig. 12-6 except for the presence of the second harmonic component at the detector output.

push-pull
synchronous
detector

A more complete synchronous detector used for direct vector display on an oscilloscope is illustrated in Fig. 12-6. The principal difference between the circuit illustrated in Fig. 12-6 and the simplified detector illustrated in Fig. 12-5 is the addition of a second synchronous detector to form a push-pull circuit. The total combination, however, is *not* a balanced detector for either the carrier or the sidebands. The push-pull arrangement is used mainly because the detected color difference signals will be amplified and applied directly to the deflection plates of a CRT. In addition, the push-pull demodulator is nearly 100% efficient because chrominance is being detected on both the positive and the negative half of the subcarrier sinewave. In contrast, a color picture monitor (or television receiver) must first convert the chroma and luminance back into the original red, green, and blue picture signals before a three-color picture can be correctly displayed.

circuit
limitations

While this circuit is quite suitable for measurement applications in an oscilloscope, the circuit has one limitation. The subcarrier signal appears at the output mixed with the chroma and must be filtered to make the chroma information useful. Filtering will easily remove subcarrier, and can remove part of the signal components as well. For some measurement applications, observation and measurement of the transitions from one color to another are desirable. Excessive filtering removes much of these transitions -- in effect, distorting part of the chrominance signal.

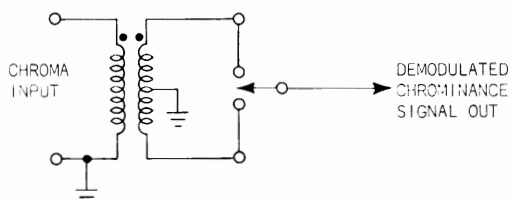


Fig. 12-7. Chrominance-switch demodulator. Switch is operated at subcarrier repetition rate.

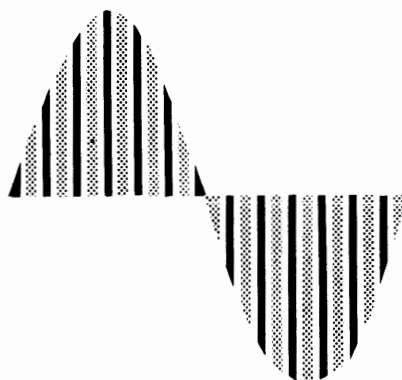


Fig. 12-8. Demodulator output from typical full-wave switch. The envelope shape represents the color difference signal. The bands represent one-half cycle of the 3.58 subcarrier.

additional requirements Chrominance demodulation without waveshape distortion of the recovered chroma signal imposes several requirements:

1. Demodulator sampling time (percentage of one cycle of subcarrier the chrominance signal is actually "seen" coupled to the output) should be as high as possible.
2. The demodulator voltage output should be linear with respect to the varying amplitude of the chrominance input signal. Therefore, the demodulator output should not be affected by amplitude variations of the subcarrier.
3. Minimum amount of subcarrier filtering at the output to prevent removal of the higher frequency chrominance components.

carrier-operated switch demodulator Fig. 12-7 illustrates a simplified circuit which fulfills most of the requirements previously listed.

The configuration is fundamentally a full-wave carrier-operated switch with an output waveform similar to that shown in Fig. 12-8.

In contrast to amplitude-selective demodulators, switch circuits are basically time-selection devices. The principal advantage of the switch circuit is that there is no critical dependence on the amplitude or shape of the switching waveform. In addition, since color difference signals do have high-frequency transient information nearly equal to one-half of the carrier frequency, the rapid response of the time selector makes it particularly useful as a demodulator of an amplitude-modulated waveform.

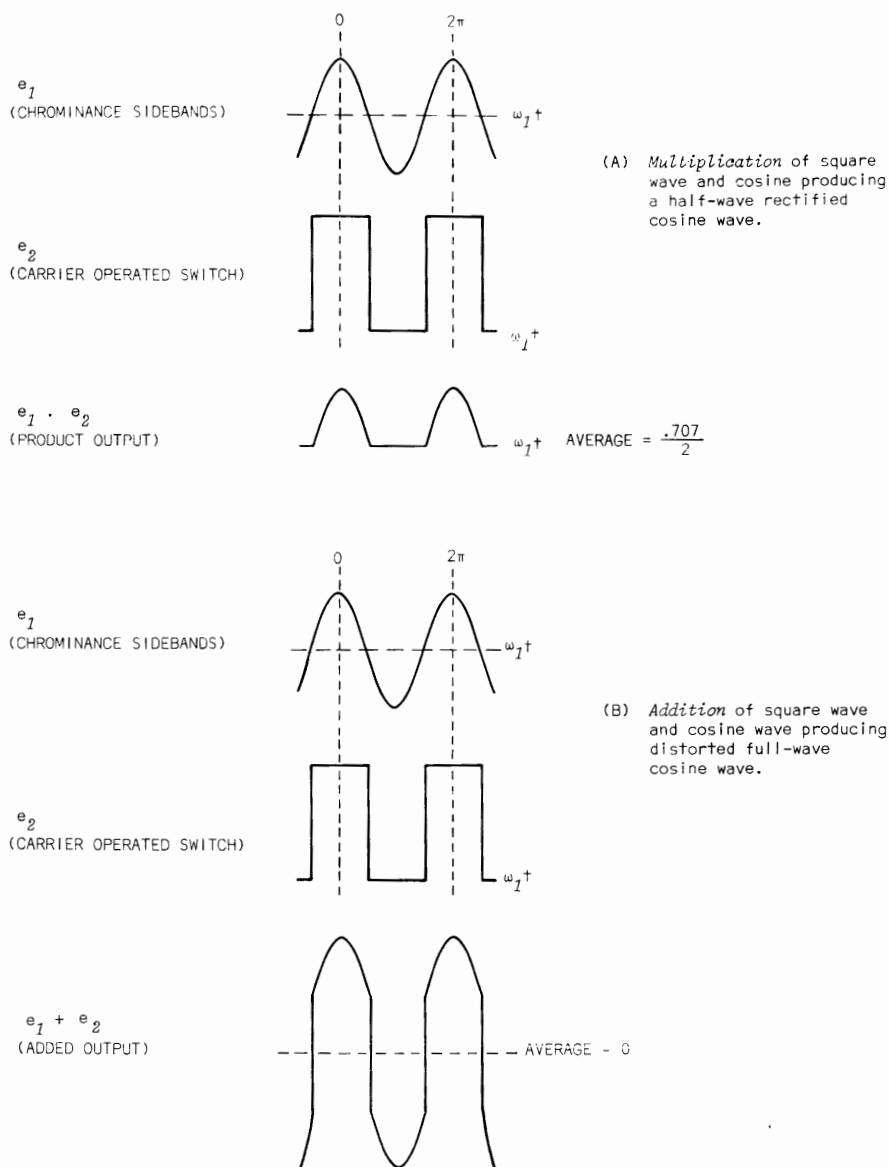


Fig. 12-9. A better understanding of the principle of demodulation can be seen by the graphical display contrasting two waveforms; (A) multiplied together and (B) added together.

The product of the carrier-operated switch waveform and the chroma sideband information can be graphically seen in Fig. 12-9A. Contrast the nonlinear (product) waveform with the linear (resistively added) waveform of Fig. 12-9B.

The simplification of the switch circuit is limited by the nonavailability of mechanical switches that will operate at the subcarrier frequencies. However, by slightly modifying the switch configuration as shown in Fig. 12-10, a transistor switch can be used.

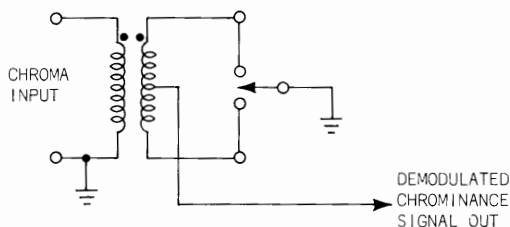


Fig. 12-10. Modified chrominance-switch demodulator. Switch is operated at subcarrier repetition rate.

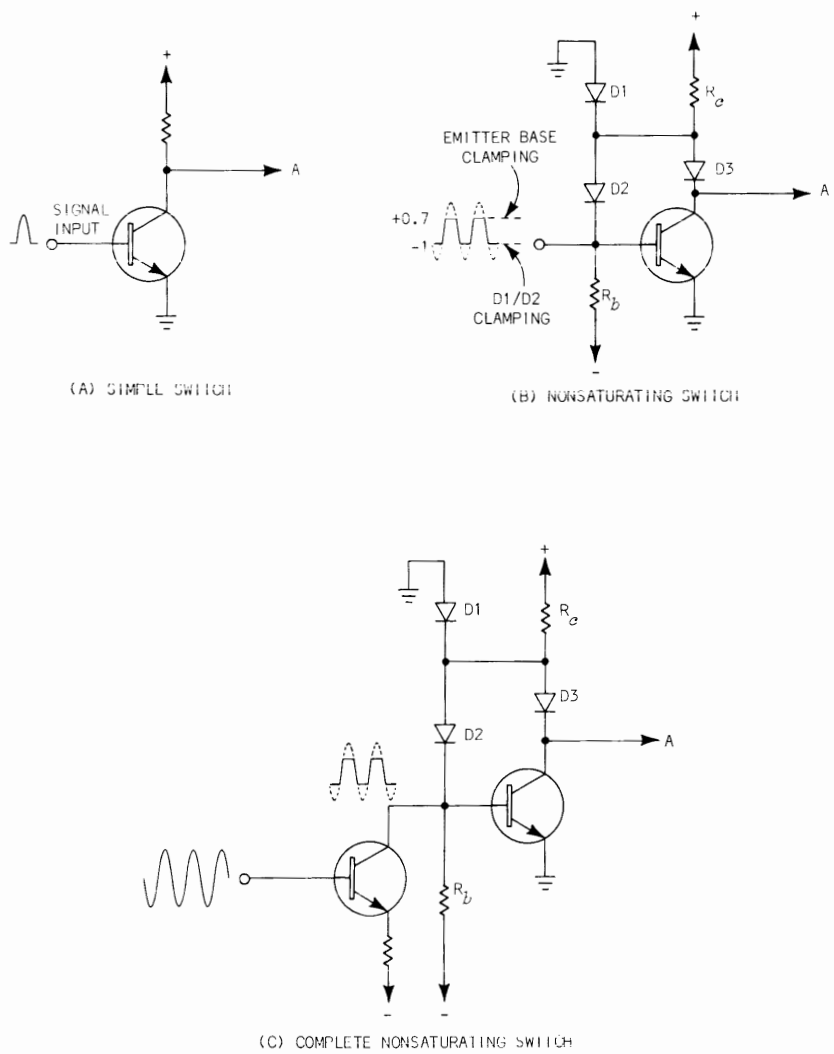


Fig. 12-11. Transistorized carrier-operated switch.

transistor
switch

An equivalent transistor switch is shown in Fig. 12-11A. When a positive transition such as a positive-going portion of the subcarrier sinewave is applied to the transistor base, the transistor conducts -- essentially grounding point "A." However, though transistor off-to-saturated switches are easy to turn on they have one major disadvantage: they take time to switch off again, making the sampling period somewhat difficult to control.

The addition of three diodes as shown in Fig. 12-11B effectively prevents transistor saturation by limiting the transistor collector from becoming too negative and forward-biasing the base/collector junction. However, point "A" is still very close to ground when the transistor is conducting.

transistor
switch
circuit

The operation of the switch circuit is as follows: A positive-going portion of the subcarrier sinewave forward-biases the transistor. The sinewave amplitude is limited by the forward-biased emitter/base junction.

Quiescently, D1 and D3 are not conducting; D2 is always conducting. The diode junction of D2 serves only to provide a DC voltage difference between the transistor base and anode of D3 so the waveform on the transistor base will also appear unattenuated, at the anode of D3. When a positive-going portion of the sinewave is applied to the transistor base, the signal will not only cause the transistor to conduct, but will also cause D3 to conduct as well. The conduction of D3 clamps the collector just above the base voltage, thus preventing saturation.

Clamping both the positive and negative portions of the sinewave applied to the base has several advantages:

1. The switch is "on" or "off" for nearly 50% of the sinewave instead of 25 to 30%.
2. Switching time differences are less subject to component variations.

The clamping of the positive and negative portions of the sinewave requires an isolation stage between the subcarrier source and the transistor switch to prevent loading and distorting the source. The added isolation stage is shown in Fig. 12-11C.

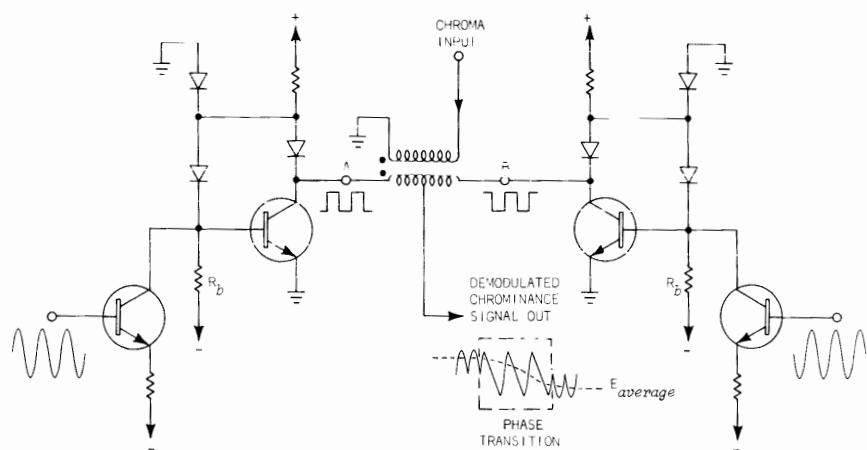


Fig. 12-12. Complete full-wave switching demodulator.

switch
demodulator
circuit

The ideal switch of Fig. 12-10 has a single-pole double-throw action; the transistor is a single-pole single-throw switch. The addition of a second switch transistor as shown in Fig. 12-12, alternating with the switching action of the first transistor, will provide the necessary double-throw action.

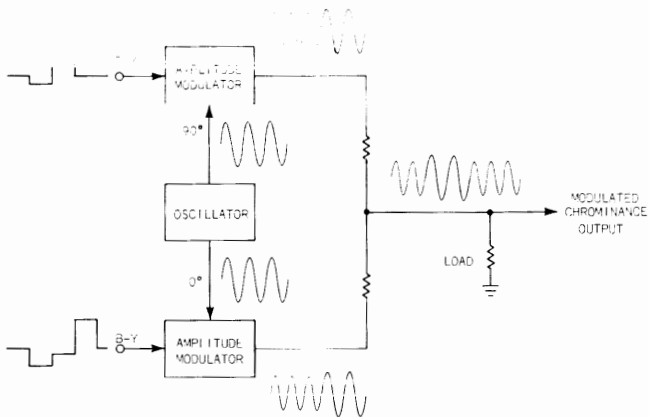


Fig. 13-1. Simplified block diagram of the color-difference signal modulator system.

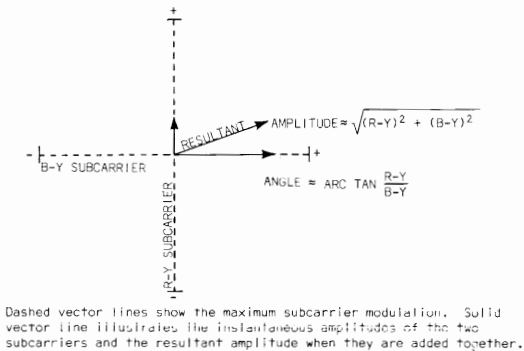


Fig. 13-2. Vector diagram of the two chrominance subcarriers.

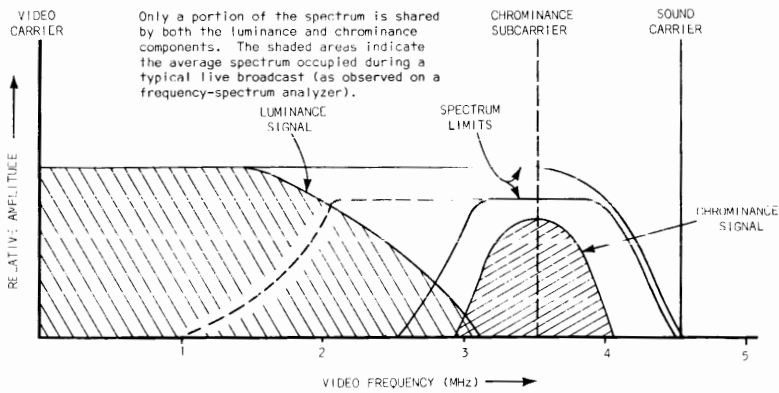


Fig. 13-3. Video frequency spectrum occupied by luminance and chrominance components of a complete color picture signal.

13

CHROMINANCE SUBCARRIER
REGENERATORS

After the two chroma difference signals have modulated the subcarriers, the two subcarriers are added together to form a resultant sinewave as illustrated in Fig. 13-1. If the resultant sinewave is visualized by the vector diagram shown in Fig. 13-2, the resultant sinewave, when compared to the reference oscillator, will have an instantaneous amplitude proportional to the square root of the sum of the squares of the two added signals and a resultant phase relationship to a single reference oscillator, proportional to $\arctan (R-Y)/(B-Y)$. The resultant sinewave is then added directly to the video luminance waveform.

frequency
multiplexing

The process of adding the chroma subcarrier sidebands directly to the luminance is often called "frequency multiplexing" (the simultaneous transmission of separate information) and in this case two separate bits of information are carried on one waveform. The luminance and chrominance information occupy the same frequency spectrum. (See Fig. 13-3.) When both luminance and chrominance information occupy the same frequency area, crosstalk between the two signals is possible unless preventative steps are taken. Simply adding the chroma sidebands to the video waveform without crosstalk is not possible without evaluating the possible effects of the chrominance sidebands on the luminance portion of the signal. Therefore, careful consideration and field testing was given to the frequency choice of the color subcarrier. A conventional receiver or picture monitor interprets the chrominance information as standard video (amplitude of the chroma is the same as brightness information and the subcarrier frequency is equivalent to picture detail information).

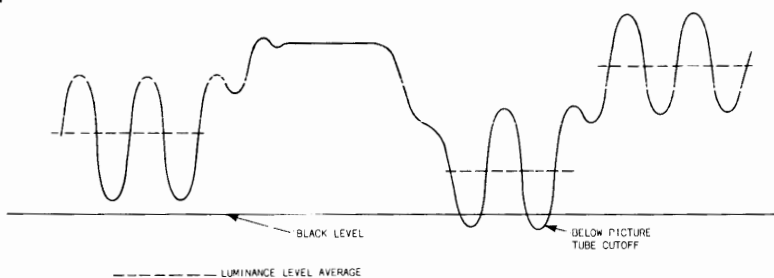


Fig. 13-4. Portion of luminance signal with the modulated chrominance signal added.

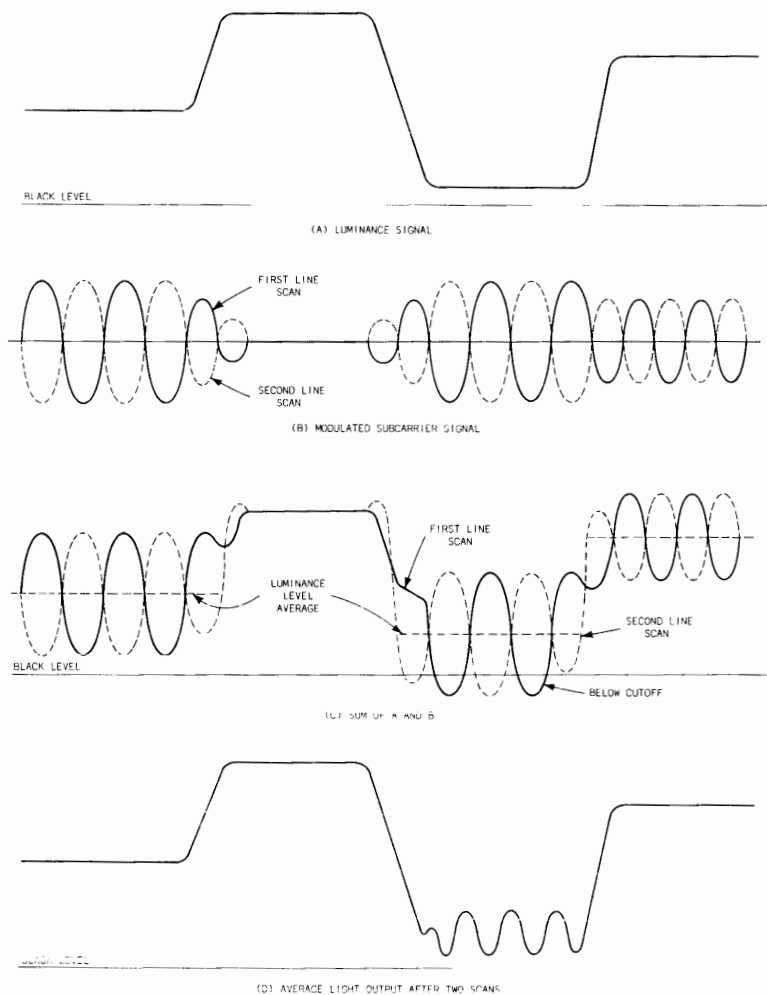


Fig. 13-5. Time-displacement of the chrominance subcarrier signal on successive horizontal scanning lines.

choosing subcarrier frequency

Assuming for the moment that the presence of the color subcarrier on the luminance signal will produce a visible interference pattern on the displayed picture, the interference will be much less objectionable if it is stationary. Therefore, as the first requirement, the subcarrier frequency must be harmonically related to the line and scanning frequencies. As a second requirement, if the subcarrier frequency is high enough, the interference pattern will be less perceptible to the eye at normal viewing distances -- particularly on black and white receivers which normally have reduced high-frequency bandwidth. However, at the same time the frequency must be low enough to provide sufficient bandwidth for the chroma information.

frequency interlace

By adding a third requirement the visible effects of chroma sidebands on luminance information can be reduced still further by cancelling the *effect* of the subcarrier at the picture tube. The requirement, called frequency interlace, is illustrated in Fig. 13-4 which shows a portion of video with the subcarrier added. If the subcarrier is made an odd multiple of the line scanning frequency, a *time displacement* of one-half cycle will occur on each successive scanning line as shown in Fig. 13-5.

half-cycle displacement

An intensity-modulated dot structure will exist when chroma sidebands are present which will not be too noticeable to the eye if:

chroma-sideband effects

1. The pattern is stationary.
2. The frequency is high enough to make a closely spaced dot structure.

The final requirement is to choose a frequency whose beat with the sound carrier also cancels. Since the color subcarrier is closer to the sound carrier than to the picture carrier, any beats between the sound and color subcarrier will be relatively low frequency, producing a coarse interference pattern on the picture. As it turns out, the sound carrier is frequency modulated so the odd-multiple subcarrier frequency chosen is referenced to the unmodulated sound carrier. Maintaining an odd-multiple relationship between both the unmodulated sound carrier and the scanning frequencies (line and field) requires that either the sound carrier frequency or the scanning rates be changed slightly. Field tests

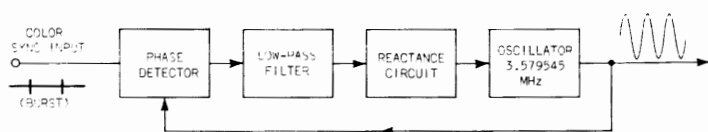


Fig. 13-6. Block diagram of basic subcarrier regenerator system.

frequency
subcarrier

indicated that a slight change of the scanning rates was a better compromise. Three different frequencies have been used in the history of color television, but these four requirements plus a few additional considerations led to a final choice of the subcarrier frequency standardized at 3.579545 MHz with the line and field scanning frequencies adjusted to 59.94 Hz and 15.734 kHz respectively.

subcarrier
regenerator
elements

As mentioned earlier, the subcarrier is not included in the composite video signal, therefore, the 3.579545-MHz subcarrier must be regenerated by a local oscillator for use in the chroma synchronous detectors. A variety of subcarrier regenerator circuits have been developed, but all the circuits have four basic components in common: (Fig. 13-6.)

1. Phase detector to compare the oscillator frequency (initially) and phase to the synchronizing color burst.
2. Oscillator circuit to regenerate a continuous 3.579545-MHz sinewave (chrominance subcarrier).
3. Low-pass filter to integrate (average) the phase detector output over a large number of burst samples. (Reduce the effects of noisy signals.)
4. A circuit whose capacitance is varied by the output voltage of the phase detector. (Oscillator frequency control)

The complexity of each of the four components just outlined is determined by the relative importance of the following objectives:

subcarrier
regenerator
requirements

1. *Pull-in range:* Since the oscillator must eventually lock the exact frequency of the incoming signal, the subcarrier regenerator system should have sufficient pull-in frequency range, typically from 100 Hz to around 500 Hz.

2. *Phase stability*: Phase stability is classified two ways; long-term phase stability called static phase, and short-term stability called dynamic phase. After the subcarrier regenerator oscillator is synchronized to the incoming sync-burst signal, good phase stability must be maintained. In the literature references are made to subjective tests indicating $2.5^\circ - 5^\circ$ as a tolerable stability for picture monitors. For vectorscopes much greater resolution is needed when accurate measurements are desired, requiring short-term phase stability of less than 0.5° .
3. *Pull-in time*: Frequency pull-in time of the oscillator should be as short as practical (usually one to five seconds).
4. *Constant amplitude*: Maintain a constant amplitude sinewave.

All four objectives are, of course, intended to permit carrier reinsertion and eventual recovery of the original video (chroma) without adding additional signals or distorting the original information. The most important objectives for vectorscope measurement applications are dynamic phase stability and constant amplitude sinewave output.

The first of the four operational blocks in the subcarrier regenerator system is the phase-detection circuit shown in Fig. 13-7. The primary purpose of the phase detector is to precisely compare the burst frequency to the oscillator frequency, and provide an output voltage proportional to any frequency difference. When describing precise frequency comparison, the detector is usually thought of as performing two functions -- frequency comparison for

precise
frequency
comparison

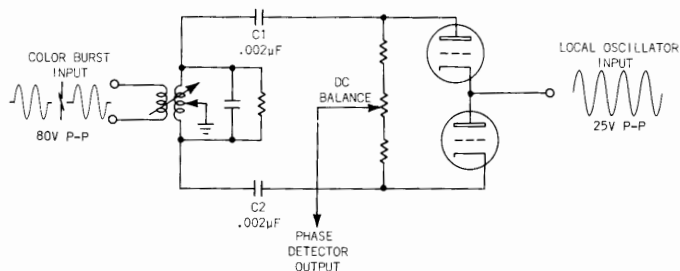


Fig. 13-7. Phase Detector circuit.

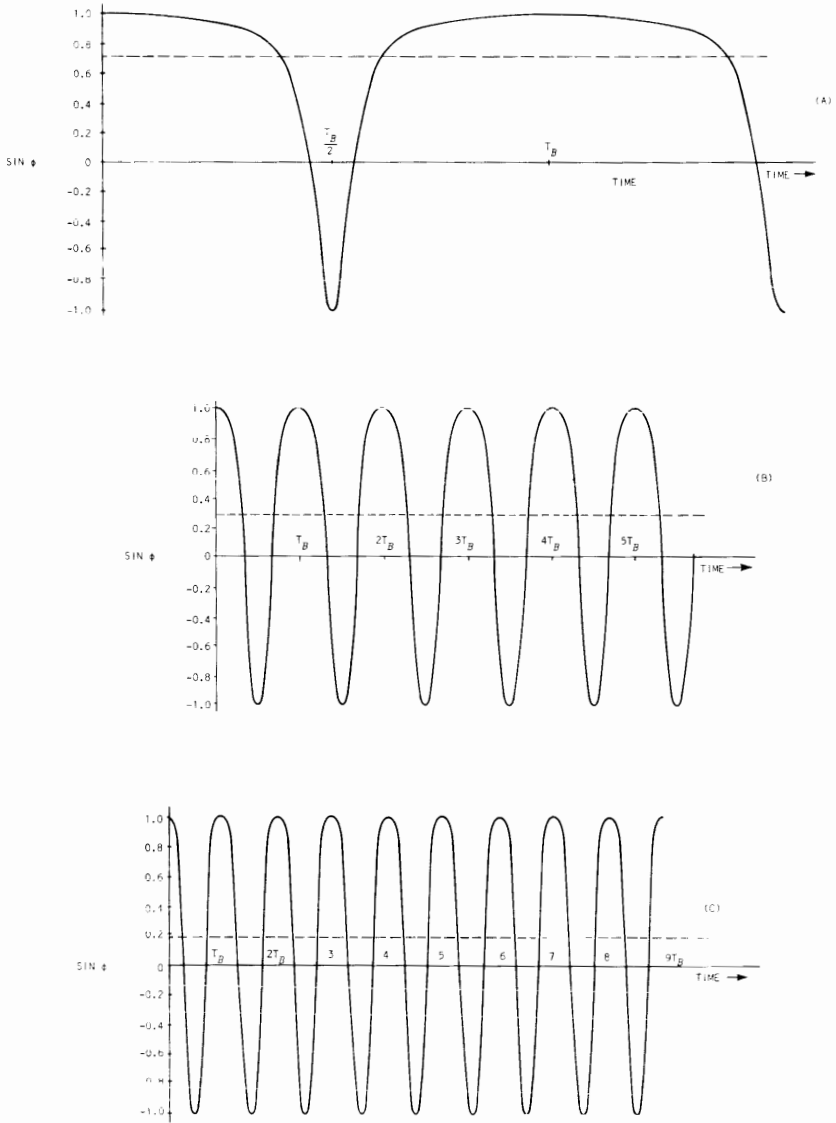


Fig. 13-8. Phase-detector output beatnote waveform for various tuning errors. T_B is the beatnote period in seconds. DC components are shown by the dotted line.

differences greater than one Hz, and phase comparison for differences less than one Hz -- principally because mathematical solutions exist for each condition.

error-signal
characteris-
tics

The phase-detector output voltage is proportional to the cosine of the phase angle *difference* between the burst input and the oscillator input and is, therefore, 0 volts when the two signals are 90° apart. If the two signals are not 90° apart a DC voltage will be developed, with a polarity determined by the direction of the phase error. As might be concluded, if the frequency difference is greater than one Hz the phase angle between the two signals will be constantly changing, producing a sinusoidal output. Such is the case if oscillator frequency lock-in does not occur. The phase detector output will be a sinewave with a frequency equal to the difference between the burst frequency and the oscillator frequency if the feedback loop is not functioning to reduce the frequency difference.

When the feedback loop is properly functioning, the oscillator frequency will continue to change until it is the same as the burst frequency. Now, the difference-frequency waveform at the phase detector output will *not* be sinusoidal, but will instead have a low-frequency beatnote and a high-frequency beatnote similar to the waveshape illustrated in Fig. 13-8. The unequal areas under the positive and negative half-cycles form a DC component used to change the oscillator frequency via the reactance circuit.

Noting whether the detector output voltage is sinusoidal, nonsinusoidal or DC can provide a valuable clue when analyzing the operation of the phase detector loop. It should also be noted that the burst repetition rate is equal to the horizontal line scanning frequency of approximately 15,734 kHz, producing the beatnote waveforms of Fig. 13-8 consisting of "samples" at the line rate.

phase
detector
sensitivity

With the waveshape output of the phase detector in mind, the next consideration is the magnitude of the detector output error voltage. Phase detector efficiency is determined by its sensitivity in volts/radian and is numerically equal to the amplitude of the smaller signal applied to the phase detector. Therefore, the burst signal applied to the detector is over twice the amplitude of the oscillator signal

to make the output sensitivity proportional to the oscillator drive, which is generally more free of noise than the incoming burst. Typical signal amplitudes for the circuit illustrated in Fig. 13-7 are 80-120 V peak burst and about 35 volts peak reference oscillator, providing a detector sensitivity of about 35 V/radian.

static
phase
error

The phase detector is an important part of the subcarrier regenerator system. The phase detector affects both the static phase error (long-term phase stability) and frequency pull-in time -- two factors considered in the circuit design. Both factors must also be verified in the actual circuit operation. The first consideration is static phase error. Accuracy of the oscillator frequency (static phase) in addition to being a function of the oscillator drift is also dependent on the accuracy of the phase detector transformer winding symmetry and the balance of the resistive network associated with the detector. Both conditions are compensated for, by adjustments in the actual circuit. The resistive balance can be verified by checking for zero volts DC at the center arm of the adjustment potentiometer with the burst signal removed and the oscillator still applied to the detector.

The impedance of the two halves of the circuit should be equal for minimum static phase error. Impedance balance can be checked by removing the oscillator signal with the burst applied to the phase detector and noting whether the detector output voltage is still close to 0 volts or at least within 1% of the smallest signal normally applied to the phase detector. Usually, though, the impedance balance is not checked because:

1. The .002 μ F C1 and C2 capacitors are matched to 1%,
2. The transformer, when adjusted for maximum amplitude, is very close to ideal balance, and
3. Any balance error that then exists causes only slight static phase error.

Both adjustments in the phase-detector circuit are intended primarily to offset the static phase error, because these errors affect the symmetry of the frequency pull-in range. The frequency pull-in range should be approximately the same above and below the

correct frequency of the oscillator. As will be observed in subsequent discussions, static phase errors as such are not of great concern since front-panel control compensation is available.

phase
detector
sensitivity

Phase-detector sensitivity is the second major consideration. The phase-detector sensitivity partially determines the frequency-holding range, but more important almost entirely determines frequency pull-in *time*. This pull-in time is effective only when the difference between the two signals is less than about 100 Hz. A variable inductor in the oscillator output adjusts a tuned circuit to resonance to ensure maximum drive to the phase detector. The frequency pull-in time should then be one to two seconds. If the difference frequency is greater than 100 Hz (approximately), the low-pass filter begins to affect the amplitude of the phase detector output signal as applied to the reactance circuit, extending the pull-in time considerably under extreme conditions.

The second of the four operational blocks in the subcarrier regenerator system is the low-pass filter shown in Fig. 13-9. The principal function of the low-pass filter is to minimize the phase jitter (dynamic phase error) that can result from noise on the incoming burst signal. The intent of the low-pass filter is to "integrate out" the random-occurring noise and average the frequency and phase information from many successive burst samples.

While integration may be a familiar process, an analogy may serve to give a clearer mental picture of the importance of averaging the information contained in the color burst. If a triggered sinewave is displayed on an oscilloscope and random background noise is added until the sinewave is no longer visible, the desired sinewave appears to be lost in the background. If a camera is attached to the oscilloscope and the shutter opened for a long period

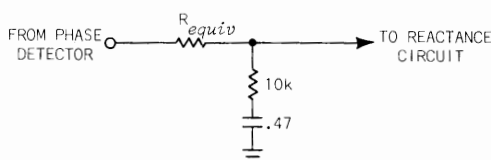


Fig. 13-9. Low-pass filter.

Integrator

of time, the original sinewave will be visible on the photograph against a fogged background. The film (ideal, nonexistent film) acts to integrate the random noise into a uniform background and integrate the nonrandom sinewave to produce the original waveshape. The longer the film is exposed, the better the sinewave visibility.

The low-pass filter acts in a similar manner by averaging the repetitive burst frequency and phase from a large number of burst samples. Because of other effects of averaging, the integration time (maximum number of burst samples needed) does have practical limitations.

The major limitation is the relationship between filter response and pull-in frequency range. The pull-in frequency range is determined by the bandwidth of the total phase control loop. The low-pass filter to a very large extent determines that bandwidth. If the integration time is too long, the bandwidth -- and therefore the pull-in range -- will be too narrow, requiring either a very stable oscillator or separate frequency- and phase-detector loops.

The efficiency of the filter is described by its phase-transfer ratio -- the ratio of AC transmission to DC transmission through the filter as shown graphically in Fig. 13-10. Notice the double curve of the filter effected by the addition of a 10-k Ω resistor in series with a 0.47- μ F capacitor. Addition of the 10-k Ω resistor permits independent choice of the two parameters affecting (1) static phase error and (2) noise bandwidth (dynamic phase

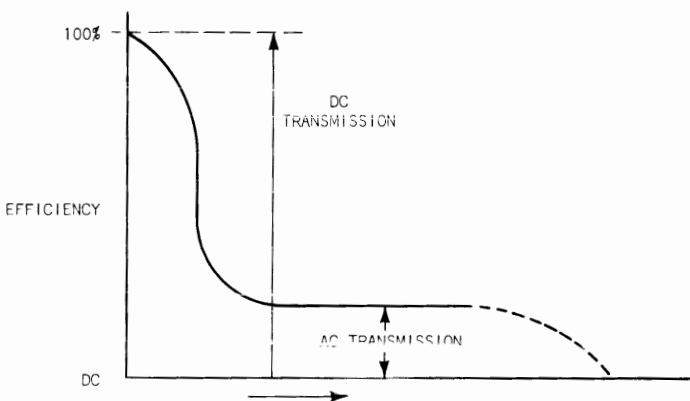


Fig. 13-10. Low-pass filter response.

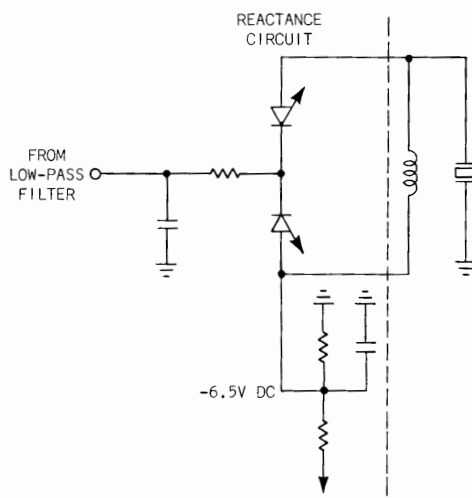


Fig. 13-11. Reactance circuit.

crystal-
controlled
oscillator
needed

error). The 10-k Ω resistor specifically affects the noise bandwidth and, therefore, the dynamic phase instability originating from noisy burst waveforms. Usually a compromise choice between the practical frequency pull-in range and the desired noise bandwidth is made. Even if a reasonable amount of phase instability can be tolerated, the compromise usually requires a more stable oscillator to keep the pull-in range requirement down to a few hundred hertz or less, therefore a crystal-controlled oscillator is commonly used.

capacitive-
sensitive
diodes

The output voltage from the filter is applied to a reactance circuit -- the third operational block -- to provide a voltage to capacitance conversion. The capacitance forms part of the oscillator resonant circuit. The reactance-circuit efficiency is described in terms of sensitivity in hertz/volt; in other words, the oscillator output frequency change for each volt applied to the reactance circuit. In the interest of economy vacuum tubes are usually used for the reactance circuit, however vacuum tubes are susceptible to filament hum. The filament hum can and sometimes does modulate the oscillator phase at a 60- or 120-Hz rate. Capacitance-sensitive diodes provide an effective solution to hum modulation as well as providing circuit simplification as can be seen in Fig. 13-11.

The reactance circuit's major effect on overall system performance is the frequency-holding range (maintaining frequency lock). The frequency-holding range in Hz is determined by phase detector sensitivity multiplied by the sensitivity of the reactance circuit. Since the reactance circuit actually forms part of the oscillator itself, reactance-circuit sensitivity is dependent not only on the capacitance-sensitive diodes, but on the oscillator circuit as well -- especially the crystal. The reactance sensitivity is dependent on two factors:

reactance
sensitivity

1. Variations between crystals.
2. Nonlinear capacitance/voltage ratio of the diodes. The capacitance change is approximately logarithmic.

Within the limits of the previous two factors, the reactance sensitivity is about 200 Hz/V. Temperature stability typically is optimized by careful choice of DC bias on the diodes.

The oscillator circuit itself is a tuned-grid, tuned-cathode type. The screen grid of the oscillator tube acts as a plate for the oscillator with electron-coupling to the actual plate of the tube. This arrangement is equivalent to an oscillator and buffer amplifier combined in one tube, making the oscillator frequency essentially independent of the load impedance without additional circuitry.

The oscillator circuit is shown in Fig. 13-12 with the reactance circuit included. The main objective of the oscillator is to provide a sinewave output at 3.579545 MHz with constant amplitude and the desired frequency stability. For increased oscillator frequency stability a high capacitance-to-inductance ratio is desirable.

The addition of a crystal will make the oscillator more stable than a conventional L/C resonant circuit for several reasons:

crystal
oscillator
characteristics

1. The electrical equivalent of the crystal provides the high capacitance-to-inductance ratio (Fig. 13-13).
2. The Q of the crystal is very high. The effective Q of the resonant circuit in the oscillator is important. The percentage change in frequency required to produce a

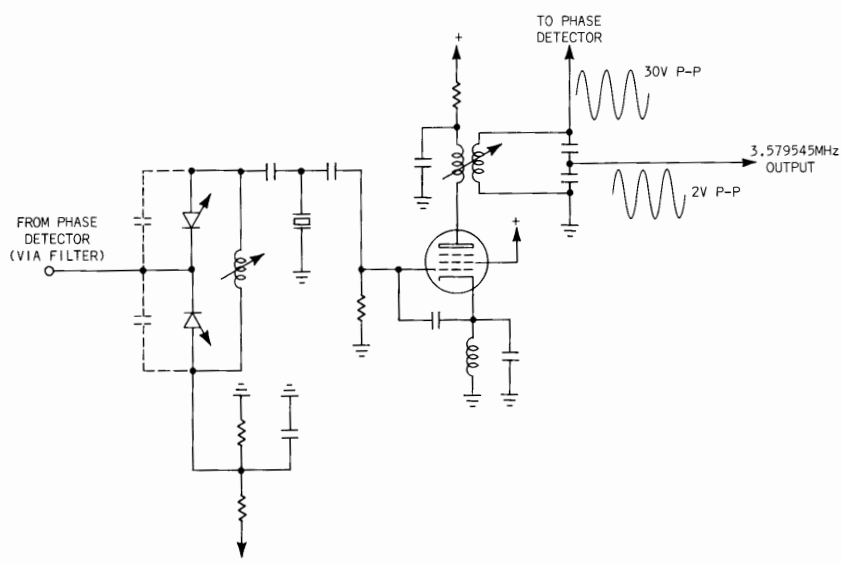


Fig. 13-12. Oscillator and reactance circuit.

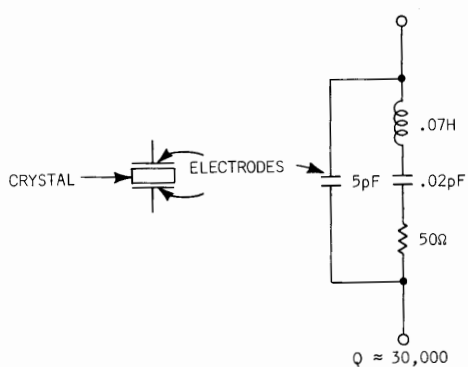


Fig. 13-13. 3.579545-MHz crystal and electrical equivalent.

correcting phase shift (due to circuit drift) between input and output voltages of the oscillator tube is inversely proportional to the effective Q. The high-Q crystal reduces the required correction to a small percentage, making the oscillator frequency more stable.

Reviewing briefly, crystals are usually operated in typical oscillator circuits either one of two ways:

1. *Series resonance*: At series resonance the reactive components of C and L are equal and the equivalent circuit becomes resistive.
2. *Antiresonance*: The antiresonant mode, commonly called the parallel resonant mode, provides a very high impedance at the resonant frequency. As the name implies, this operating mode is used in a parallel resonant circuit such as shown in Fig. 13-12 primarily because the frequency of oscillation can be adjusted by varying the load capacitance (capacitance-sensitive diodes) into which the crystal is looking. Increasing the load capacity lowers the operating frequency, while reducing the capacity raises the frequency.

The phase detector, low-pass filter, reactance circuit, and oscillator all combine to form the complete automatic phase-control loop. The entire system can be considered a dynamic integrator, because the 8-11 cycle "burst" sample of the original reference sinewave is applied to the system input and a continuous sinewave, frequency- and phase-locked to the burst sample, is generated at the output -- essentially "filling the gaps" between bursts. A locally generated replica of the original subcarrier reference frequency is then available.

Summarizing the operation of the subcarrier regenerator system, two main functions are required:

1. Oscillator frequency pull-in to the burst frequency, and
2. Stable phase-lock after frequency pull-in.

The degree to which the two functions are performed is determined by the static and dynamic phase error limits that can be tolerated in measurement vectorscope applications. Compromises of static and dynamic phase-error limits are in turn determined by the required pull-in frequency range and tolerable pull-in time.

dynamic
phase
error

Fig. 13-14 illustrates the block diagram of another subcarrier-regenerator system in which the emphasis between static and dynamic phase compromises has been shifted. The principle objective of this subcarrier-regenerator system is to achieve as low a dynamic phase error as possible. Minimum dynamic phase error (usually called jitter or phase-modulation) is very desirable in applications where high-resolution chrominance demodulators are driven by the subcarrier output because any time-modulation of the subcarrier will be observed at the demodulator output.

As mentioned earlier in the chapter, subcarrier regenerators have four basic components in common:

1. The phase detector
2. Low pass filter
3. Oscillator
4. Oscillator frequency control in the form of a reactance circuit.

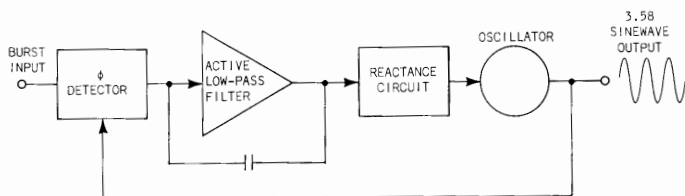


Fig. 13-14. Block diagram of subcarrier regenerator employing low-pass feedback amplifier.

Since the objective of the circuit illustrated in Fig. 13-14 is low dynamic-phase error, special attention is devoted to that portion of the overall system which has the greatest affect on dynamic phase error -- the low-pass filter.

Phase-modulating noise may exist on the burst sample which can adversely affect the phase of the oscillator during the time the burst is being sampled. The effects of noise can be reduced by narrowing the bandwidth of the low-pass filter to only pass DC and very low frequency information from the phase detector output.

However, the low-pass filter does form part of the regenerator system, so compromises of the subcarrier regenerator's overall performance must be made to achieve minimum phase jitter. The two principal compromises are:

1. Less frequency-pull-in range. The smaller pull-in range is a result of the narrower filter bandwidth, in turn requiring a very stable crystal-controlled oscillator.
2. Longer pull-in time. In effect, the reduced filter bandwidth increases the burst-sampling integration time, requiring more burst samples and therefore more time for frequency-lock.

low-pass
amplifier

Note in the block diagram of Fig. 13-14 that the low-pass filter element is a high-gain negative feedback amplifier (integration amplifier) instead of a passive RC filter. The use of a feedback amplifier as a low-pass filter minimizes to some extent the undesirable compromises previously listed.

Use of an amplifier for filtering provides the following advantages:

1. Filter efficiency (illustrated in Fig. 13-10) is greater than 100%. The additional gain at DC reduces the static phase error to a negligible figure.
2. The frequency pull-in time (within the bandwidth of the filter) is greatly reduced since the phase detector sensitivity in volts/radian is multiplied by the gain of the integrator amplifier.

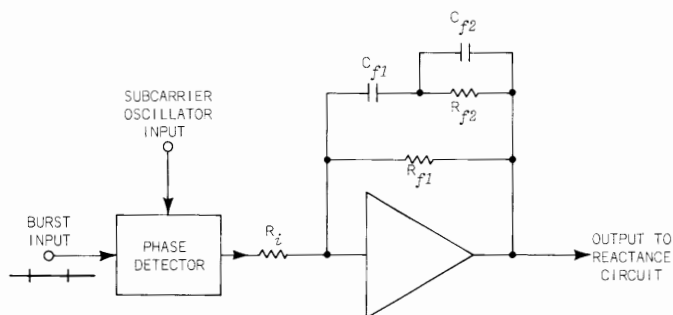


Fig. 13-15. Block diagram of phase detector and low-pass amplifier.

The block diagram of the low-pass amplifier can be seen in Fig. 13-15 in more detail. The output amplitude is essentially determined by the ratio of R_f/R_i at DC; the gain is reduced at very low frequencies by the presence of C_{f1} in series with R_{f2} . The ratio between C_{f1} and C_{f2} is about 20:1 so the amplifier gain will be relatively constant for intermediate low frequencies -- determined by the ratio of R_{f2}/R_f .

The results of the complex feedback network can graphically be seen in the illustration of Fig. 13-16. Notice that the rate-of-change from the AC transmission efficiency to the DC transmission efficiency is linear. Contrast to the conventional passive filter can be seen by observing the response characteristic of Fig. 13-10.

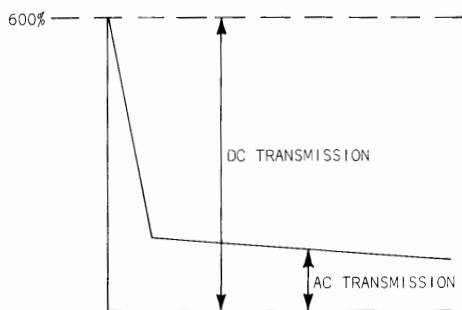
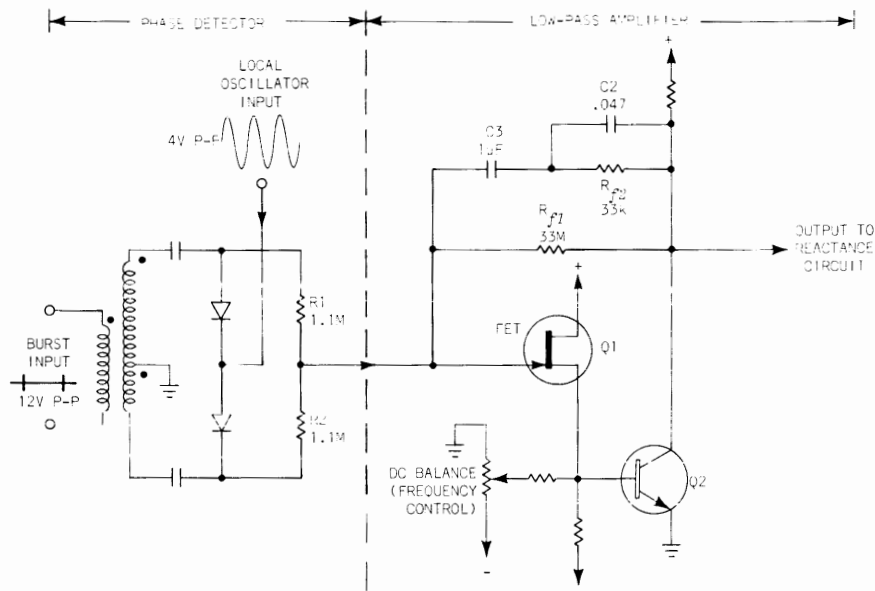
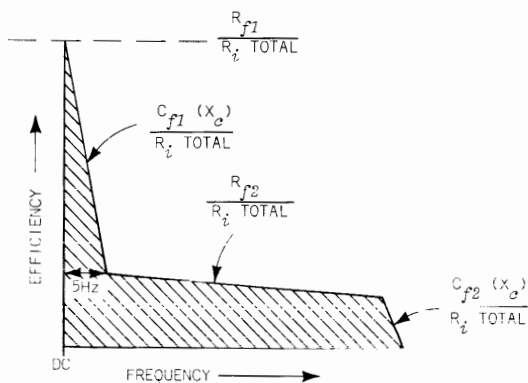


Fig. 13-16. Frequency response characteristic to low-pass amplifier.



(A) CIRCUIT



(B) GRAPHIC EFFECTS OF FEEDBACK ELEMENTS ON AMPLIFIER FREQUENCY RESPONSE

Fig. 13-17. Low-pass amplifier.

Fig. 13-17A illustrates the combined phase detector and low-pass amplifier block diagram. Except for value differences, the phase detector is essentially the same, operationally, as the phase detector described earlier in the chapter. The effectiveness of the phase detector, however, is modified by the presence of the low-pass amplifier.

Recall that the phase-detector sensitivity in volts/radian is approximately equal to the peak amplitude of the oscillator sinewave applied to the phase detector. The addition of the low-pass DC-coupled amplifier increases the phase-detector sensitivity by the DC gain of the amplifier, improving both frequency pull-in time and static phase error (vector displacement on the CRT as a result of oscillator frequency drift).

Fig. 13-17A illustrates the actual circuit of the phase detector and low-pass amplifier. The FET input stage of the amplifier provides very high input impedance for the feedback network, and the grounded-emitter transistor stage (Q2) provides high voltage gain. The phase-detector resistors, R1 and R2 in parallel, form the R_i of the feedback amplifier. The effect of the various components on the feedback loop are shown on the response curve of Fig. 13-17B.

The effective phase-detector sensitivity with the amplifier is about 120 volts/radian resulting in a static phase error of about $0.01^\circ/\text{Hz}$.

The low-pass filter efficiency of the feedback amplifier is about 600% resulting in a dynamic phase error of less than 1° of jitter with 30% white noise on the burst.

These figures have more meaning when contrasted to the conventional figures of over 2° static phase error and over 1° of dynamic phase error in most conventional designs employed in picture monitors.

The reactance circuit is similar to the one described earlier in the chapter except for a few minor differences:

1. Only one variable-capacitance diode is used.
2. An inductor is not used in parallel with the diode.

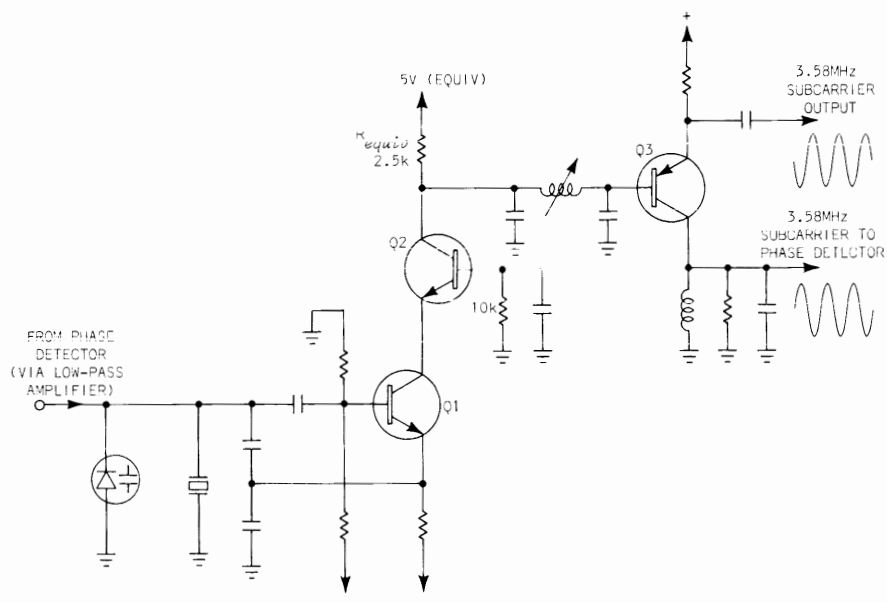


Fig. 13-18. 3.579545-MHz oscillator circuit.

capacitive-
sensitive
diodes

The variable-capacitance diode is about 6 pF/volt, making the reactance sensitivity about 50 Hz/volt. While the reactance sensitivity is comparatively low, the high phase-detector sensitivity multiplied by the reactance sensitivity results in a respectable frequency-holding range of about 100 Hz.

The oscillator is a Colpitts configuration operating Class C. Notice in the circuit illustration of Fig. 13-18 the oscillator parallel-resonant circuit does not contain a variable inductor. The Q of the crystal is extremely high (an order of magnitude higher than the crystal used in the previous circuit) so the capacitance-to-inductance ratio is sufficient without an external inductor.

The oscillator drives a buffer stage to provide isolation between the oscillator and the load. Since the oscillator is operating Class C waveform distortion exists, generating harmonics of the oscillating frequency. A π -filter between the buffer amplifier and the output emitter follower removes the harmonic content from the waveform.

The output amplifier Q3 is split loaded, performing as an emitter follower and a parallel-resonant collector amplifier.

The emitter follower provides a low-impedance output to drive the subcarrier-processing circuitry while the collector drives the subcarrier-regenerator phase detector.

The parallel-resonant collector load of Q3 is more desirable than a resistive load because:

1. A DC voltage close to ground is needed.
2. AC voltage gain is also needed since the detector sensitivity is directly related to the amplitude of the signal at the collector of Q3.

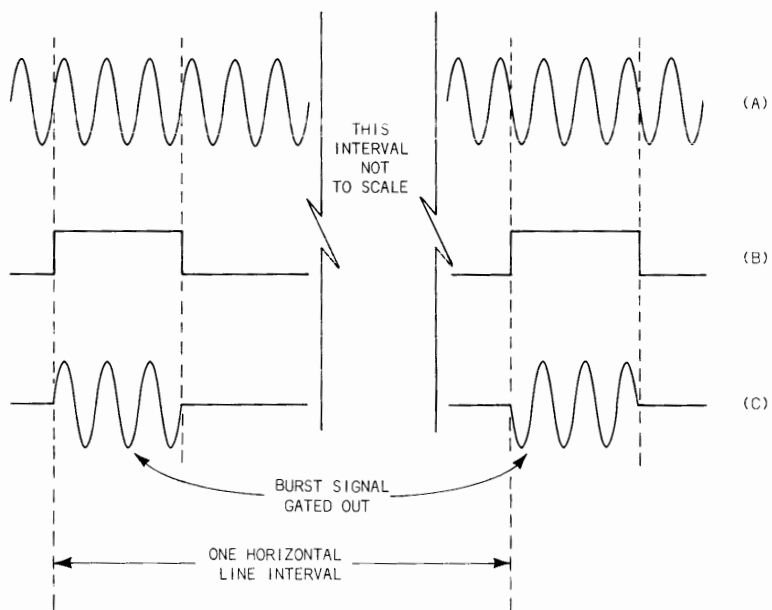


Fig. 14-1. Formation of color burst waveform.

14

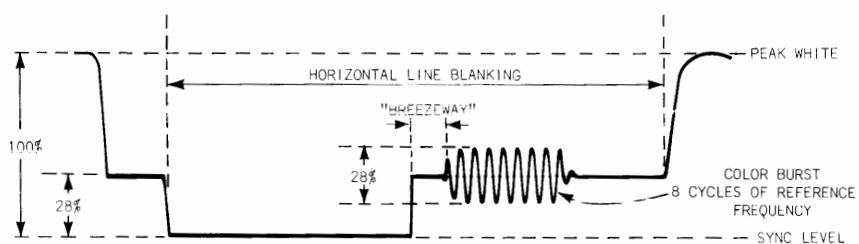
CHROMINANCE SYNCHRONIZING

reference
oscillator
function

The reference oscillator must produce a sinewave output locked to the burst, and the burst only. Therefore, some method must be provided to exclude all other chrominance information from the phase detector. The method used to exclude chrominance from the phase detector will depend to a certain extent on the information contained in the burst.

Basically, the burst is a gated-in sample of the original reference generator used to develop the two phase-displaced color subcarriers as shown in Fig. 14-1. The two subcarriers are, of course, the same frequency since they are both generated from the reference oscillator. A definite phase relationship between the phase angle of the reference sinewave and either of the two subcarriers exists. The burst must then convey several useful bits of information:

1. Being a sample of the original reference oscillator, the *frequency* of the burst is 3.579545 MHz.
2. The *phase* must then be the same as the original reference oscillator, since the burst is a gated sample of the reference oscillator. The previous statement would be true under ideal theoretical conditions, but for practical reasons not important to this discussion the burst phase is minus 180° . The effect on the vectorscope is a burst reference vector on the minus-X axis; the effect on a color picture tube is to produce green vertical lines if line-blanking is not adequate. Therefore, when choosing the ideal color burst phase for both maximum convenience and minimum luminance, a 180° phase shift from the reference sinewave was found to be the best compromise. Even so,



The burst is not transmitted during field sync and equalizing pulses.

Fig. 14-2. Location of the color sync burst on the back porch of the line sync pulse.

color-burst
phasing

the burst will appear on the picture monitor as an unexpected low-luminance green color under very adverse conditions. In practice, however, the luminance produced by burst is not a practical problem. The phase of the burst with respect to the original reference oscillator is 180° .

burst and
harmonics

3. The peak-to-peak *amplitude* of the burst is made equal to the peak amplitude of the sync pulse (Fig. 14-2). Since the chroma information is essentially a "difference" signal, an absolute amplitude reference at the receiving end is needed for comparison to the chroma difference signals.
4. The *repetition rate* of the burst sample is one H or 15734 Hz. Since the burst waveform is gated at the line-scanning rate, harmonic multiples of the line frequency are produced. The result is a burst of sinewaves at 3.579545 MHz with sidebands spaced 15734 H apart. In the strict sense, sidebands of the burst exist down to the field frequency (since burst is omitted on field sync pulses and equalizing pulses), but generally 85% of the total energy of the burst is contained in the sub-carrier component and the first twenty sidebands on each side of the sub-carrier.

The question may arise: Where is the spectrum space for all the sideband information of the burst when frequency spectrum conservation is the objective in the color system? The answer lies in the fact that the spectrum is "time-shared" or "time-multiplexed" with the chroma information. Burst always occurs in time when chroma is never transmitted by being gated onto the back porch as illustrated in Fig. 14-2.

The considerations of the burst waveform itself outline the needs of the circuitry required to exclude the chroma information from the phase detector. Since the burst is gated onto the composite video initially, the circuitry at the receiver end of the system commonly takes a similar form.

Therefore, signal-gating the subcarrier signal into the reference oscillator system only during the occurrence of burst is the normal technique used. Signal gating can present problems, however, unless several precautions are observed:

signal
gating,
ringing,
noise

1. Since inductive circuits are usually used, care must be taken that the gating waveform does not ring the resonant circuits, otherwise phase errors -- either static errors, or dynamic errors (jitter) -- may result. The gate pulse rise and fall time should be slow enough so that the harmonic content of the pulse in the 3.58 MHz region is at a low energy level.
2. The gate should remain "open" long enough to allow most of the burst to pass, but the gate should not be too long otherwise a greater amount of noise (if present) will be allowed through the gate along with burst.

Since the burst must eventually be applied to the subcarrier-regenerator phase-detector circuit, voltage amplification will be necessary. The main requirement of the burst amplifier is that the amplifier bandwidth be sufficient to amplify the burst without distortion. As shown in Fig. 14-3, the burst gating can be done in the burst-amplifier stage. The gate circuit is arranged to return the grid circuit of the burst amplifier to ground through a low impedance source rather than alternately turn the amplifier on and off. As a further precaution, the gate pulses are applied push-pull to prevent any part of the gate pulse (particularly the leading and trailing edges of the gate pulse) to be "seen" by the burst amplifier.

gating
circuit

The gating circuit is shown in Fig. 14-4. The gate pulse is simultaneously applied to one-half of the diode-bridge gate and an inverting amplifier. Quiescently, R1 and R2 forward-bias the diode bridge to maintain a low impedance path to ground (essentially the forward resistance of the diodes). When the burst gate reverse-biases the diode bridge coincident with the arrival of burst, the signal at point A is developed across the coil, since the diodes in parallel with the coil are now reverse-biased.

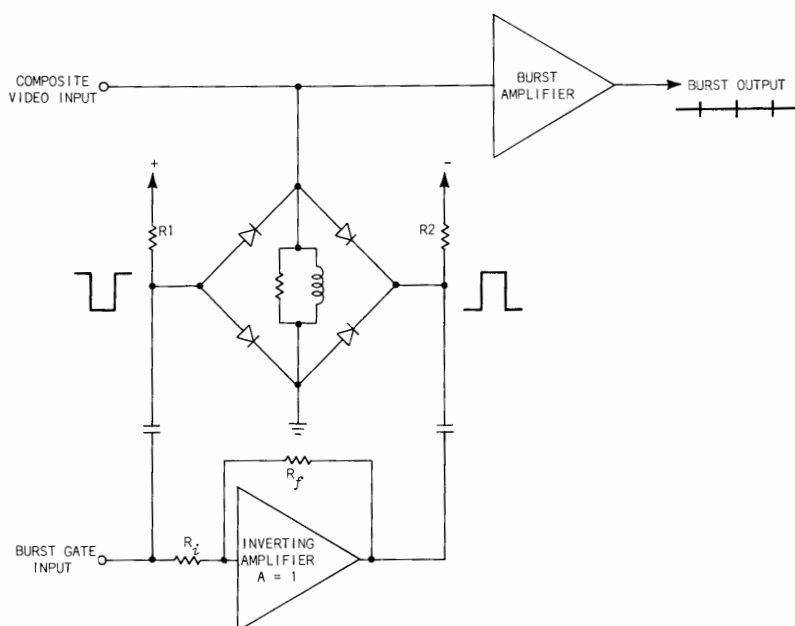


Fig. 14-3. Functional block diagram of burst gating and amplifier circuit.

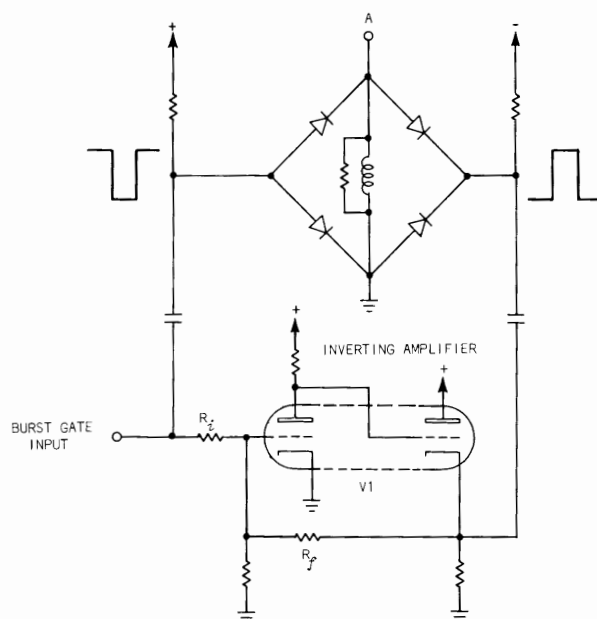


Fig. 14-4. Burst gating circuit.

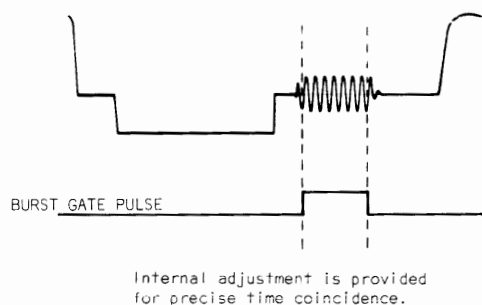


Fig. 14-5. Coincidence of the color burst and the burst gate.

gate
pulse

The gate pulse itself is a delayed pulse approximately $2\text{-}\mu\text{s}$ wide with a risetime of about $0.2\text{ }\mu\text{s}$, slow enough to prevent ringing if pulse amplitude unbalance occurs. The coincidence of the burst gate pulse to the occurrence of the color burst is illustrated in Fig. 14-5.

The pulse-inverting amplifier V1 has essentially unity gain which is determined almost entirely by the ratio of R_f to R_i , therefore tube aging does not affect pulse amplitude unbalance.

SUBCARRIER PROCESSING CIRCUITS

After regeneration of a clean subcarrier sinewave, further processing is required before carrier reinsertion into the chrominance sidebands can take place in the product demodulators. The principal objective of the subcarrier processing circuits is to derive *two* phase-displaced subcarriers precisely 90° apart. The secondary objective of the subcarrier processing circuits is to provide calibrated phase variations for vector-measurement purposes.

Two types of phase adjustments are needed:

quadrature
adjustment

1. Adjustment of the phase displacement *between* the two subcarriers is necessary. Normally, the two subcarriers are 90° apart and, therefore, are in quadrature. Measurement accuracy of the resultant vectors displayed on the CRT is dependent on the quadrature relationship of the two subcarriers. A front-panel phase control is provided to precisely adjust the relative phase of the two subcarriers to 90°. The phase adjustment affects only *one* subcarrier; the second subcarrier remains at a fixed phase with respect to the incoming color signal.

subcarrier
adjustment

2. Simultaneous phase adjustment of *both* subcarriers with respect to the incoming composite color signal is necessary. This phase adjustment has the CRT display effect of a "rotational positioning control" because the phase of both vertical and horizontal demodulator subcarriers is changed simultaneously. The adjustment is also provided on the front panel for precise vector positioning, especially the color burst reference vector.

A color picture monitor has a similar front-panel control labeled "hue" or "tint."

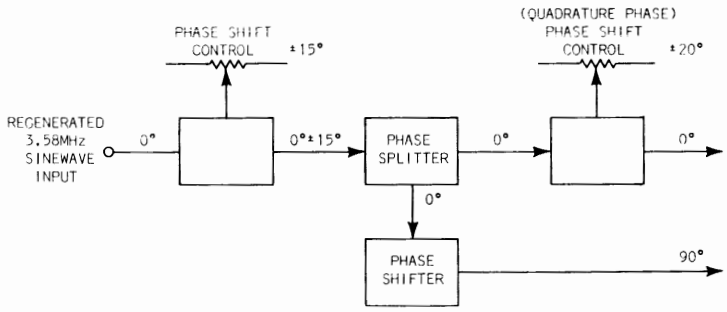


Fig. 15-1. Simplified block diagram of subcarrier phase-processing (phase splitting) system.

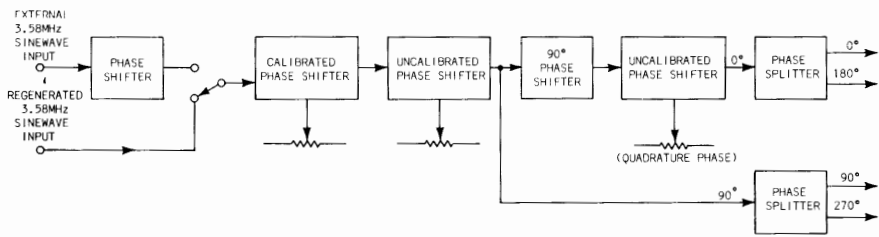


Fig. 15-2. Block diagram of complete subcarrier processing system illustrating phase adjustments.

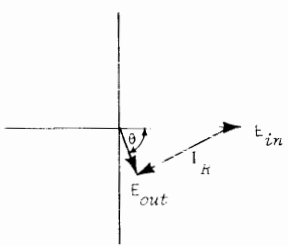


Fig. 15-3. RC phase shift network.

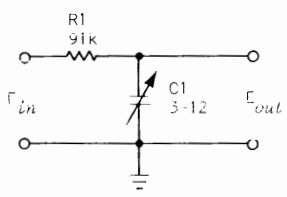


Fig. 15-4. RC phase shift network.

The basic block diagram illustrating the concept of operation of the subcarrier processing circuit is shown in Fig. 15-1. Notice the two phase-shift control blocks; both controls are identical in operation but different in their effect on the displayed waveform, primarily because of their location in the circuitry. One phase-shift control is prior to the phase splitter; any phase-change adjustment affects both the 0° and the 90° subcarrier. The other phase shift control is after the phase splitter to adjust the 0° subcarrier relative to the 90° subcarrier.

subcarrier
processing

A more complete subcarrier processing system is shown in the block diagram of Fig. 15-2. To avoid confusion, reference to the 0° phase shift control will be subsequently called the Quadrature control. Notice the four different phase outputs. An earlier chapter made reference to the push-pull demodulators which require subcarriers 180° apart for each half of the push-pull demodulator. Therefore, the 0° and 90° subcarriers are further split into two 180° components.

Four different types of phase shift networks are used in the subcarrier-regenerator system, each with a different phase shift requirement.

The four network types are:

phase
shift
networks

1. Simple RC network.
2. Resonant circuit.
3. Artificial transmission (delay) line.
4. Inductive goniometer.

The operation of the RC phase-shift network can be described any number of ways, but the simplest way can be illustrated as shown in Fig. 15-3. R and X_C are vectorially drawn to show the two components of the RC network shown in Fig. 15-4.

Since X_C is much smaller than R at the subcarrier frequency, the phase shift between the input and output will be large -- numerically, the phase angle of E_{out} will be $\arctan X_C/R$. The output amplitude will be small because the signal is taken across the capacitor and will, therefore, be $E_{in} \frac{X_C}{\sqrt{R^2 + X_C^2}}$.

quadrature

Of course the phase shift will not be precisely 90° because the ratio between R and X_C can never be

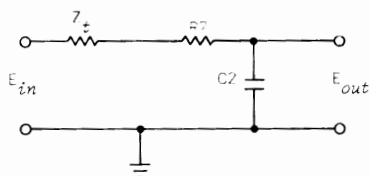


Fig. 15-5. Equivalent circuit of cascaded RC network.

infinite but the angle will approach 90° . In order to more closely approach 90° a second RC phase shift network is used in series with the first network. The equivalent circuit of the second phase-shift network is modified by the presence of the first phase network as shown in Fig. 15-5. The final result is a phase shift of 90° , adjustable $\pm$ about 3° , between E_{in} and E_{out} . The actual circuit is shown in Fig. 15-6. The principal advantage of the RC phase-shift network is simplicity; the principal disadvantage is the amplitude attenuation by the RC network and the amplitude variation when either component of the RC network is adjusted. As a result, use of the RC network is usually limited to applications not requiring front panel control of the RC network.

resonant
phase-
shift
network

The parallel resonant phase-shift network is useful for phase-shift changes when adjustments from the front panel are needed because:

1. The amplitude of the output sinewave is not as seriously altered when the resonant circuit is slightly detuned.
2. The circuit can provide a symmetrical phase shift (either a leading or a lagging phase change).

Looking at the universal curve of Fig. 15-7, the phase shift at 3.58 MHz can be seen for the parallel resonant circuit shown in Fig. 15-8.

Observing the universal chart in Fig. 15-7, notice that a $\pm 20^\circ$ phase change will also cause about a 20% amplitude change. An amplitude change of the subcarrier will affect the amplitude of the CRT display. Therefore, if the parallel resonant phase shift network is to be useful the amplitude variation due to detuning the resonant circuit will have to be minimized. Reducing the Q of the resonant

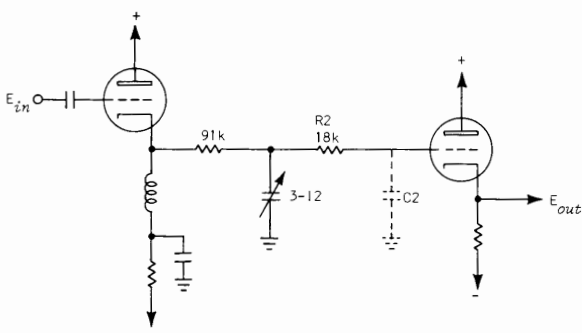


Fig. 15-6. Actual circuit of RC phase-shift network.

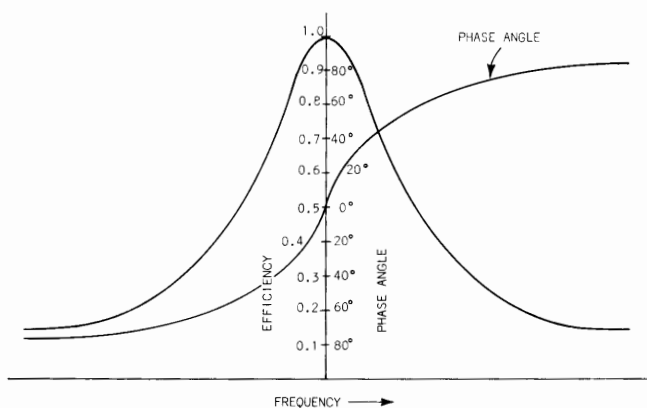


Fig. 15-7. Universal curve of parallel-resonant circuit.

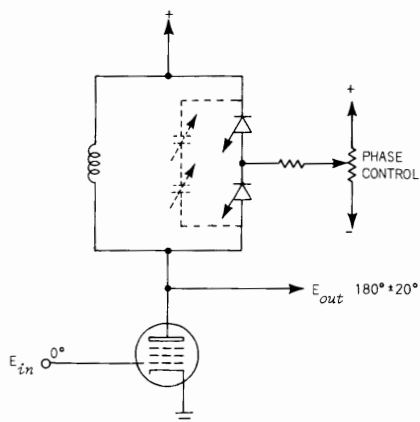


Fig. 15-8. Parallel-resonant phase shift circuit.

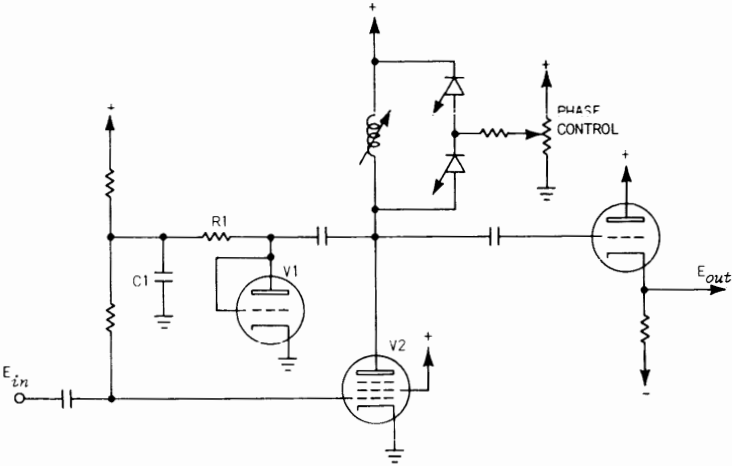


Fig. 15-9. Adjustable-phase parallel-resonant circuit.

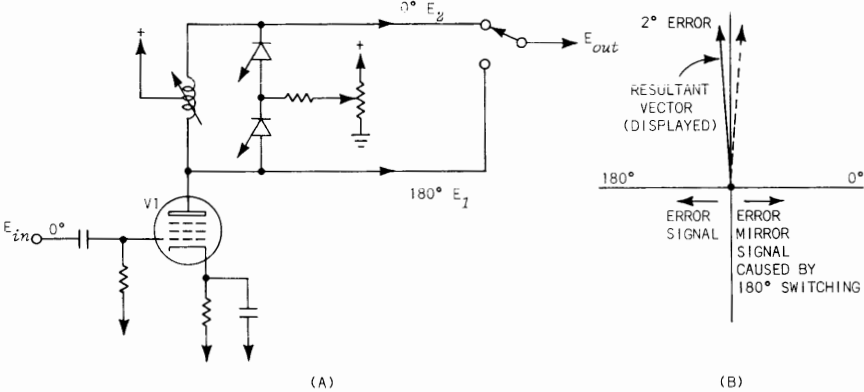


Fig. 15-10. Adjustable-phase parallel-resonant circuit with inverted or noninverted output phase selection.

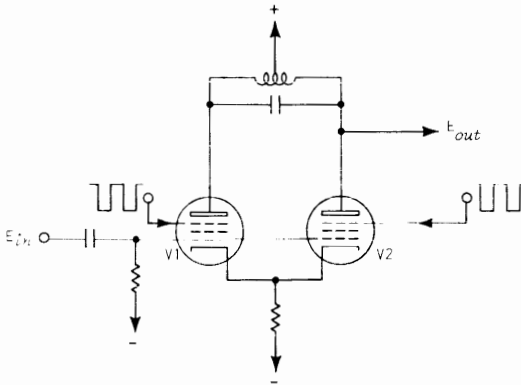


Fig. 15-11. Electronic switch for output phase inversion selection.

circuit will reduce the amplitude change, along with a reduction in phase change. Therefore, to maintain sufficient phase change with minimum amplitude variation, AGC is used as shown in the actual circuit of Fig. 15-9. With the AGC loop the amplitude variation is reduced to a less objectional $\pm 5\%$.

Looking at the circuit of Fig. 15-9, V1, a diode-wired triode, clamps the positive peaks of the sinewave to ground. The sinewave then develops a negative *average* voltage at the junction of R1 and C1 which is used to control the bias of V2. When the peak-to-peak amplitude of the sinewave appearing at the plate of V2 is reduced (by detuning the resonant circuit), the grid bias is made more positive to increase the plate current.

The use of variable-capacitance diodes to form part of the resonant circuit facilitates the use of potentiometers at the front panel, rather than using mechanical means to change the resonant circuit.

A similar circuit is illustrated in Fig. 15-10A except that provision is made for selecting either an inverted or noninverted output signal.

The principal advantage of alternating the output phase 180° is a substantial resolution increase when verifying the absolute phase of the $E_{out\ subcarrier}$ with respect to the quadrature (90°) subcarrier. (See Fig. 15-1.) For example: A one-degree absolute phase error in the subcarrier processing circuitry when displayed on the CRT can be difficult to detect when comparing to a 90° coordinate. Alternately switching one subcarrier 180° results in a mirror display on the CRT. Any phase-quadrature error is then displayed *differentially*, making observation and corrective adjustment relatively easy as shown in Fig. 15-10B. The phase of E_1 is 180° with respect to the E_{in} ; E_2 is 0° with the input signal. The phase of E_{out} with respect to E_{in} then can be selected either in-phase or 180° out-of-phase with the input signal.

The phase of E_{out} can be *electronically* switched by adding a second vacuum tube as shown in Fig. 15-11. By alternating the plate current from V1 to V2, the

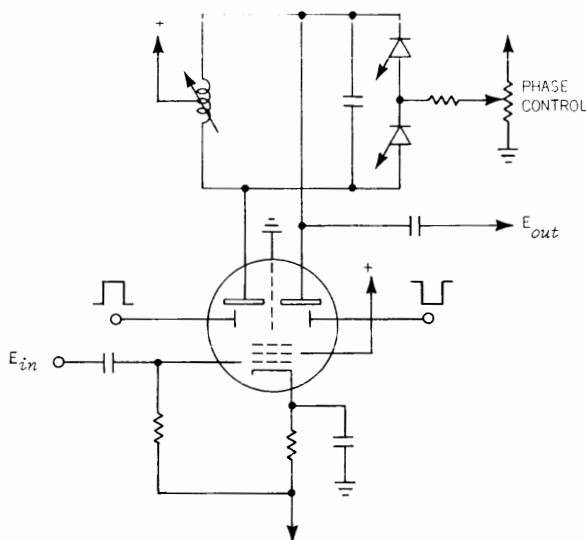


Fig. 15-12. Actual circuit of 180° phase switching.

phase of E_{out} can be alternately switched 180°. E_{out} will be 0° when V1 is conducting and 180° when V2 is conducting. The 180° switching is intended only as a test procedure; the actual circuit is arranged to switch only when an internal test signal is applied to the CRT.

The actual circuit is shown in Fig. 15-12. The circuit is simplified by the use of a special-application pentode vacuum tube. Rather than using two vacuum tubes, the plate is split into two parts and a pair of deflection plates added to deflect the electron stream from one plate to the other.

A push-pull squarewave is applied to the deflection plates when the test signal is being used to verify the quadrature phase relationship of the two 3.58 MHz subcarriers. A differential DC voltage (one plate positive and the other plate negative) is applied to the deflection plates when a composite color signal is displayed on the CRT. Under the latter condition, 180° phase switching does not take place and the operation of the phase-shift network is virtually the

same as the circuit described in Fig. 15-9. When a composite color signal is displayed, the circuit is arranged to allow plate current on the right hand plate only (Fig. 15-12).

delay-line
phase shift
network

The third form of phase-shift network is the artificial transmission (delay) line.

The principal advantages of the delay line when compared to either the RC network or the parallel-resonant circuit are:

1. Phase shifts up to 360° .
2. Virtually constant output amplitude.
3. Either continuous or "stepped" phase shifts.

Fig. 15-13 shows a transmission line used as a phase-shift network. The total electrical length of the transmission closely approximates the time of one complete cycle at 3.579545 MHz, the subcarrier frequency. The parameters of each T-section of the line are arranged so that the velocity of propagation is about 23 ns for each section or 30° at 3.58 MHz. By selecting the appropriate tap along the delay line, any desired phase shift between E_{in} and E_{out} is electrically possible. Tapping a delay line can cause undesirable phase shifts if steps are not taken to make the tap look resistive. The phase shifts due to circuit loading are not too critical for the particular circuit application since the 30° phase steps are intended to roughly locate the displayed

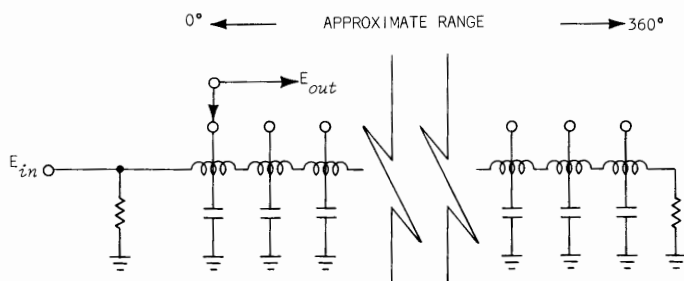


Fig. 15-13. Artificial transmission line phase-shift network.

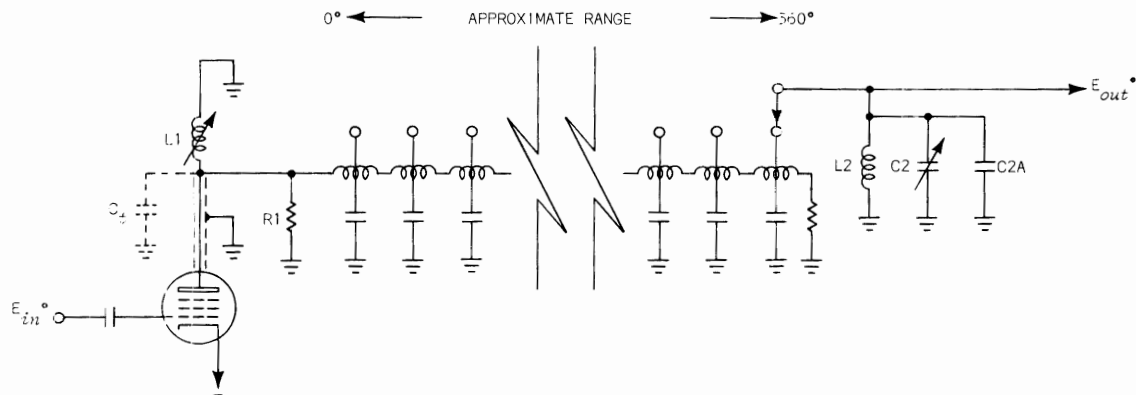


Fig. 15-14. Actual circuit of phase-shift system.

vectors. Nevertheless, as shown in the actual circuit of Fig. 15-14, the phase shifting delay line is isolated from external circuitry. L_1 and C_t , the combined vacuum tube and coaxial cable capacitance, form a resonant circuit to "look" resistive into the delay line. The vacuum tube amplifier isolates the external incoming signal (3.58 MHz sinewave) from the delay line. The parallel-resonant circuit, L_2 and C_2 , is intended to make the delay line center tap essentially resistive at 3.58 MHz.

A continuously variable phase-shifting delay line operates basically the same as the fixed delay line except that the switchable taps are replaced with a wiper arm. The line is spirally wound so that mechanical rotation of a shaft will move the tap from one end of the delay line to the other. For measurement purposes the variable delay line is equipped with a phase-readout dial knob, therefore, an RC phase-shift network is added *in series* with the delay line to match the actual total end-to-end electrical delay to the readout dial. The simplified circuit is shown in Fig. 15-15 -- a delay-line phase-shift network in series with an RC phase-shift network. The unique part of the circuit is the fact that the capacitive portion of the RC network is made to vary with the positioning of the delay-line wiper. In effect the RC network is nonlinear; the network is made purely resistive (at 3.58 MHz) at one end of the delay line, progressing to a maximum phase-shift effect at the other end of the delay line. If the

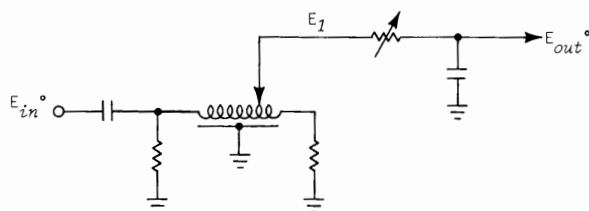
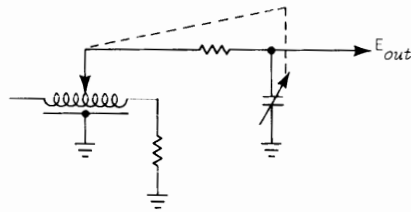
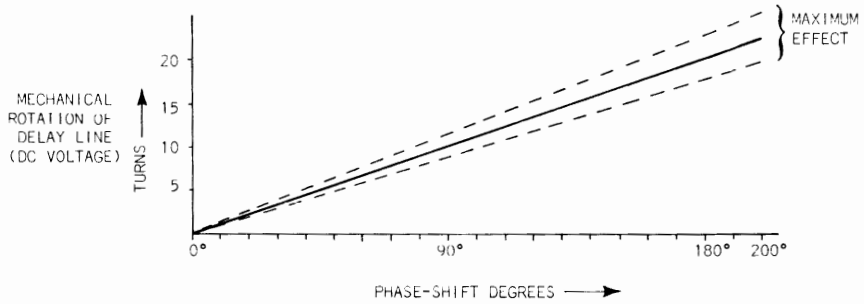
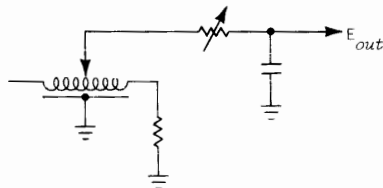
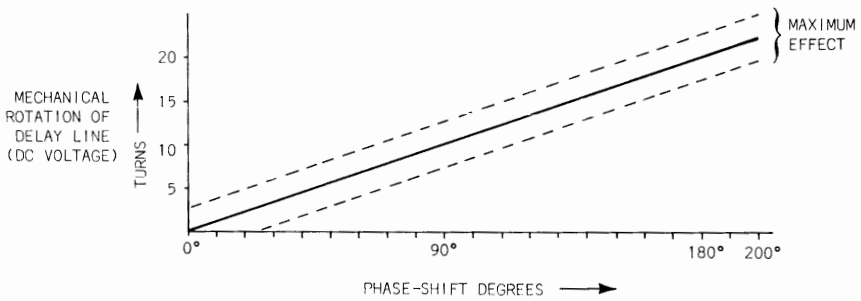


Fig. 15-15. Equivalent circuit of delay-line phase-shift network in series with a conventional RC network.



(A) EFFECT OF A NONLINEAR RC NETWORK



(B) EFFECT OF A CONVENTIONAL RC NETWORK

Fig. 15 16. Simplified diagram of RC delay-line network and graphical effect of the network.

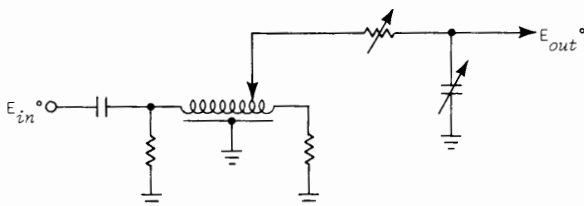


Fig. 15-17. Equivalent circuit of delay line phase-shift network in series with a nonlinear RC network.

RC network were not made resistive at one end of the delay line, adjustment of the *total* phase shift by means of the RC network would be difficult. (See Fig. 15-16.) As shown in the equivalent circuit of Fig. 15-17, *both* components of the RC network are made variable; the resistance variable is a mechanical adjustment and the capacitance is made variable electrically. The capacitor is actually a variable-capacitance diode biased by the differential DC voltage which is developed across the tapped portion of the delay line. The equivalent DC circuit is shown in Fig. 15-18. Since the delay line has resistance (about $10\ \Omega$), a usable voltage drop will exist if enough DC current flows through the delay line. The maximum differential voltage across the delay line is relatively small (about 200 mV DC) but enough to detune C1 from resonance when the wiper arm of the delay line is moved from one end to the other. The resonant circuit C1-L1, when

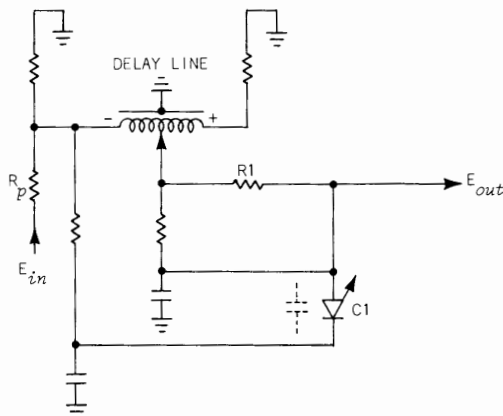


Fig. 15-18. Equivalent DC circuit of calibrated phase-shift network.

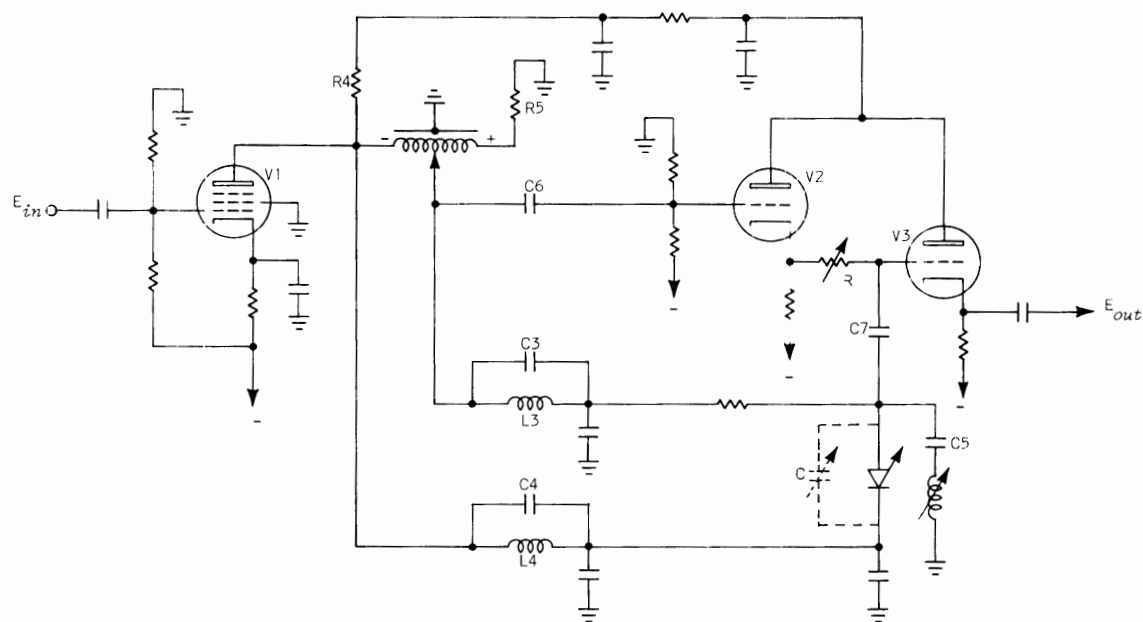


Fig. 15-19. Actual circuit of complete calibrated phase-shift network.

progressively detuned, becomes capacitive to form the "capacitor" portion of the nonlinear phase-shift network R1-C1. The nonlinear RC network is intended only to provide a means of easily adjusting the total electrical phase shift to match the mechanical rotation of the readout dial. The circuit does not make phase comparisons to optimize the incremental phase linearity (small phase-angle intervals).

Highlights of the actual circuit shown in Fig. 15-19, are:

1. The amplifier, V1, provides isolation between the input signal and the delay line.
2. R4 and R5 terminate the delay line in its characteristic impedance to prevent standing waves (and resultant phase error).
3. C5, C6, and C7 are essentially coupling capacitors which, for practical purposes, can be excluded from the equivalent circuit.
4. L4, C4, and L3-C3 form resonant circuits to remove all 3.58 MHz from the DC voltage applied to the varactor C1. Additional 3.58 MHz filtering in the form of a π -filter (C10-R10-C11) is added in series with L4 and C4.
5. The cathode-followers V2 and V3 provide isolation and a low-impedance source to drive the RC network and circuit output respectively. Notice in Fig. 15-19 that DC plate current of the cathode-followers is through the delay line to provide additional end-to-end DC voltage across the delay line.
6. Phase shift with respect to the input is 180° at minimum delay, progressing to 360° at maximum delay. The total phase shift is adjustable about $\pm 3^\circ$ by the RC network. The dial actually reads 0° at *maximum* delay so that the displayed vectors will rotate clockwise. (Normally time delay is a phase lag, -- causing the vectors to rotate counterclockwise.)

parallel-
resonant
calibrated
phase
shifter

Another calibrated phase shift network utilizes the parallel-resonant circuit rather than the combination delay line and RC network.

As noted earlier in the chapter, the resonant circuit does not have the phase-shift range in

degrees afforded by the delay line but the mechanical variation of one of the resonant circuit components can be made to conform linearly to the readout dial.

The numerical reading at any point on the dial can then be made to correspond with the actual electrical phase shift.

Fig. 15-20 illustrates how the calibrated phase shift network and the dial are coupled. The mechanical rotation of a capacitor is attached to a readout dial; rotating the capacitor alters the dial reading, detunes the resonant circuit, and shifts the phase of the output sinewave.

Assuming that the dial and capacitance vary directly, if the dial is matched to the actual phase shift at three points the dial reading will conform to the electrical phase shift at any other point. The three initial points of reference on the dial are:

1. One end of the dial.
2. Middle of the dial.
3. Other end of the dial.

construction
techniques

To facilitate making accurate measurements directly from the dial, one desirable feature is to have small increments on the dial. However, a large phase shift range in small increments requires a very large dial. To conserve valuable panel space, a "ribbon dial" is used.

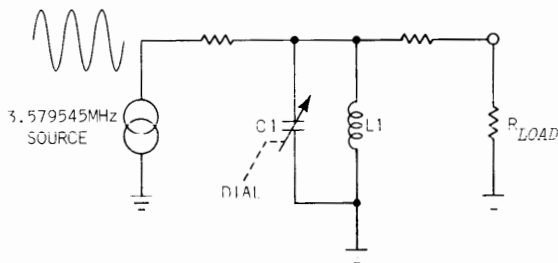


Fig. 15-20. Equivalent circuit of parallel-resonant calibrated phase network.

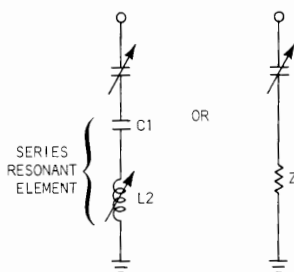
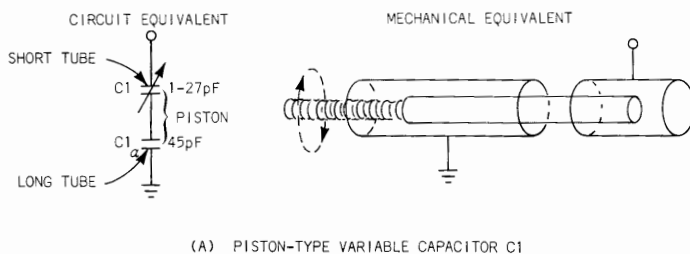


Fig. 15-21. Calibrated phase-shift network.

The ribbon dial requires more than a 360° -rotation capacitor so a piston-type variable capacitor is used. The equivalent capacitance consists of a 45-pF fixed capacitor in series with a 1-27 pF variable capacitor as shown in Fig. 15-21A. The two capacitors in series cause the total capacitance change to be nonlinear when the variable portion of the capacitor is changed, making direct coupling with a linear dial difficult. The addition of an inductor in series with the capacitor, as shown in Fig. 15-21B, forms a series-resonant circuit with the 45-pF fixed capacitor. The 45 pF then becomes resistive at 3.58 MHz, eliminating the undesirable effects of nonlinearity.

The phase-shift *range* is determined by two factors:

phase-shift
range

1. The desired phase change is $\pm 15^\circ$; therefore, a fixed capacitor is added in parallel with the variable capacitor to make the variable piston capacitor a small portion of the total capacitance.

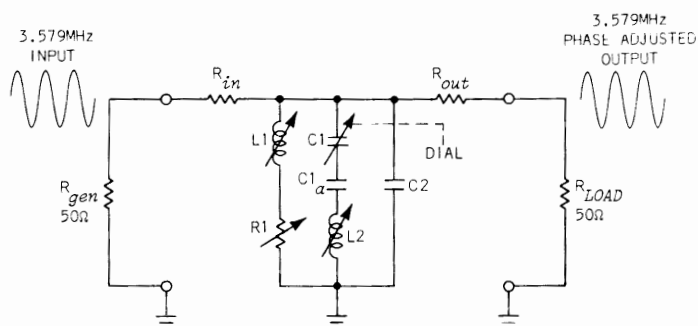


Fig. 15-22. Calibrated parallel-resonant phase-shift network.

2. Since the phase-shift range must be calibrated to match the dial, the Q of the resonant circuit is precisely established by a variable resistance placed in series with the inductor. An adjustable series resistance is more effective than a resistor placed in parallel with the inductor because the stray capacitance of the potentiometer does not affect the resonant circuit.

practical
circuit

In contrast to the equivalent circuit of Fig. 15-20, the addition of components needed to satisfy these practical considerations of the calibrated parallel-resonant circuit can be seen in the complete circuit of Fig. 15-22.

The three circuit adjustments calibrate the actual electrical phase shift of the circuit to the dial.

Observing the universal curve of the parallel-resonant circuit in Fig. 15-23 the adjustments affect the following:

1. $L1$ adjusts the circuit to resonance at 3.58 MHz; the graphical effect is a horizontal displacement of the phase-shift curve. The dial is positioned to the center (0°) for this adjustment.

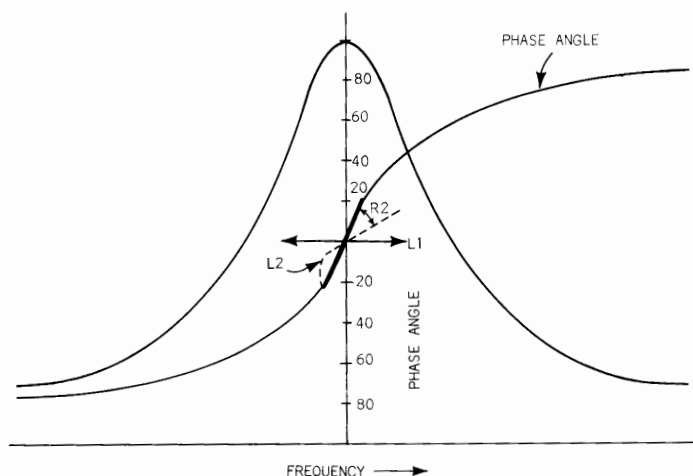


Fig. 15-23. Graphical effects of the circuit adjustments shown in Fig. 15-22.

2. R_1 adjusts the Q of the resonant circuit; the graphical effect is a slight rotation of the phase angle curve about the zero axis. In practice, R_1 is adjusted at the minimum capacitance end of the dial ($+15^\circ$).
3. L_2 cancels the effects of C_{1A} , making the shape of the phase-angle curve symmetrical on both sides of the zero axis. L_2 is adjusted at the maximum capacitance end of the dial (-15°).

Adjustment of the circuit in the order described will result in minimum adjustment interaction.

The overall Q of the circuit is about 2.5, instead of 25 to 50, so the circuit stability is very good; the low Q also reduces the amplitude variation to less than 5%. The Q is determined largely by the series combination $R_{in}/R_{generator}$ in parallel with R_{out}/R_{load} .

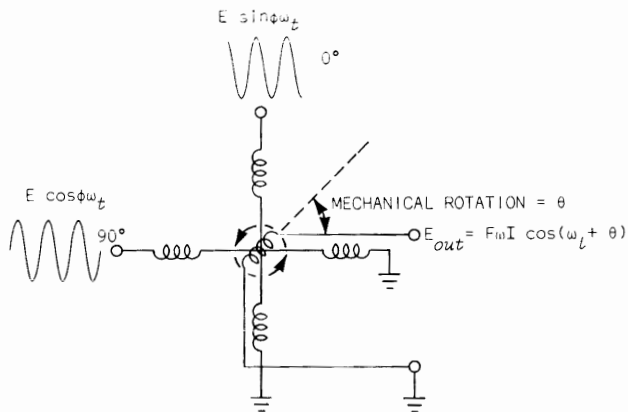


Fig. 15-24. Simplified inductive goniometer.

inductive
goniometer

The fourth type of phase shift network is the inductive goniometer. The term "gonia" is a Greek word meaning "phase." Goniometers have been used through history to mechanically measure phase angles, but they have been used very little to electronically measure phase angles until recent innovations simplified the circuit techniques.

goniometer
advantage

The principal advantage of the inductive phase-shift network is the availability of a 360° continuous phase shift which, for practical purposes, can be made to vary linearly with the angular rotation of a mechanical shaft.

The inductive phase-shift network operates similar to any inductively coupled transformer. Recall that in a conventional transformer, a current will be induced into the secondary winding reversed in phase. However, if two pairs of coils are mechanically and electrically arranged in quadrature (90° apart) as shown in Fig. 15-24, a *rotating* magnetic field will be formed.

A pickup coil (secondary winding), which can be mechanically rotated, is placed in the center of the rotating magnetic field. The phase angle of the pickup-coil output voltage will be determined by the angular position of the pickup coil with respect to the two quadrature coils.

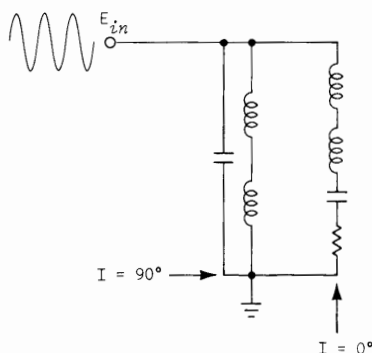


Fig. 15-25. 90° phase-shifting network of the inductive goniometer.

The circuit illustrated in Fig. 15-25 shows how the 90° phase shifting of the two coil-pairs is accomplished. The two sets of coils are resonant circuits. One pair of coils is series-resonant; the current is in phase with the driving voltage. The second pair of coils form a parallel-resonant circuit with the current shifted 90° from the driving voltage.

The actual coupling between the quadrature coils and the pickup coil is very low, but high coupling has more disadvantages than the addition of an external amplifier.

To eliminate the use of contact wipers on the rotating coil, a tightly coupled transformer is used to transfer the phase-shifted voltage to external circuitry. The pickup coil and the coupling transformer form a series-resonant circuit as shown in the complete circuit of Fig. 15-26. The series resistor R_1 limits the Q to a relatively low value to make the output amplitude more uniform as the coil is rotated, as well as to eliminate the need for optimum-value or adjustable components.

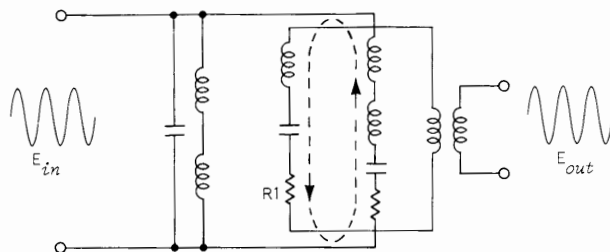


Fig. 15-26. Actual circuit of inductive goniometer.

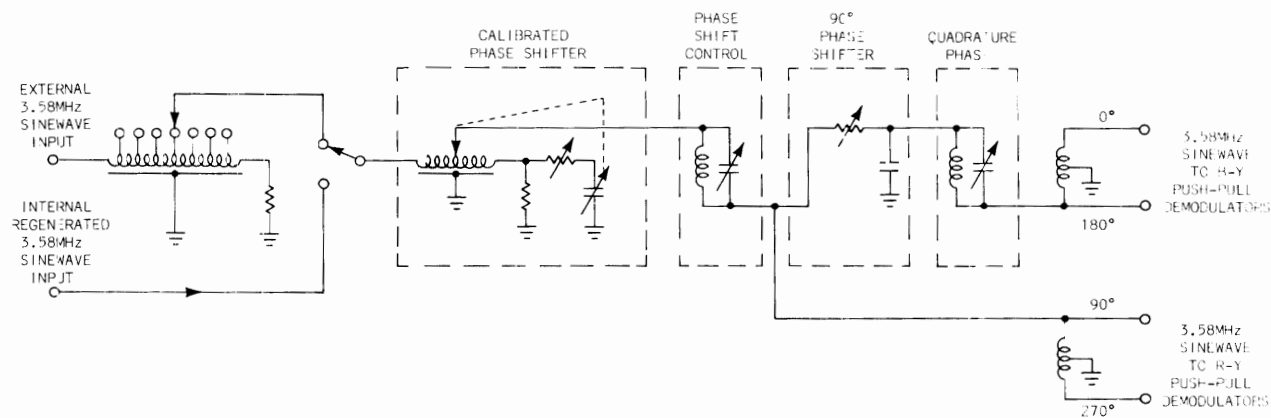


Fig. 15-27. Function block of subcarrier phase-processing system.

Fig. 15-27 shows how the phase-shifting networks are arranged in the subcarrier processing system to process the 3.58-MHz input sinewave into four phase-related sinewaves at the output. (See also Fig. 15-2.)

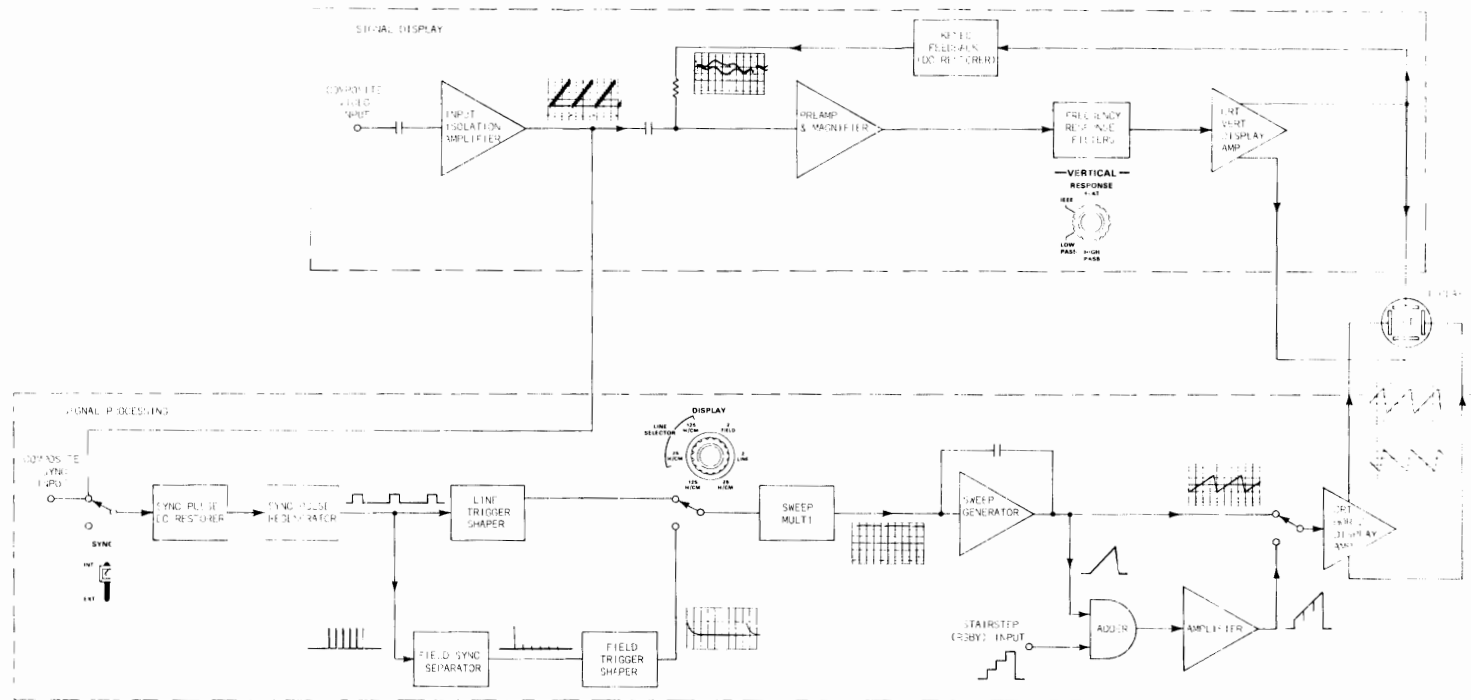


Fig. 16-1. Block diagram of typical composite video waveform oscilloscope.

16

CIRCUIT CONCEPTS IN THE FUNCTIONAL BLOCK SYSTEM

At this point the functional block diagram of the waveform monitor can now be discussed in more complete detail in terms of the circuit concepts. Waveform monitor oscilloscopes, that is, oscilloscopes designed specifically to process and display composite video waveforms, fall into two general categories:

1. Conventional waveform oscilloscopes.
2. Composite chroma vector display oscilloscopes or, simply, vectorscopes.

The primary difference between the two types of oscilloscopes is the measurement reference. The conventional waveform oscilloscope displays the composite video against a time reference on the horizontal axis. The vectorscope is an X-Y oscilloscope which displays the relationship between the two decoded color-difference signals. The measurement reference is not a horizontal time base but instead a selected segment of the waveform itself (burst).

Fig. 16-1 illustrates the block diagram of a typical waveform oscilloscope. The general functional blocks fall into two classes corresponding to the vertical and horizontal axes of the display CRT.

These two general blocks are:

1. Signal-Display Block. The principal objective of the signal-display block is to process and amplify the composite video waveform with minimum distortion of the waveform.
2. Signal-Processing Block. The principal objective of the signal-processing block is not necessarily minimum waveform distortion but minimum interaction between the separated components of the composite video waveform.

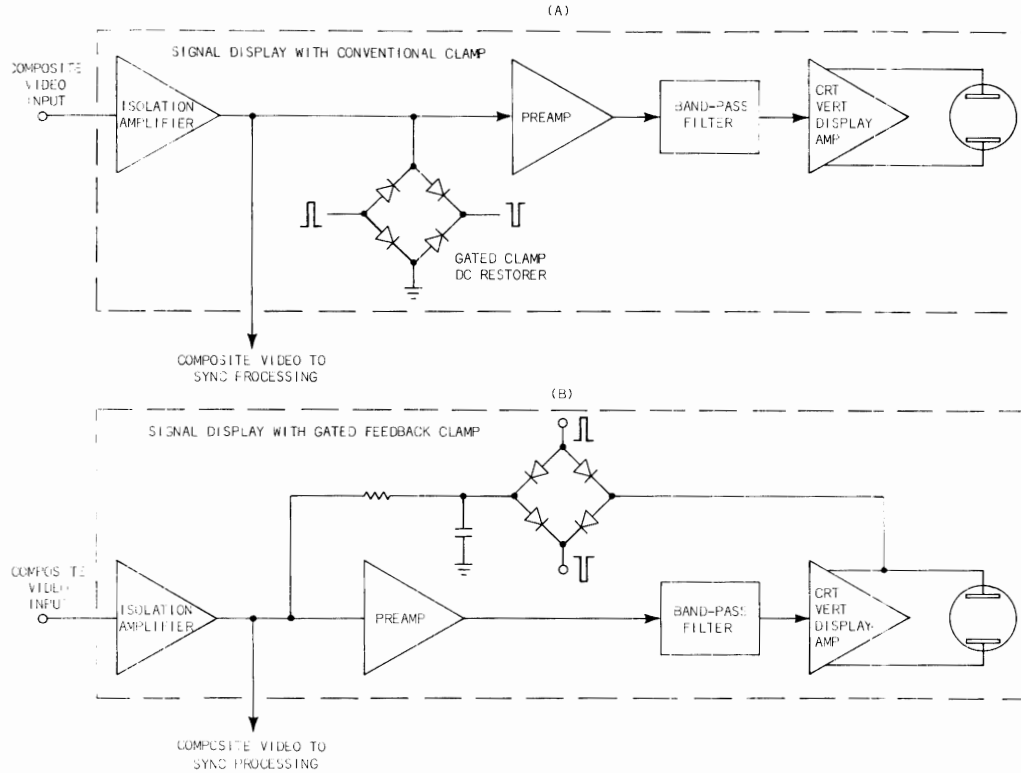


Fig. 16-2. Block diagram of two alternative approaches to signal displaying.

The signal-display block is very similar to the vertical amplifier of any conventional laboratory-grade oscilloscope with two additional functions:

1. A DC-restorer circuit to reinsert the composite-video DC reference.
2. Selectable frequency-response filters to facilitate specific measurements of the composite-video waveform. The filter designs are conventional and for that reason have not been covered in detail.

Observing the signal-display block more closely in Fig. 16-2, two *commonly used* forms of DC restoration can be seen.

Fig. 16-2A illustrates the gated-clamp-type DC restorer. Any flat recurrent portion of the composite-video waveform can selectively be clamped to ground or a fixed DC voltage. The two alternative portions of the composite-video waveform usually selected for clamping are the tip of the sync pulse or the back porch of the sync pulse. For some measurement applications, the gated clamp circuit has the advantage of removing 60-Hz and impulse noise in the sync region.

The more commonly used gated-feedback DC restorer is illustrated in Fig. 16-2B. The feedback DC restorer, like the gated-clamp restorer, produces minimum sync-tip distortion compared to the simple sync-tip clamp (peak-detector type) DC restorer.

The feedback restorer has several additional advantages compared to the gated clamp circuit:

1. The gated feedback circuit can be arranged to either: (A) Completely remove 60-Hz abnormalities which might exist on the composite video, or; (B) leave the 60 Hz (so its presence can be observed or measured) and only compensate for average picture level changes.
2. Amplifier stability. The keyed-feedback network, when applied around an amplifier system, not only establishes a selected portion of the composite-video waveform to a reference point, but also eliminates the effects of amplifier drift.

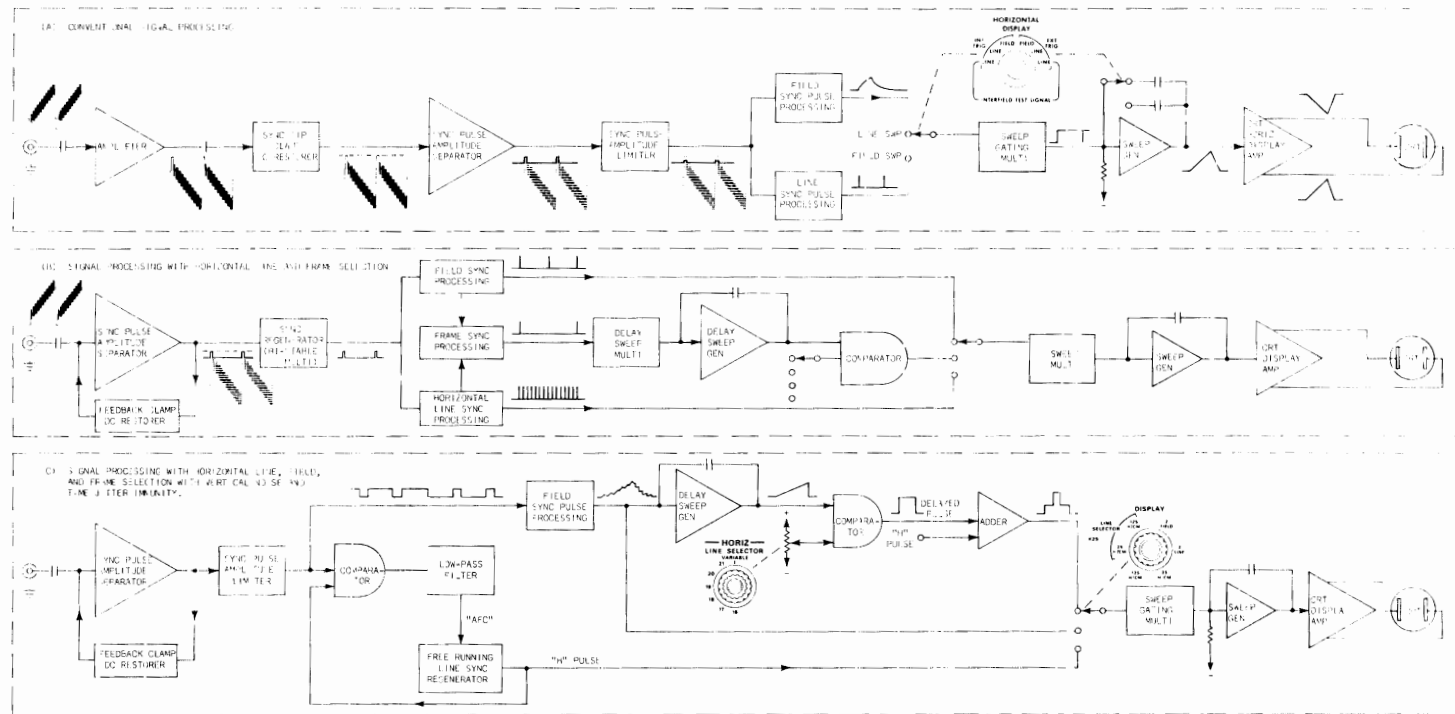


Fig. 16-3. Waveform monitor oscilloscope sync pulse separation and processing system variations.

In the second functional block -- the signal-processing block -- more circuit variety exists because the specific measurement requirements of each waveform monitor vary. The greatest circuit variation occurs in the sync pulse separation and processing circuitry, as shown in Fig. 16-3. If the waveform monitor is only intended to display one horizontal line or one field of composite video on the CRT, the sync pulse itself can be used to trigger a time base sweep as shown in Fig. 16-3A.

If the waveform monitor must display a small portion of one horizontal video line (magnified sweep), the time jitter of the CRT display must be minimized. A sync regenerator, free of noise and formed by the original sync pulse, is then used, similar to the block shown in Fig. 16-3B.

If the waveform monitor must display a selected line of one field (only one of the 525 lines), additional sync processing similar to the system illustrated in Fig. 16-3C, is used. A delaying sweep generator, triggered at a field or frame rate, is then added to prevent the display sweep generator from being triggered by each line sync pulse. The delaying sweep generator when used in conjunction with the normal display sweep generator is called the Line Selector Mode.

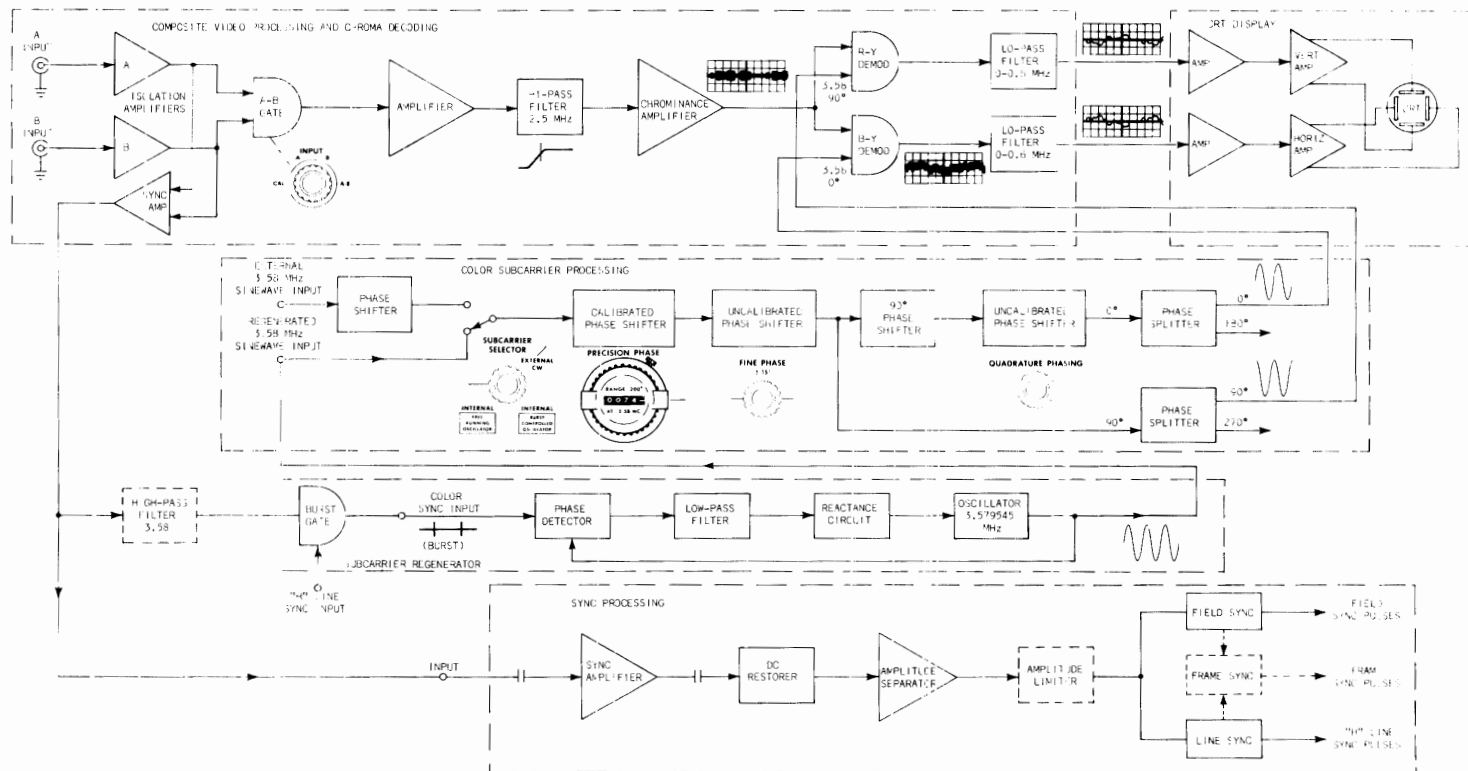


Fig. 16-4. Vectorscope functional block.

In addition to the two functional blocks of the conventional oscilloscope, the vectorscope has three additional functional blocks as illustrated in the diagram of Fig. 16-4. The additional three blocks of the vectorscope are needed to separate and decode the two color-difference signals from the composite video.

The five functional blocks of the vectorscope are:

1. X-Y CRT display (signal display)
2. Color subcarrier regenerator
3. Subcarrier processing
4. Chroma decoding
5. Sync processing.

The signal-display block consists of two identical CRT-display amplifier systems with special attention given to maintaining identical gain and frequency response of both systems. The CRT is arranged to have the same scanning area on both the vertical and horizontal axes.

Before the two color-difference signals can be amplified and displayed on the CRT, the composite color video waveform must first be processed by signal-decoding block.

The chroma-decoding block has three essential functions:

1. Removal of the chrominance information from the rest of the composite-video waveform. The chrominance separation process is usually carried out with a bandpass filter or a bandpass amplifier which rejects the lower-frequency video and sync information.
2. Chrominance demodulation -- the chrominance information is applied to two separate demodulators. The two regenerated subcarriers, identical in frequency but different in phase, are applied to the respective demodulators along with the composite chrominance signal.
3. The demodulated color-difference signals are then each passed through a low-pass filter to remove any remaining subcarrier component and second harmonics generated by the demodulator.

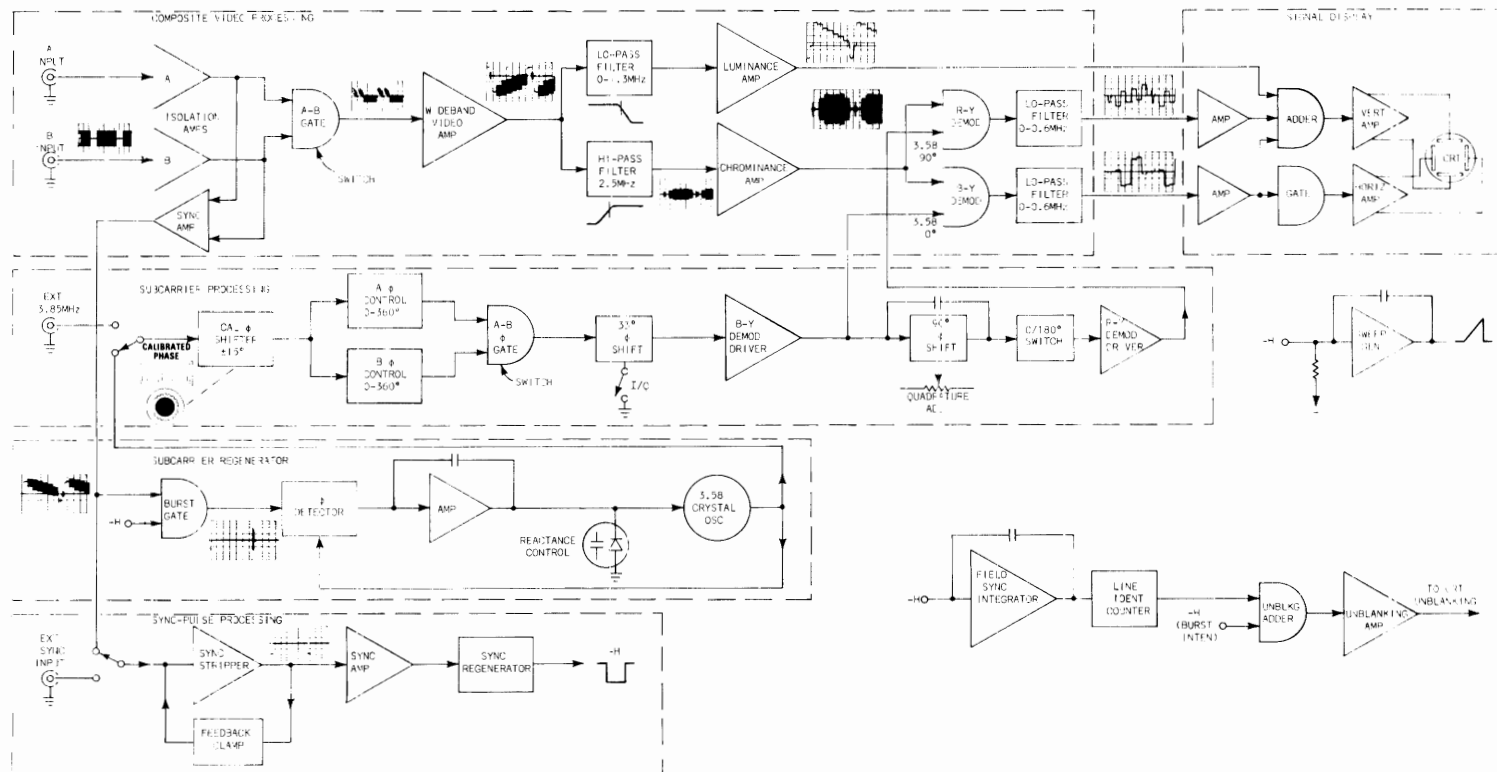


Fig. 16-5. Complete signal-decoding vectorscope.

The subcarrier-regenerator block functions as a controlled oscillator which must regenerate a frequency-stable constant-amplitude 3.579545-MHz sinewave for eventual carrier reinsertion at the demodulators.

A comparator in the form of a phase detector produces a correction voltage (if necessary) proportional to the phase difference between the color-burst sample of the original subcarrier and the regenerated subcarrier.

The correction voltage is used to control the frequency of the oscillator, ensuring an identical and constant phase between the original and regenerated subcarriers.

The subcarrier regenerator output sinewave is applied to the subcarrier processing circuitry. The subcarrier processing block has two principal functions:

1. Derive two phase-displaced subcarriers precisely 90° apart.
2. Provide calibrated phase adjustments in the form of front panel controls for vector-measurement purposes.

The subcarrier-processing circuitry is essentially a series of phase-splitting and phase-shifting networks. Additional circuit complexity is determined by the required phase-angle accuracy and long-term stability.

The final functional block is the sync processing circuit. Since the timebase reference in the vectorscope is used primarily for waveform display rather than as a reference, the requirements of the sync block are simpler than the conventional waveform monitor. The sync system is similar to the diagram of Fig. 16-3A. The composite-video sync pulses are used to trigger a delayed-pulse generator which in turn is used to selectively remove the color-sync burst from the composite video.

The block diagram of Fig. 16-5 illustrates the combined elements of a conventional waveform monitor and the chroma decoding vectorscope. The result displayed on the CRT is the combined characteristics of colored light -- not only the hue and saturation

but the brightness as well. By combining the luminance waveform with the chrominance waveform, the various colors are then displayed in terms of the three primary colors optically derived at the camera input and displayed at the picture tube output.

The entire coding and decoding system from the camera to the picture monitor can then be evaluated. As a result, system errors can be measured and analyzed to insure correct color rendition.

Since the eye is sensitive to logarithmic light changes, the colored luminance errors must be relatively large before they can be visually detected. In addition, the variation in color response between observers makes subjective evaluation of picture monitor displays somewhat inconsistent. The additional decoding process in the vectorscope illustrated in the block diagram of Fig. 16-5 eliminates the need to subjectively evaluate the decoded color signal on a picture monitor.

REFERENCES

- Carnt and Townsend. *Color Television*.
 Cutler. *Electronic Circuit Analysis*.
 Fink. *Television Engineering*.
 Grob. *Basic Television*.
 Grob. *Basic Television Principles and Service*.
 Hazeltine. *Principles of Color Television*.
Electronics. Volume 23.
NAB Engineering Handbook.
MIT Radiation Laboratory Series. Volumes 19, 21, 22.
Proceedings of the IRE. Volumes 29, 31, 37, 39, 40, 41, 42.
RCA Review. Volumes 9, XV.

INDEX

- Average picture level, 19-20
- Back porch, 11-13
- Black level, 11-13
- Brightness, 1, 99
- Burst, 11, 145-150, 175
- Carrier-balanced modulator, 110-111
- Carrier-operated switch demodulator, 114-121
- Carrier suppression, 108
- Chroma, 1, 15
- Chromaticity, 101-108
- Chrominance signal, 104-108
- Color burst, 11, 145-150
- Color demodulation, 109-121
- Colorimetry, 97-108
- Color sync, 11-13
- Composite sync, 43
- Composite video, 9, 15
- Contrast, 1
- Crystal oscillator, 133-136, 143
- DC restorer,
 - feedback, 35-42, 63
 - gated clamp, 176-177
 - keyed clamp, 28-42, 63
 - positive bias, 27-42, 65
 - simple, 24-25, 47
- Delay-line phase-shift network, 159-165
- Demodulators, 113-121
- Detector,
 - phase, 127-130, 141
 - synchronous, 113
- Doubly-balanced modulator, 110-111
- Dynamic phase error, 131, 137
- Equalizing pulses, 13, 66
- Field, 5, 89-91
- Field pulse differentiator, 82-87
- Field pulse discrimination, 82-86
- Field identification, 43, 89-96
- Field pulse integrator, 79-81
- Field sync, 11-13, 43, 47-66, 79-92
- Flicker, 4
- Frame, 5
- Frame rate, 4
- Frame sync, 43
- Frequency interlace, 125
- Frequency multiplexing, 125
- Frequency pull-in time, 130-131
- Front porch, 11
- Goniometer, 170-173
- H pulses, 12, 66
- Hue, 11-13, 99, 151
- Inductive goniometer, 170-173
- Interlaced scanning, 5
- Jitter,
 - phase, 131
 - time, 51-53, 61-62, 75, 80, 179
- Line sync, 11, 43-66, 77
- Luminance, 99, 104-108
- Phase detector, 127-130, 141
- Phase error,
 - dynamic, 131, 137
 - static, 130-131, 137
- Phase jitter, 131
- Phase-shifting networks
 - delay line, 159-165
 - inductive goniometer, 170-173
 - RC, 153-154
 - resonant, 154-159, 165-169
- Phase-transfer ratio, 132
- Resolution, 5-8
- Resonant phase-shift network, 154-159, 165-169, 170-173
- Saturation, 11-13, 99
- Set-up (black) level, 11
- Static phase error, 130-131, 137
- Subcarrier regenerators, 123-144, 183-184

- Suppressed carrier, 108
- Sync,
 - color, 11-13
 - composite, 43-44, 82
 - field, 11-13, 43, 47-66, 77-92
 - frame, 43, 89
 - line, 11, 43-66, 77-79
- Synchronous detector, 113
- Sync pulse regenerator, 55,
65-75
- Sync separator, 47-75
 - requirements, 65
- Sync-tip clamping
 - (DC restoration), 23
- Timing jitter, 51-53, 61-62, 75,
80, 179
- Video-balanced modulator,
110-111
- Vertical interval testing (VIT),
89
- Visual perception, 1

CIRCUIT CONCEPTS BY INSTRUMENT

Many concepts discussed in this book are applicable to all Tektronix products designed for television waveform processing. Therefore, the concepts listed are used in only certain products. Applicability is indicated by a ●.

CIRCUIT CONCEPTS:	INSTRUMENTS:						
	124	520	524	525	526	527	529
BURST GATING (p148)					●		
DC RESTORATION with:							
feedback clamp (p36)						●	
keyed clamp (p28)				●			
modified feedback clamp (p41)							●
positive bias (p27)						●	
transistor keyed clamp (p32)		●					
DEMODULATION with:							
synchronous demodulation (p113)					●		
synchronous switch (p115)		●					
FIELD IDENTIFICATION by:							
time coincidence and delayed pulse generator (p91)							●
time coincident sync pulses (p94)		●					
FIELD SYNC PROCESSING,							
by integration (p79)	●		●		●		
with differentiating network(p84)							●
with high-pass differentiating amplifier (p86)							●
with transmission-line reflection (p80)						●	
OSCILLATOR FREQUENCY CONTROL (pp133,138)		●			●		
PHASE DETECTION (p127)					●		
PHASE JITTER REDUCTION (p131)					●		
SUBCARRIER PROCESSING (p153)					●		
SUBCARRIER REGENERATION (pp134,143)		●			●		

015-0062-00

CIRCUIT CONCEPTS:	INSTRUMENTS:						
	124	520	524	525	526	527	529
SUBCARRIER PHASE SHIFTING,							
calibrated (pp161,166)		●			●		
inductive goniometer (p170)		●					
parallel resonant (p157)					●		
RC (pp153,156)		●			●		
switchable parallel resonant(p157)					●		
transmission line (p159)					●		
variable delay line (p161)					●		
SYNC AMPLITUDE SEPARATION (pp49,51)	●		●		●		
by nonlinear feedback amplitude stripper (p56)							●
by nonlinear feedback amplifier with limiter amplifier (p61)							●
by nonlinear feedback amplifier with limiter amplifier(p62)		●					
with limiter (p53)				●			
SYNC REGENERATION (pp55,59)						●	●
with free-running system (p66)		●					

015-0062-00

NOTES

NOTES

*This book is but one of a series.
The series consists of two groups,
circuits and measurements. These
texts present a conceptual approach
to circuits and measurements which
apply to Tektronix products.
Several "concept books" are in
preparation; those now available
are:*

Power Supply Circuits

Oscilloscope Cathode-Ray Tubes

Storage Cathode-Ray Tubes and Circuits

Television Waveform Processing Circuits