

SUNSTAR 商斯达实业集团是集研发、生产、工程、销售、代理经销、技术咨询、信息服务等为一体的高科技企业，是专业高科技电子产品生产厂家，是具有10多年历史的专业电子元器件供应商，是中国最早和最大的仓储式连锁规模经营大型综合电子零部件代理分销商之一，是一家专业代理和分销世界各大品牌IC芯片和电子元器件的连锁经营综合性国际公司，专业经营进口、国产名厂名牌电子元件，型号、种类齐全。在香港、北京、深圳、上海、西安、成都等全国主要电子市场设有直属分公司和产品展示展销窗口门市部专卖店及代理分销商，已在全国范围内建成强大统一的供货和代理分销网络。我们专业代理经销、开发生产电子元器件、集成电路、传感器、微波光电元器件、工控机/DOC/DOM电子盘、专用电路、单片机开发、MCU/DSP/ARM/FPGA软件硬件、二极管、三极管、模块等，是您可靠的一站式现货配套供应商、方案提供商、部件功能模块开发配套商。商斯达实业公司拥有庞大的资料库，有数位毕业于著名高校——有中国电子工业摇篮之称的西安电子科技大学（西军电）并长期从事国防尖端科技研究的高级工程师为您精挑细选、量身订做各种高科技电子元器件，并解决各种技术问题。

微波光电部专业研制、代理经销高频、微波、光纤、光电元器件、组件、部件、模块、整机；电磁兼容元器件、材料、设备；微波CAD、EDA软件、开发测试仿真工具；微波、光纤仪器仪表。欢迎国外高科技微波、光纤厂商将优秀产品介绍到中国、共同开拓市场。长期大量现货专业批发高频、微波、卫星、光纤、电视、CATV器件：晶振、VCO、连接器、PIN开关、变容二极管、开关二极管、低噪晶体管、功率电阻及电容、放大器、功率管、MMIC、混频器、耦合器、功分器、振荡器、合成器、衰减器、滤波器、隔离器、环行器、移相器、调制解调器；光电子元件和组件：红外发射管、红外接收管、光电开关、光敏管、发光二极管和发光二极管组件、半导体激光二极管和激光器组件、光电探测器和光接收组件、光发射接收模块、光纤激光器和光放大器、光调制器、光开关、DWDM用光发射和接收器件、用户接入系统光收发器件与模块、光纤连接器、光纤跳线/尾纤、光衰减器、光纤适配器、光隔离器、光耦合器、光环行器、光复用器/转换器；无线收发芯片和模组、蓝牙芯片和模组。

更多产品请看本公司产品专用销售网站：[欢迎索取免费详细资料、设计指南和光盘](#)；产品凡多，未能尽录，欢迎来电查询

商斯达中国传感器科技信息网：<http://www.sensor-ic.com/>

商斯达工控安防网：<http://www.pc-ps.net/>

商斯达电子元器件网：<http://www.sunstare.com/>

商斯达微波光电产品网：[HTTP://www.rfoe.net/](http://www.rfoe.net/)

商斯达消费电子产品网：<http://www.icasic.com/>

商斯达实业科技产品网：<http://www.sunstars.cn/> 微波元器件销售热线：

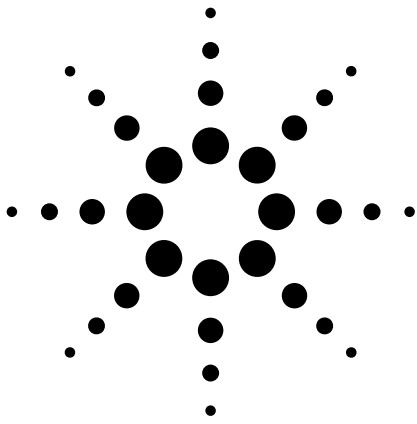
地址：深圳市福田区福华路福庆街鸿图大厦1602室

电话：0755-82884100 83397033 83396822 83398585

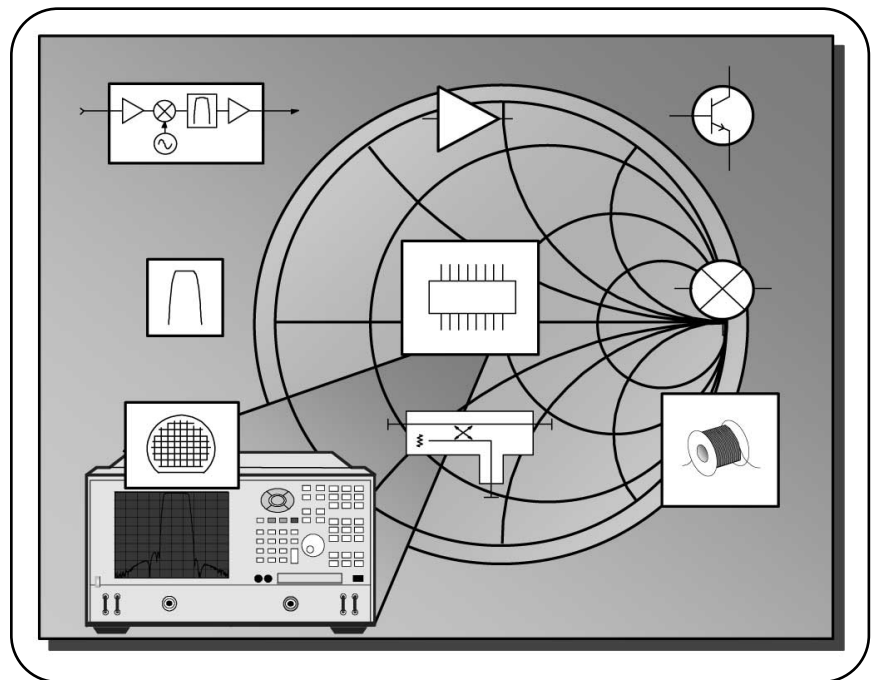
传真：0755-83376182 (0) 13823648918 MSN: SUNS8888@hotmail.com

邮编：518033 E-mail:szss20@163.com QQ: 195847376

技术支持: 0755-83394033 13501568376



Agilent Network Analyzer Basics

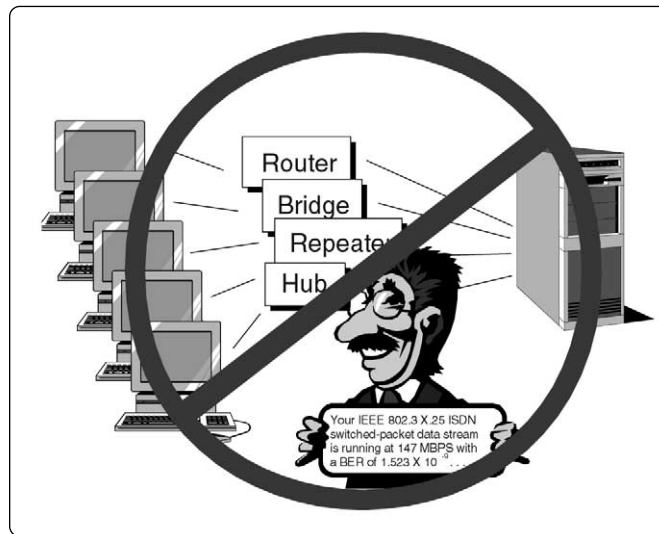


Agilent Technologies

Abstract

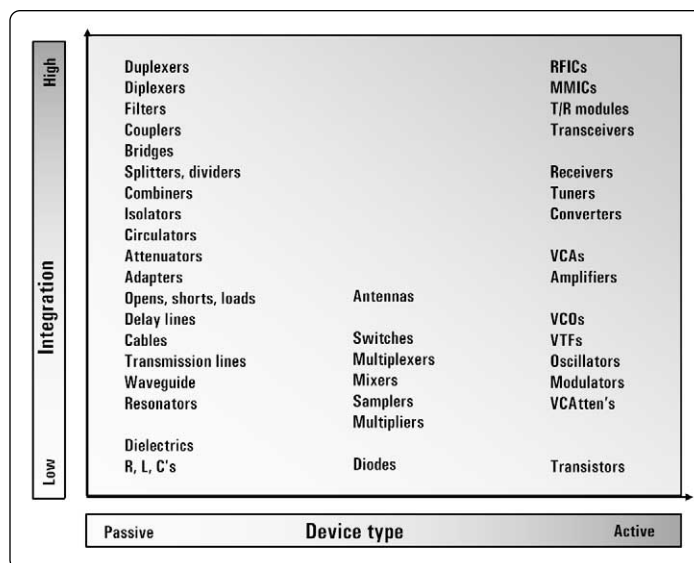
This presentation covers the principles of measuring high-frequency electrical networks with network analyzers.

You will learn what kinds of measurements are made with network analyzers, and how they allow you to characterize both linear and nonlinear behavior of your devices. The session starts with RF fundamentals such as transmission lines and the Smith chart, leading to the concepts of reflection, transmission and S-parameters. The next section covers the major components in a network analyzer, including the advantages and limitations of different hardware approaches. Error modeling, accuracy enhancement, and various calibration techniques will then be presented. Finally, some typical swept-frequency and swept-power measurements commonly performed on filters and amplifiers will be covered. An appendix is also included with information on advanced topics, with pointers to more information.



Network Analysis is Not....

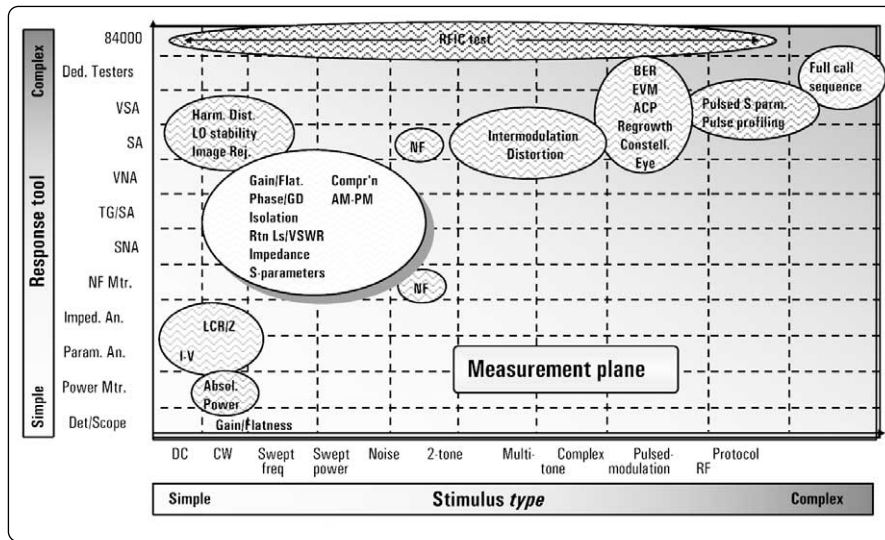
This module is not about computer networks! When the name "network analyzer" was coined many years ago, there were no such things as computer networks. Back then, networks always referred to *electrical* networks. Today, when we refer to the things that network analyzers measure, we speak mostly about devices and components.



What Types of Devices Are Tested?

Here are some examples of the types of devices that you can test with network analyzers. They include both passive and active devices (and some that have attributes of both). Many of these devices need to be characterized for both linear and nonlinear behavior. It is not possible to completely characterize all of these devices with just one piece of test equipment.

The next slide shows a model covering the wide range of measurements necessary for complete linear and nonlinear characterization of devices. This model requires a variety of stimulus and response tools. It takes a large range of test equipment to accomplish all of the measurements shown on this chart. Some instruments are optimized for one test only (like bit-error rate), while others, like network analyzers, are much more generalpurpose in nature. Network analyzers can measure both linear and nonlinear behavior of devices, although the measurement techniques are different (frequency versus power sweeps for example). This module focuses on swept-frequency and swept-power measurements made with network analyzers



Device Test Measurement Model

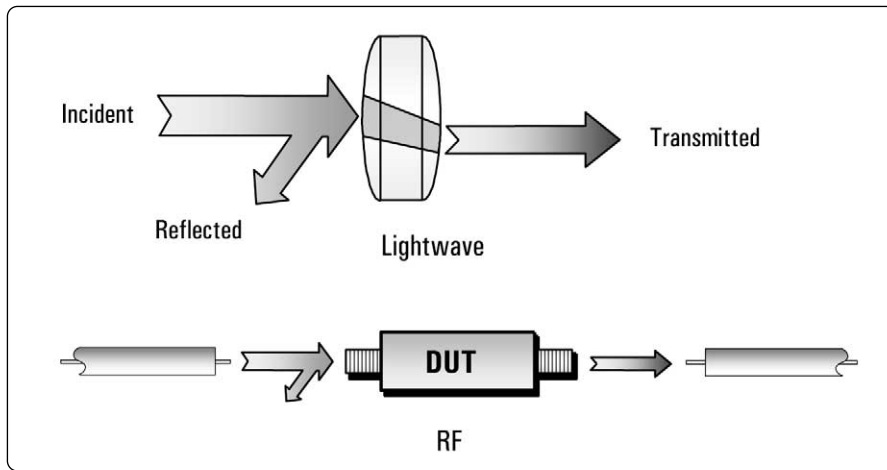
Here is a key to many of the abbreviations used at right:

Response

84000	8400 series high-volume RFIC tester
Ded. Testers	Dedicated (usually one-box) testers
VSA	Vector signal analyzer
SA	Spectrum analyzer
VNA	Vector signal analyzer
TG/SA	Tracking generator/spectrum analyzer
SNA	Scalar network analyzer
NF Mtr.	Noise-figure meter
Imped. An.	Impedance analyzer (LCR meter)
Power Mtr.	Power meter
Det./Scope	Diode detector/oscilloscope

Measurement

ACP	Adjacent channel power
AM-PM	AM to PM conversion
BER	Bit-error rate
Compr'n	Gain compression
Constell.	Constellation diagram
EVM	Error-vector magnitude
Eye	Eye diagram
GD	Group delay
Harm. Dist.	Harmonic distortion
NF	Noise figure
Regrowth	Spectral regrowth
Rtn Ls	Return loss
VSWR	Voltage standing wave ratio



Lightwave Analogy to RF Energy

One of the most fundamental concepts of high-frequency network analysis involves incident, reflected and transmitted waves traveling along transmission lines. It is helpful to think of traveling waves along a transmission line in terms of a lightwave analogy. We can imagine incident light striking some optical component like a clear lens. Some of the light is reflected off the surface of the lens, but most of the light continues on through the lens. If the lens were made of some lossy material, then a portion of the light could be absorbed within the lens. If the lens had mirrored surfaces, then most of the light would be reflected and little or none would be transmitted through the lens. This concept is valid for RF signals as well, except the electromagnetic energy is in the RF range instead of the optical range, and our components and circuits are electrical devices and networks instead of lenses and mirrors.

Network analysis is concerned with the accurate measurement of the ratios of the reflected signal to the incident signal, and the transmitted signal to the incident signal.

- Verify specifications of “building blocks” for more complex RF systems
- Ensure distortionless transmission of communications signals
 - linear: constant amplitude, linear phase / constant group delay
 - nonlinear: harmonics, intermodulation, compression, AM-to-PM conversion
- Ensure good match when absorbing power (e.g., an antenna)



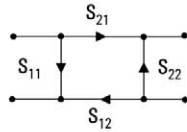
Why Do We Need to Test Components?

Components are tested for a variety of reasons. Many components are used as “building blocks” in more complicated RF systems. For example, in most transceivers there are amplifiers to boost LO power to mixers, and filters to remove signal harmonics. Often, R&D engineers need to measure these components to verify their simulation models and their actual hardware prototypes. For component production, a manufacturer must measure the performance of their products so they can provide accurate specifications. This is essential so prospective customers will know how a particular component will behave in their application.

When used in communications systems to pass signals, designers want to ensure the component or circuit is not causing excessive signal distortion. This can be in the form of linear distortion where flat magnitude and linear phase shift versus frequency is not maintained over the bandwidth of interest, or in the form of nonlinear effects like intermodulation distortion.

Often it is most important to measure how reflective a component is, to ensure that it absorbs energy efficiently. Measuring antenna match is a good example.

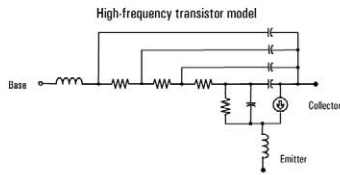
1. Complete characterization of linear networks



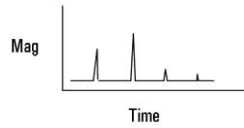
2. Complex impedance needed to design matching circuits



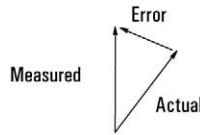
3. Complex values needed for device modeling



4. Time-domain characterization



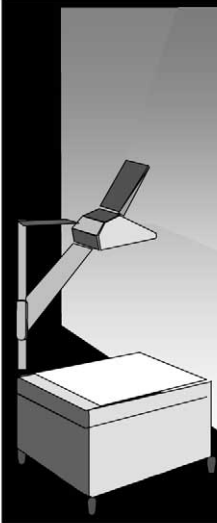
5. Vector-error correction



The Need for Both Magnitude and Phase

In many situations, magnitude-only data is sufficient for our needs. For example, we may only care about the gain of an amplifier or the stop-band rejection of a filter. However, as we will explore throughout this paper, measuring phase is a critical element of network analysis.

Complete characterization of devices and networks involves measurement of phase as well as magnitude. This is necessary for developing circuit models for simulation and to design matching circuits based on conjugate matching techniques. Time-domain characterization requires magnitude and phase information to perform the inverse-Fourier transform. Finally, for best measurement accuracy, phase data is required to perform vector error correction.



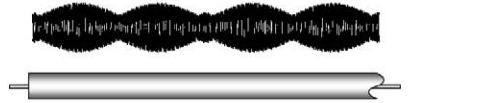
- **What measurements do we make?**
 - Transmission-line basics
 - Reflection and transmission parameters
 - S-parameter definition
- **Network analyzer hardware**
 - Signal separation devices
 - Detection types
 - Dynamic range
 - T/R versus S-parameter test sets
- **Error models and calibration**
 - Types of measurement error
 - One- and two-port models
 - Error-correction choices
 - Basic uncertainty calculations
- **Example measurements**
- **Appendix**

Agenda

In this section we will review reflection and transmission measurements. We will see that transmission lines are needed to convey RF and microwave energy from one point to another with minimal loss, that transmission lines have a characteristic impedance, and that a termination at the end of a transmission line must match the characteristic impedance of the line to prevent loss of energy due to reflections. We will see how the Smith chart simplifies the process of converting reflection data to the complex impedance of the termination. For transmission measurements, we will discuss not only simple gain and loss but distortion introduced by linear devices. We will introduce S-parameters and explain why they are used instead of h-, y-, or z-parameters at RF and microwave frequencies.

Low frequencies

- wavelengths $\gg$ wire length
- current (I) travels down wires easily for efficient power transmission
- measured voltage and current not dependent on position along wire



High frequencies

- wavelength $\approx$ or $<$ length of transmission medium
- need transmission lines for efficient power transmission
- matching to characteristic impedance (Z_0) is very important for low reflection and maximum power transfer
- measured envelope voltage dependent on position along line

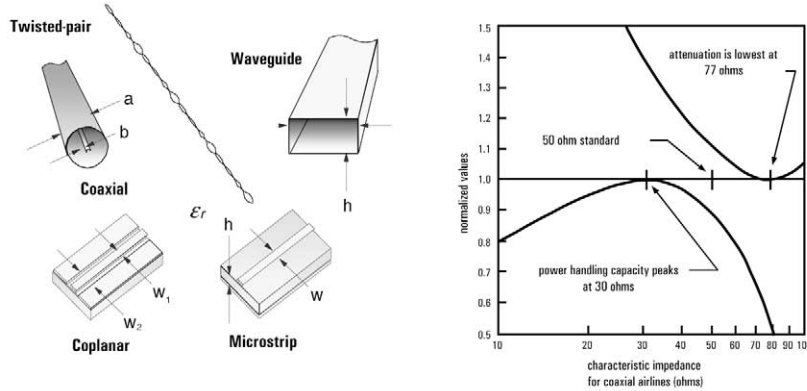
Transmission Line Basics

The need for efficient transfer of RF power is one of the main reasons behind the use of transmission lines. At low frequencies where the wavelength of the signals are much larger than the length of the circuit conductors, a simple wire is very useful for carrying power. Current travels down the wire easily, and voltage and current are the same no matter where we measure along the wire.

At high frequencies however, the wavelength of signals of interest are comparable to or much smaller than the length of conductors. In this case, power transmission can best be thought of in terms of traveling waves.

Of critical importance is that a lossless transmission line takes on a characteristic impedance (Z_0). In fact, an infinitely long transmission line appears to be a resistive load! When the transmission line is terminated in its characteristic impedance, maximum power is transferred to the load. When the termination is not Z_0 , the portion of the signal which is not absorbed by the load is reflected back toward the source. This creates a condition where the envelope voltage along the transmission line varies with position. We will examine the incident and reflected waves on transmission lines with different load conditions in following slides

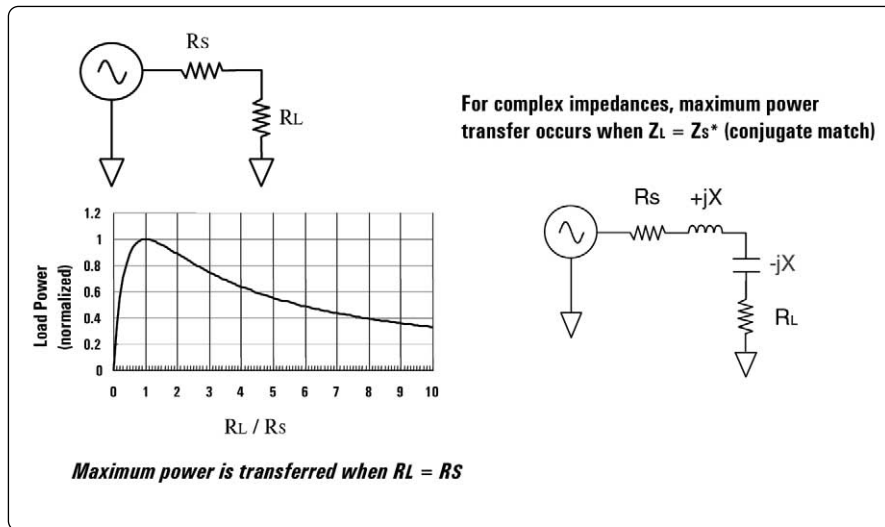
- Z_0 determines relationship between voltage and current waves
- Z_0 is a function of physical dimensions and ϵ_r
- Z_0 is usually a real impedance (e.g. 50 or 75 ohms)



Transmission Line Z_0

RF transmission lines can be made in a variety of transmission media. Common examples are coaxial, waveguide, twisted pair, coplanar, stripline and microstrip. RF circuit design on printed-circuit boards (PCB) often use coplanar or microstrip transmission lines. The fundamental parameter of a transmission line is its characteristic impedance Z_0 . Z_0 describes the relationship between the voltage and current traveling waves, and is a function of the various dimensions of the transmission line and the dielectric constant (ϵ_r) of the non-conducting material in the transmission line. For most RF systems, Z_0 is either 50 or 75 ohms.

For low-power situations (cable TV, for example) coaxial transmission lines are optimized for low loss, which works out to about 75 ohms (for coaxial transmission lines with air dielectric). For RF and microwave communication and radar applications, where high power is often encountered, coaxial transmission lines are designed to have a characteristic impedance of 50 ohms, a compromise between maximum power handling (occurring at 30 ohms) and minimum loss.



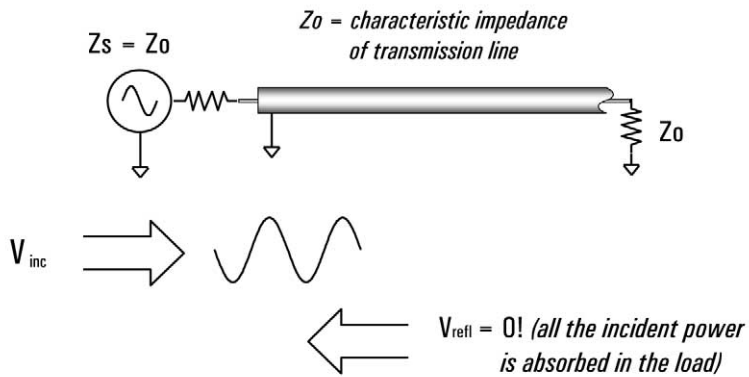
Power Transfer Efficiency

Before we begin our discussion about transmission lines, let us look at the condition for maximum power transfer into a load, given a source impedance of R_s . The graph above shows that the matched condition ($R_L = R_S$) results in the maximum power dissipated in the load resistor. This condition is true whether the stimulus is a DC voltage source or an RF sinusoid.

For maximum transfer of energy into a transmission line from a source or from a transmission line to a load (the next stage of an amplifier, an antenna, etc.), the impedance of the source and load should match the characteristic impedance of the transmission line. In general, then, Z_0 is the target for input and output impedances of devices and networks.

When the source impedance is not purely resistive, the maximum power transfer occurs when the load impedance is equal to the complex conjugate of the source impedance. This condition is met by reversing the sign of the imaginary part of the impedance. For example, if $R_S = 0.6 + j0.3$, then the complex conjugate $R_S^* = 0.6 - j0.3$.

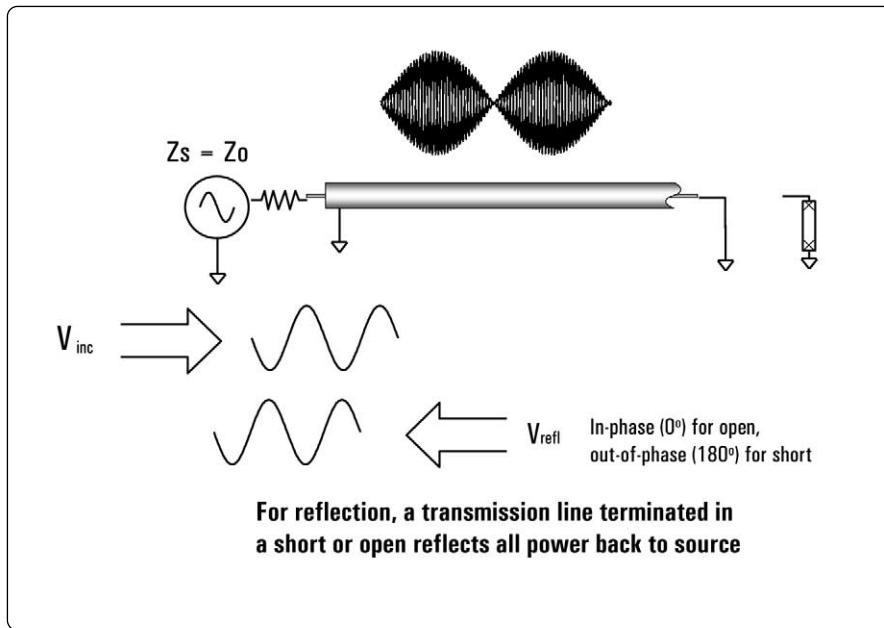
Sometimes the source impedance is adjusted to be the complex conjugate of the load impedance. For example, when matching to an antenna, the load impedance is determined by the characteristics of the antenna. A designer has to optimize the output match of the RF amplifier over the frequency range of the antenna so that maximum RF power is transmitted through the antenna.



For reflection, a transmission line terminated in Z_o behaves like an infinitely long transmission line

Transmission Line Terminated With Z_o

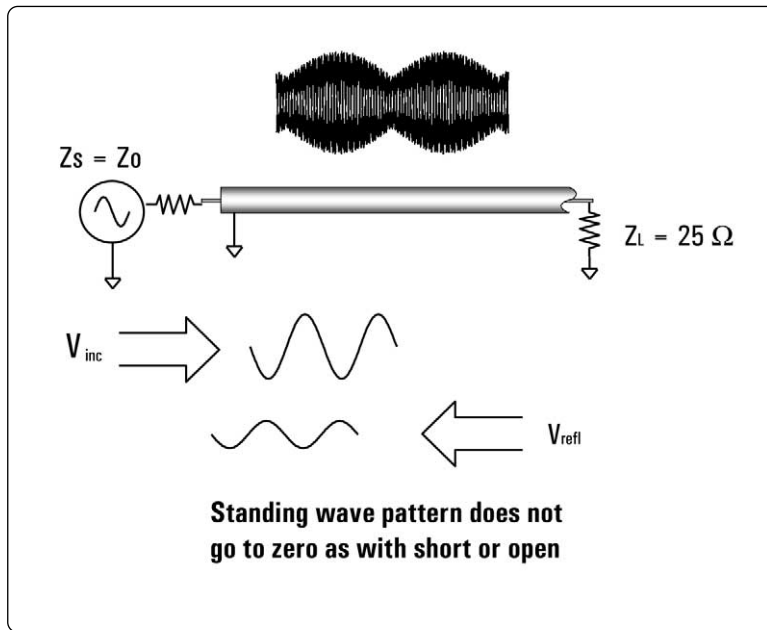
Let's review what happens when transmission lines are terminated in various impedances, starting with a Z_o load. Since a transmission line terminated in its characteristic impedance results in maximum transfer of power to the load, there is no reflected signal. This result is the same as if the transmission line was infinitely long. If we were to look at the envelope of the RF signal versus distance along the transmission line, it would be constant (no standing-wave pattern). This is because there is energy flowing in one direction only.



Transmission Line Terminated with Short, Open

Next, let's terminate our line in a short circuit. Since purely reactive elements cannot dissipate any power, and there is nowhere else for the energy to go, a reflected wave is launched back down the line toward the source. For Ohm's law to be satisfied (no voltage across the short), this reflected wave must be equal in voltage magnitude to the incident wave, and be 180° out of phase with it. This satisfies the condition that the total voltage must equal zero at the plane of the short circuit. Our reflected and incident voltage (and current) waves will be identical in magnitude but traveling in the opposite direction.

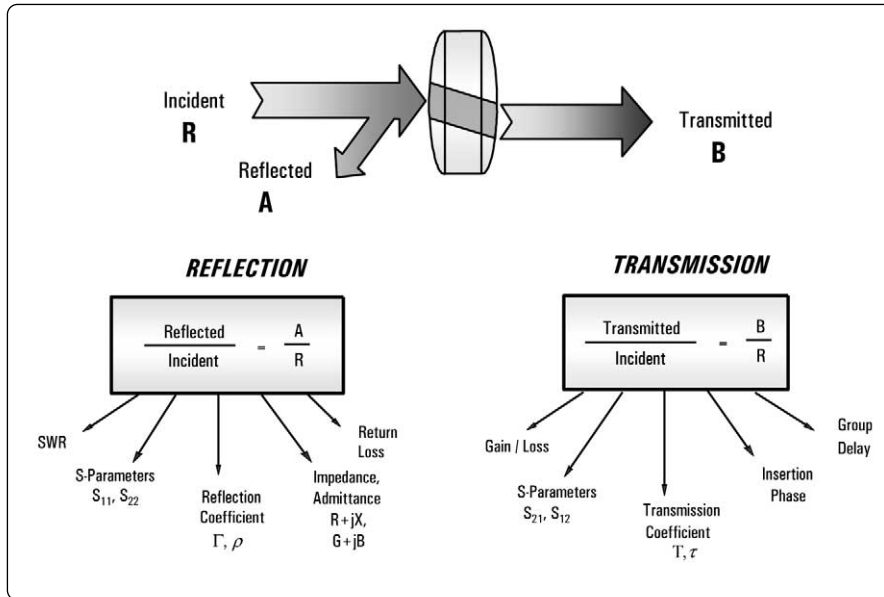
Now let us leave our line open. This time, Ohm's law tells us that the open can support no current. Therefore, our reflected current wave must be 180° out of phase with respect to the incident wave (the voltage wave will be in phase with the incident wave). This guarantees that current at the open will be zero. Again, our reflected and incident current (and voltage) waves will be identical in magnitude, but traveling in the opposite direction. For both the short and open cases, a standing-wave pattern will be set up on the transmission line. The valleys will be at zero and the peaks at twice the incident voltage level. The peaks and valleys of the short and open will be shifted in position along the line with respect to each other, in order to satisfy Ohm's law as described above.



Transmission Line Terminated with $25\ \Omega$

Finally, let's terminate our line with a $25\ \Omega$ resistor (an impedance between the full reflection of an open or short circuit and the perfect termination of a $50\ \Omega$ load). Some (but not all) of our incident energy will be absorbed in the load, and some will be reflected back towards the source. We will find that our reflected voltage wave will have an amplitude $1/3$ that of the incident wave, and that the two waves will be 180° out of phase at the load. The phase relationship between the incident and reflected waves will change as a function of distance along the transmission line from the load. The valleys of the standing-wave pattern will no longer be zero, and the peak will be less than that of the short/open case.

The significance of standing waves should not go unnoticed. Ohm's law tells us the complex relationship between the incident and reflected signals at the load. Assuming a 50-ohm source, the voltage across a 25-ohm load resistor will be two thirds of the voltage across a 50-ohm load. Hence, the voltage of the reflected signal is one third the voltage of the incident signal and is 180° out of phase with it. However, as we move away from the load toward the source, we find that the phase between the incident and reflected signals changes! The vector sum of the two signals therefore also changes along the line, producing the standing wave pattern. The apparent impedance also changes along the line because the relative amplitude and phase of the incident and reflected waves at any given point uniquely determine the measured impedance. For example, if we made a measurement one quarter wavelength away from the 25-ohm load, the results would indicate a 100-ohm load. The standing wave pattern repeats every half wavelength, as does the apparent impedance.



High-Frequency Device Characterization

Now that we fully understand the relationship of electromagnetic waves, we must also recognize the terms used to describe them. Common network analyzer terminology has the incident wave measured with the R (for reference) receiver. The reflected wave is measured with the A receiver and the transmitted wave is measured with the B receiver. With amplitude and phase information of these three waves, we can quantify the reflection and transmission characteristics of our device under test (DUT). Some of the common measured terms are scalar in nature (the phase part is ignored or not measured), while others are vector (both magnitude and phase are measured). For example, return loss is a scalar measurement of reflection, while impedance results from a vector reflection measurement. Some, like group delay, are purely phase-related measurements.

Ratioed reflection is often shown as A/R and ratioed transmission is often shown as B/R, relating to the measurement receivers used in the network analyzer

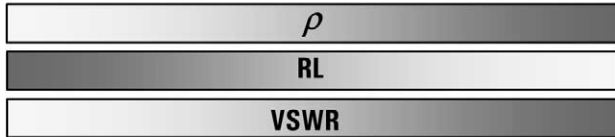
Reflection Coefficient

$$\Gamma = \frac{V_{\text{reflected}}}{V_{\text{incident}}} = \rho \angle \Phi = \frac{Z_L - Z_0}{Z_L + Z_0}$$

$$\text{Return loss} = -20 \log(\rho), \quad \rho = |\Gamma|$$

**Voltage Standing Wave Ratio**

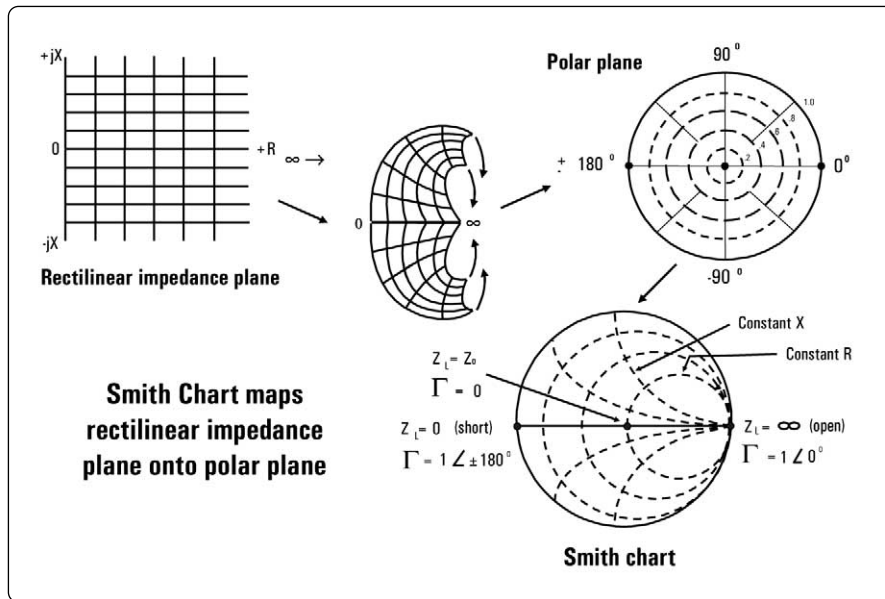
$$\text{VSWR} = \frac{E_{\text{max}}}{E_{\text{min}}} = \frac{1 + \rho}{1 - \rho}$$

No reflection
($Z_L = Z_0$)0
 ∞ dB
1**Full reflection**
($Z_L = \text{open, short}$)1
0 dB
 ∞ **Reflection Parameters**

Let's now examine reflection measurements. The first term for reflected waves is reflection coefficient gamma (Γ). Reflection coefficient is the ratio of the reflected signal voltage to the incident signal voltage. It can be calculated as shown above by knowing the impedances of the transmission line and the load. The magnitude portion of gamma is called rho (ρ). A transmission line terminated in Z_0 will have all energy transferred to the load; hence $V_{\text{refl}} = 0$ and $\rho = 0$. When Z_L is not equal to Z_0 , some energy is reflected and ρ is greater than zero. When Z_L is a short or open circuit, all energy is reflected and $\rho = 1$. The range of possible values for ρ is therefore zero to one.

Since it is often very convenient to show reflection on a logarithmic display, the second way to convey reflection is return loss. Return loss is expressed in terms of dB, and is a scalar quantity. The definition for return loss includes a negative sign so that the return loss value is always a positive number (when measuring reflection on a network analyzer with a log magnitude format, ignoring the minus sign gives the results in terms of return loss). Return loss can be thought of as the number of dB that the reflected signal is below the incident signal. Return loss varies between infinity for a Z_0 impedance and 0 dB for an open or short circuit.

As we have already seen, two waves traveling in opposite directions on the same transmission line cause a "standing wave". This condition can be measured in terms of the voltage-standing-wave ratio (VSWR or SWR for short). VSWR is defined as the maximum value of the RF envelope over the minimum value of the envelope. This value can be computed as $(1 + \rho)/(1 - \rho)$. VSWR can take ?????



Smith Chart Review

Our network analyzer gives us complex reflection coefficient. However, we often want to know the impedance of the DUT. The previous slide shows the relationship between reflection coefficient and impedance, and we could manually perform the complex math to find the impedance. Although programmable calculators and computers take the drudgery out of doing the math, a single number does not always give us the complete picture. In addition, impedance almost certainly changes with frequency, so even if we did all the math, we would end up with a table of numbers that may be difficult to interpret.

A simple, graphical method solves this problem. Let's first plot reflection coefficient using a polar display. For positive resistance, the absolute magnitude of Γ varies from zero (perfect load) to unity (full reflection) at some angle. So we have a unit circle, which marks the boundary of the polar plane shown on the slide. An open would plot at $1 \angle 0^\circ$; a short at $1 \angle 180^\circ$; a perfect load at the center, and so on. How do we get from the polar data to impedance

graphically? Since there is a one-to-one correspondence between complex reflection coefficient and impedance, we can map one plane onto the other. If we try to map the polar plane onto the rectilinear impedance plane, we find that we have problems. First of all, the rectilinear plane does not have values to infinity. Second, circles of constant reflection coefficient are concentric on the polar plane but not on the rectilinear plane, making it difficult to make judgments regarding two different impedances. Finally, phase angles plot as radii on the polar plane but plot as arcs on the rectilinear plane, making it difficult to pinpoint.

The proper solution was first used in the 1930's, when Phillip H. Smith mapped the impedance plane onto the polar plane, creating the chart that bears his name (the venerable *Smith chart*). Since unity at zero degrees on the polar plane represents infinite impedance, both plus and minus infinite reactances, as well as infinite resistance can be plotted. On the Smith chart, the vertical lines on the rectilinear plane that indicate values of constant resistance map to circles, and the horizontal lines that indicate values of constant reactance map to arcs. Z_0 maps to the exact center of the chart.

In general, Smith charts are normalized to Z_0 ; that is, the impedance values are divided by Z_0 . The chart is then independent of the characteristic impedance of the system in question. Actual impedance values are derived by multiplying the indicated value by Z_0 . For example, in a 50-ohm system, a normalized value of $0.3 - j0.15$ becomes $15 - j7.5$ ohms; in a 75-ohm system, $22.5 - j11.25$ ohms.

Fortunately, we no longer have to go through the exercise ourselves. Our network analyzer can display the Smith chart, plot measured data on it, and provide adjustable markers that show the calculated impedance at the marked point in a several marker formats.



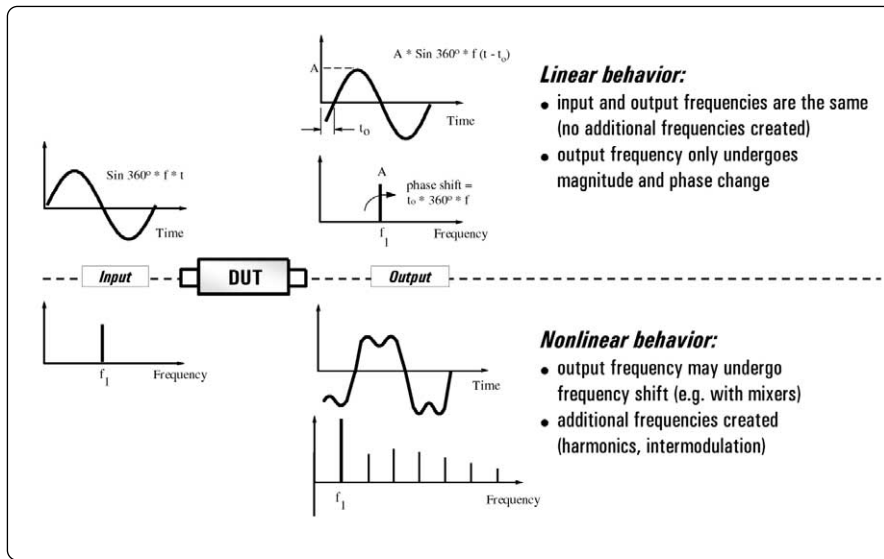
$$\text{Transmission Coefficient} = T = \frac{V_{\text{Transmitted}}}{V_{\text{Incident}}} = \tau \angle \phi$$

$$\text{Insertion Loss (dB)} = -20 \log \left| \frac{V_{\text{Trans}}}{V_{\text{Inc}}} \right| = -20 \log \tau$$

$$\text{Gain (dB)} = 20 \log \left| \frac{V_{\text{Trans}}}{V_{\text{Inc}}} \right| = 20 \log \tau$$

Transmission Parameters

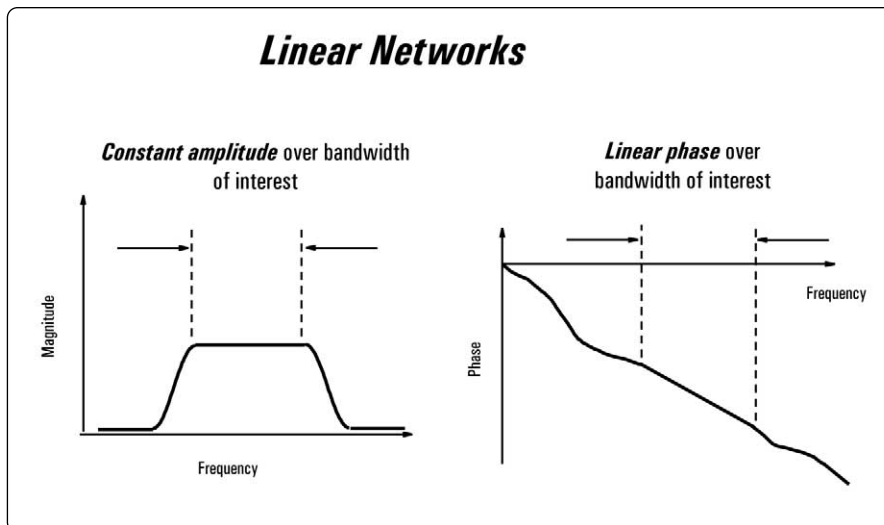
Transmission coefficient T is defined as the transmitted voltage divided by the incident voltage. If $|V_{\text{trans}}| > |V_{\text{inc}}|$, the DUT has gain, and if $|V_{\text{trans}}| < |V_{\text{inc}}|$, the DUT exhibits attenuation or insertion loss. When insertion loss is expressed in dB, a negative sign is added in the definition so that the loss value is expressed as a positive number. The phase portion of the transmission coefficient is called insertion phase. There is more to transmission than simple gain or loss. In communications systems, signals are time varying – they occupy a given bandwidth and are made up of multiple frequency components. It is important then to know to what extent the DUT alters the makeup of the signal, thereby causing signal distortion. While we often think of distortion as only the result of nonlinear networks, we will see shortly that linear networks can also cause signal distortion.



Linear Versus Nonlinear Behavior

Before we explore linear signal distortion, let's review the differences between linear and nonlinear behavior. Devices that behave linearly only impose magnitude and phase changes on input signals. Any sinusoid appearing at the input will also appear at the output at the same frequency. No new signals are created. When a single sinusoid is passed through a linear network, we don't consider amplitude and phase changes as distortion. However, when a complex, time-varying signal is passed through a linear network, the amplitude and phase shifts can dramatically distort the time-domain waveform.

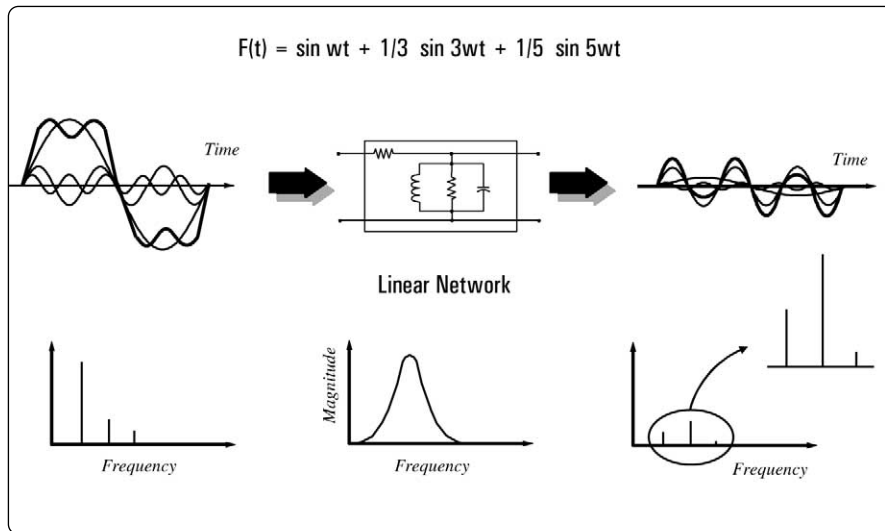
Non-linear devices can shift input signals in frequency (a mixer for example) and/or create new signals in the form of harmonics or intermodulation products. Many components that behave linearly under most signal conditions can exhibit nonlinear behavior if driven with a large enough input signal. This is true for both passive devices like filters and even connectors, and active devices like amplifiers.



Criteria for Distortionless Transmission

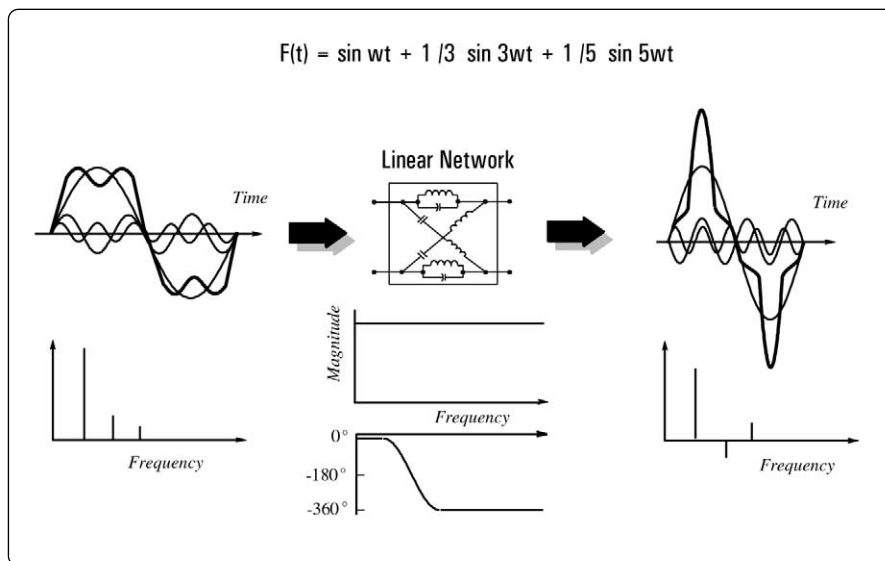
Now let's examine how linear networks can cause signal distortion. There are three criteria that must be satisfied for linear distortionless transmission. First, the amplitude (magnitude) response of the device or system must be flat over the bandwidth of interest. This means all frequencies within the bandwidth will be attenuated identically. Second, the phase response must be linear over the bandwidth of interest. And last, the device must exhibit a "minimum-phase response", which means that at 0 Hz (DC), there is 0° phase shift ($0^\circ \pm n \times 180^\circ$ is okay if we don't mind an inverted signal).

How can magnitude and phase distortion occur? The following two examples will illustrate how both magnitude and phase responses can introduce linear signal distortion.



Magnitude Variation with Frequency

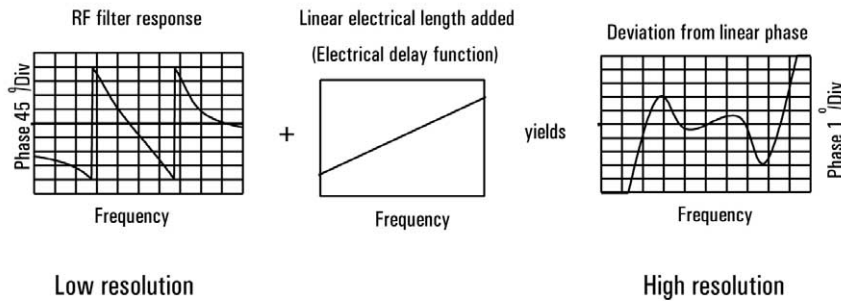
Here is an example of a square wave (consisting of three sinusoids) applied to a bandpass filter. The filter imposes a non-uniform amplitude change to each frequency component. Even though no phase changes are introduced, the frequency components no longer sum to a square wave at the output. The square wave is now severely distorted, having become more sinusoidal in nature.



Phase Variation with Frequency

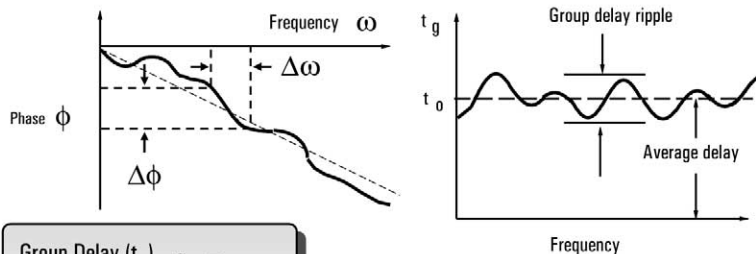
Let's apply the same square wave to another filter. Here, the third harmonic undergoes a 180° phase shift, but the other components are not phase shifted. All the amplitudes of the three spectral components remain the same (filters which only affect the phase of signals are called allpass filters). The output is again distorted, appearing very impulsive this time.

**Use electrical delay to remove
linear portion of phase response**



Deviation From Linear Phase

Now that we know insertion phase versus frequency is a very important characteristic of a component, let's see how we would measure it. Looking at insertion phase directly is usually not very useful. This is because the phase has a negative slope with respect to frequency due to the electrical length of the device (the longer the device, the greater the slope). Since it is only the deviation from linear phase which causes distortion, it is desirable to remove the linear portion of the phase response. This can be accomplished by using the electrical delay feature of the network analyzer to cancel the electrical length of the DUT. This results in a high-resolution display of phase distortion (deviation from linear phase).



$$\text{Group Delay (t)}_g = \frac{-d\phi}{d\omega} = \frac{-1}{360^\circ} * \frac{d\phi}{df}$$

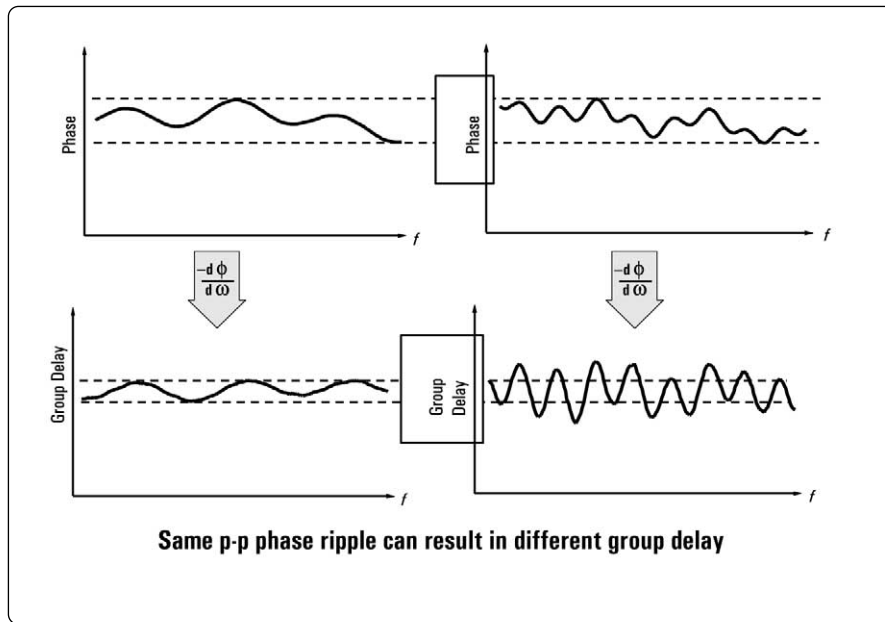
ϕ in radians
 ω in radians/sec
 ϕ in degrees
 f in Hertz ($\omega = 2\pi f$)

- group-delay ripple indicates phase distortion
- average delay indicates electrical length of DUT
- aperture of measurement is very important

Group Delay

Another useful measure of phase distortion is group delay. Group delay is a measure of the transit time of a signal through the device under test, versus frequency. Group delay is calculated by differentiating the insertion phase response of the DUT versus frequency. Another way to say this is that group delay is a measure of the slope of the transmission phase response. The linear portion of the phase response is converted to a constant value (representing the average signal-transit time) and deviations from linear phase are transformed into deviations from constant group delay. The variations in group delay cause signal distortion, just as deviations from linear phase cause distortion. Group delay is just another way to look at linear phase distortion.

When specifying or measuring group delay, it is important to quantify the aperture in which the measurement is made. The aperture is defined as the frequency delta used in the differentiation process (the denominator in the group-delay formula). As we widen the aperture, trace noise is reduced but less group-delay resolution is available (we are essentially averaging the phase response over a wider window). As we make the aperture more narrow, trace noise increases but we have more measurement resolution.



Why Measure Group Delay?

Why are both deviation from linear phase and group delay commonly measured? Depending on the device, both may be important. Specifying a maximum peak-to-peak value of phase ripple is not sufficient to completely characterize a device since the slope of the phase ripple is dependent on the number of ripples which occur over a frequency range of interest. Group delay takes this into account since it is the differentiated phase response. Group delay is often a more easily interpreted indication of phase distortion.

The plot above shows that the same value of peak-to-peak phase ripple can result in substantially different group delay responses. The response on the right with the larger group-delay variation would cause more signal distortion.

Using parameters (H, Y, Z, S) to characterize devices:

- gives linear behavioral model of our device
- measure parameters (e.g. voltage and current) versus frequency under various source and load conditions (e.g. short and open circuits)
- compute device parameters from measured data
- predict circuit performance under any source and load conditions

H parameters

$$\begin{aligned} V_1 &= h_{11}I_1 + h_{12}V_2 \\ I_2 &= h_{21}I_1 + h_{22}V_2 \end{aligned}$$

Y parameters

$$\begin{aligned} I_1 &= y_{11}V_1 + y_{12}V_2 \\ I_2 &= y_{21}V_1 + y_{22}V_2 \end{aligned}$$

Z parameters

$$\begin{aligned} V_1 &= z_{11}I_1 + z_{12}I_2 \\ V_2 &= z_{21}I_1 + z_{22}I_2 \end{aligned}$$



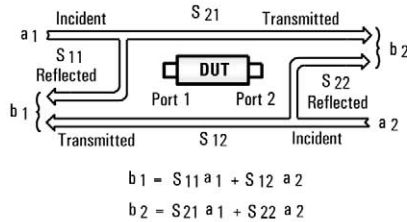
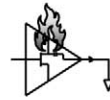
$$h_{11} = \left. \frac{V_1}{I_1} \right|_{V_2=0} \quad (\text{requires short circuit})$$

$$h_{12} = \left. \frac{V_1}{V_2} \right|_{I_1=0} \quad (\text{requires open circuit})$$

Characterizing Unknown Devices

In order to completely characterize an unknown linear two-port device, we must make measurements under various conditions and compute a set of parameters. These parameters can be used to completely describe the electrical behavior of our device (or network), even under source and load conditions other than when we made our measurements. For low-frequency characterization of devices, the three most commonly measured parameters are the H, Y and Z-parameters. All of these parameters require measuring the total voltage or current as a function of frequency at the input or output nodes (ports) of the device. Furthermore, we have to apply either open or short circuits as part of the measurement. Extending measurements of these parameters to high frequencies is not very practical.

- relatively easy to **obtain** at high frequencies
 - measure voltage traveling waves with a vector network analyzer
 - don't need shorts/opens which can cause active devices to oscillate or self-destruct
- relate to **familiar** measurements (gain, loss, reflection coefficient ...)
- can **cascade** S-parameters of multiple devices to predict system performance
- can **compute** H, Y, or Z parameters from S-parameters if desired
- can easily import and use S-parameter files in our **electronic-simulation** tools

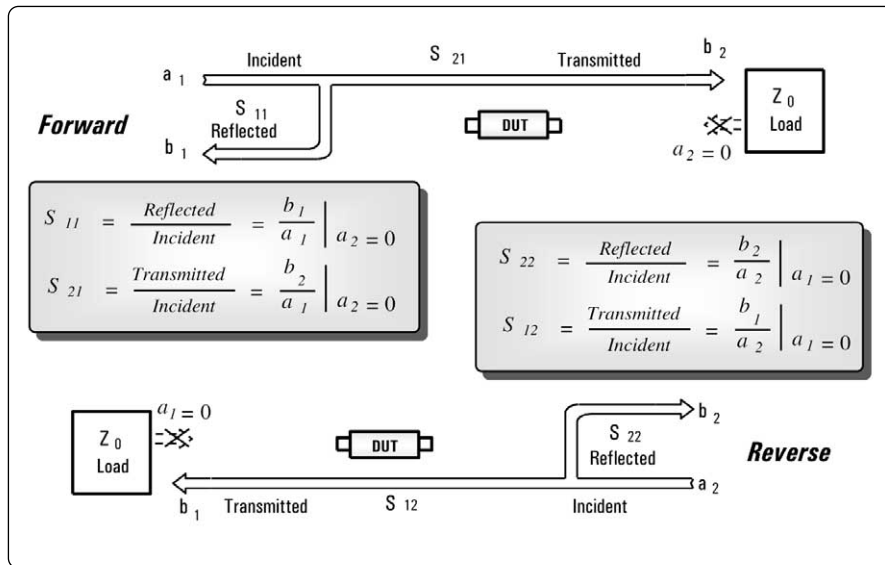


Why Use S-Parameters?

At high frequencies, it is very hard to measure total voltage and current at the device ports. One cannot simply connect a voltmeter or current probe and get accurate measurements due to the impedance of the probes themselves and the difficulty of placing the probes at the desired positions. In addition, active devices may oscillate or self-destruct with the connection of shorts and opens.

Clearly, some other way of characterizing high-frequency networks is needed that doesn't have these drawbacks. That is why scattering or S-parameters were developed. S-parameters have many advantages over the previously mentioned H, Y or Z-parameters. They relate to familiar measurements such as gain, loss, and reflection coefficient. They are defined in terms of voltage traveling waves, which are relatively easy to measure. S-parameters don't require connection of undesirable loads to the device under test. The measured S-parameters of multiple devices can be cascaded to predict overall system performance. If desired, H, Y, or Z-parameters can be derived from S-parameters. And very important for RF design, S-parameters are easily imported and used for circuit simulations in electronic-design automation (EDA) tools like Agilent's Advanced Design System (ADS). S-parameters are the shared language between simulation and measurement.

An N-port device has N^2 S-parameters. So, a two-port device has four S-parameters. The numbering convention for S-parameters is that the first number following the "S" is the port where the signal emerges, and the second number is the port where the signal is applied. So, S21 is a measure of the signal coming out port 2 relative to the RF stimulus entering port 1. When the numbers are the same (e.g., S11), it indicates a reflection measurement, as the input and output ports are the same. The incident terms (a_1 , a_2) and output terms (b_1 , b_2) represent voltage traveling waves.



Measuring S-Parameters

S_{11} and S_{21} are determined by measuring the magnitude and phase of the incident, reflected and transmitted voltage signals when the output is terminated in a perfect Z_0 (a load that equals the characteristic impedance of the test system). This condition guarantees that a_2 is zero, since there is no reflection from an ideal load. S_{11} is equivalent to the input complex reflection coefficient or impedance of the DUT, and S_{21} is the forward complex transmission coefficient. Likewise, by placing the source at port 2 and terminating port 1 in a perfect load (making a_1 zero), S_{22} and S_{12} measurements can be made. S_{22} is equivalent to the output complex reflection coefficient or output impedance of the DUT, and S_{12} is the reverse complex transmission coefficient.

The accuracy of S-parameter measurements depends greatly on how good a termination we apply to the load port (the port not being stimulated). Anything other than a perfect load will result in a_1 or a_2 not being zero (which violates the definition for S-parameters). When the DUT is connected to the test ports of a network analyzer and we don't account for imperfect test-port match, we have not done a very good job satisfying the condition of a perfect termination. For this reason, two-port error correction, which corrects for source and load match, is very important for accurate S-parameter measurements (two-port correction is covered in the calibration section).

S_{11} = forward reflection coefficient (*input match*)
 S_{22} = reverse reflection coefficient (*output match*)
 S_{21} = forward transmission coefficient (*gain or loss*)
 S_{12} = reverse transmission coefficient (*isolation*)

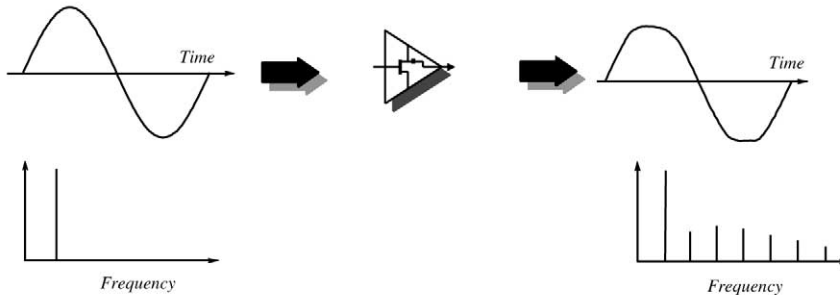
Remember, *S*-parameters are inherently complex, linear quantities -- however, we often express them in a log-magnitude format

Equating S-Parameters with Common Measurement Terms

S-parameters are essentially the same parameters as some of the terms we have mentioned before, such as input match and insertion loss. It is important to separate the fundamental definition of S-parameters and the format in which they are often displayed. S-parameters are inherently complex, linear quantities. They are expressed as real-and-imaginary or magnitude-and-phase pairs. However, it isn't always very useful to view them as linear pairs. Often we want to look only at the magnitude of the S-parameter (for example, when looking at insertion loss or input match), and often, a logarithmic display is most useful. A log-magnitude format lets us see far more dynamic range than a linear format.

Nonlinear Networks

- Saturation, crossover, intermodulation, and other nonlinear effects can cause signal distortion
- Effect on system depends on amount and type of distortion and system architecture



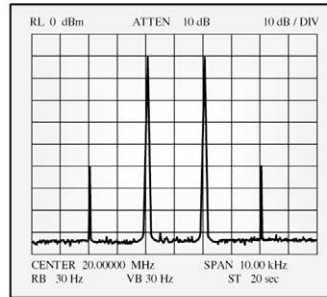
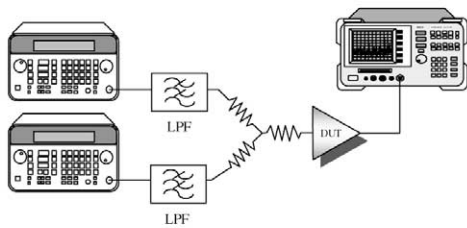
Criteria for Distortionless Transmission

We have just seen how linear networks can cause distortion. Devices which behave nonlinearly also introduce distortion. The example above shows an amplifier that is overdriven, causing the signal at the output to "clip" due to saturation in the amplifier. Because the output signal is no longer a pure sinusoid, harmonics are present at integer multiples of the input frequency.

Passive devices can also exhibit nonlinear behavior at high power levels. A common example is an L-C filter that uses inductors made with magnetic cores. Magnetic materials often display hysteresis effects, which are highly nonlinear. Another example are the connectors used in the antenna path of a cellular-phone base station. The metal-to-metal contacts (especially if water and corrosion salts are present) combined with the high-power transmitted signals can cause a diode effect to occur, producing very low-level intermodulation products. Although the level of the intermodulation products is usually quite small, they can be significant compared to the low signal strength of the received signals, causing interference problems.

Most common measurements:

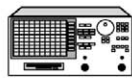
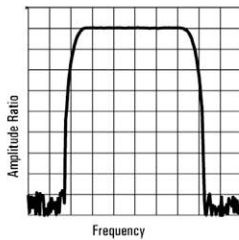
- using a **network analyzer** and power sweeps
 - gain compression
 - AM to PM conversion
- using a **spectrum analyzer** + source(s)
 - harmonics, particularly second and third
 - intermodulation products resulting from two or more RF carriers



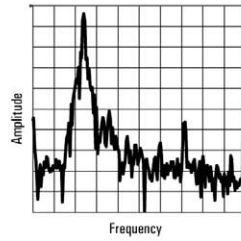
Measuring Nonlinear Behavior

So far, we've focused most of our attention on linear swept-frequency characterization, which is needed for both passive and active devices. We already know that nonlinear behavior is important to quantify, as it can cause severe signal distortion. The most common nonlinear measurements are gain compression and AM-to-PM conversion (usually measured with network analyzers and power sweeps), and harmonic and intermodulation distortion (usually measured with spectrum analyzers and signal sources). We will cover swept-power measurements using a network analyzer in more detail in the typical-measurements section of this presentation. The slide shows how intermodulation distortion is typically measured using two signal sources and a spectrum analyzer as a receiver.

Network and Spectrum Analyzers?



Measures
known signal



Measures
unknown
signals

Network analyzers:

- measure components, devices, circuits, sub-assemblies
- contain source and receiver
- display ratioed amplitude and phase (frequency or power sweeps)
- offer advanced error correction

Spectrum analyzers:

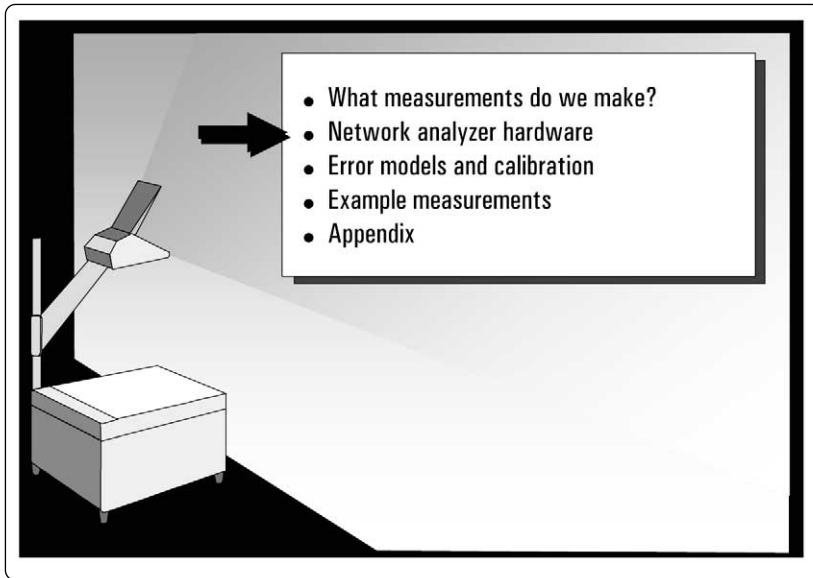
- measure signal amplitude characteristics (carrier level, sidebands, harmonics...)
- can demodulate (& measure) complex signals
- are receivers only (single channel)
- can be used for scalar component test (*no phase*) with tracking gen. or ext. source(s)

What is the Difference Between Network and Spectrum Analyzers?

Now that we have seen some of the measurements that are commonly done with network and spectrum analyzers, it might be helpful to review the main differences between these instruments. Although they often both contain tuned receivers operating over similar frequency ranges, they are optimized for very different measurement applications.

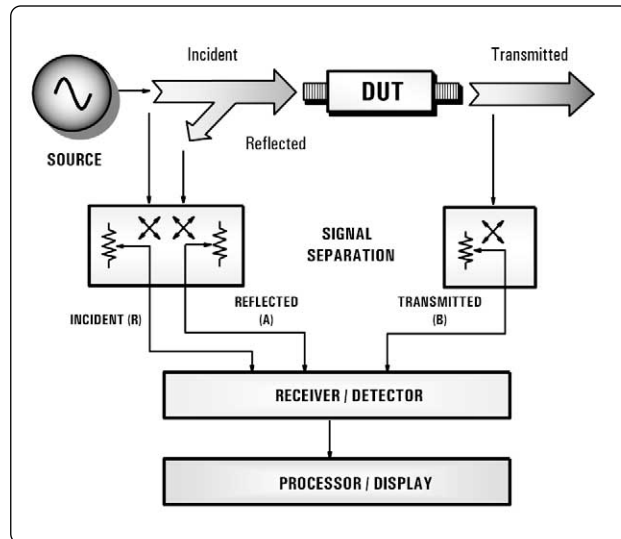
Network analyzers are used to measure components, devices, circuits, and sub-assemblies. They contain both a source and multiple receivers, and generally display *ratioed* amplitude and phase information (frequency or power sweeps). A network analyzer is always looking at a *known* signal (in terms of frequency), since it is a stimulus-response system. With network analyzers, it is harder to get an (accurate) trace on the display, but very easy to interpret the results. With vector-error correction, network analyzers provide much higher measurement accuracy than spectrum analyzers.

Spectrum analyzers are most often used to measure signal characteristics such as carrier level, sidebands, harmonics, phase noise, etc., on *unknown* signals. They are most commonly configured as a single-channel receiver, without a source. Because of the flexibility needed to analyze signals, spectrum analyzers generally have a much wider range of IF bandwidths available than most network analyzers. Spectrum analyzers are often used with external sources for nonlinear stimulus/response testing. When combined with a tracking generator, spectrum analyzers can be used for scalar component testing (magnitude versus frequency, but no phase measurements). With spectrum analyzers, it is easy to get a trace on the display, but interpreting the results can be much more difficult than with a network analyzer.



Agenda

In this next section, we will look at the main parts of a network analyzer.



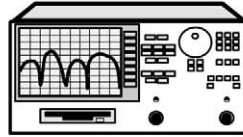
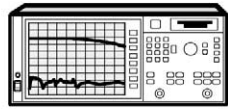
Generalized Network Analyzer Block Diagram

Here is a generalized block diagram of a network analyzer, showing the major signal-processing sections. In order to measure the incident, reflected and transmitted signal, four sections are required:

- Source for stimulus
- Signal-separation devices
- Receivers that downconvert and detect the signals
- Processor/display for calculating and reviewing the results

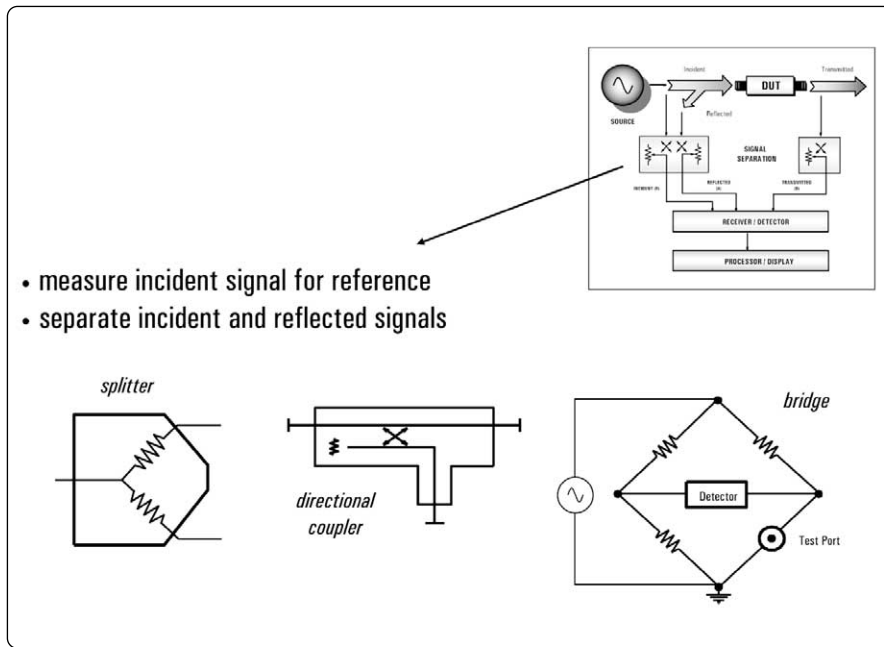
We will briefly examine each of these sections. More detailed information about the signal separation devices and receiver section are in the appendix.

- Supplies stimulus for system
- Swept frequency or power
- Traditionally NAs used separate source
- Most Agilent analyzers sold today have ***integrated, synthesized*** sources



Source

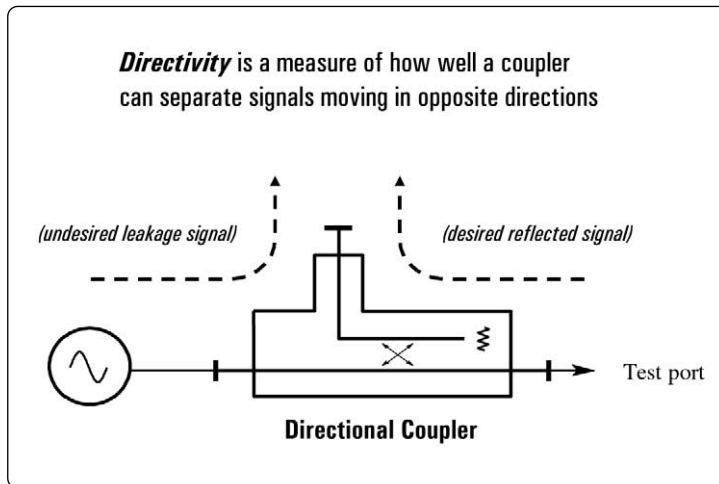
The signal source supplies the stimulus for our stimulus-response test system. We can either sweep the frequency of the source or sweep its power level. Traditionally, network analyzers used a separate source. These sources were either based on open-loop voltage-controlled oscillators (VCOs) which were cheaper, or more expensive synthesized sweepers which provided higher performance, especially for measuring narrowband devices. Excessive phase noise on open-loop VCOs degrades measurement accuracy considerably when measuring narrowband components over small frequency spans. Most network analyzers that Agilent sells today have integrated, synthesized sources, providing excellent frequency resolution and stability.



Signal Separation

The next major area we will cover is the signal separation block. The hardware used for this function is generally called the “test set”. The test set can be a separate box or integrated within the network analyzer. There are two functions that our signal-separation hardware must provide. The first is to measure a portion of the incident signal to provide a reference for ratioing. This can be done with splitters or directional couplers. Splitters are usually resistive. They are non-directional devices (more on directionality later) and can be very broadband. The trade-off is that they usually have 6 dB or more of loss in each arm. Directional couplers have very low insertion loss (through the main arm) and good isolation and directivity. They are generally used in microwave network analyzers, but their inherent high-pass response makes them unusable below 40 MHz or so.

The second function of the signal-splitting hardware is to separate the incident (forward) and reflected (reverse) traveling waves at the input of our DUT. Again, couplers are ideal in that they are directional, have low loss, and high reverse isolation. However, due to the difficulty of making truly broadband couplers, bridges are often used instead. Bridges work down to DC, but have more loss, resulting in less signal power delivered to the DUT. See the appendix for a more complete description of how a directional bridge works.



Directivity

Unfortunately, real signal-separation devices are never perfect. For example, let's take a closer look at the actual performance of a 3-port directional coupler.

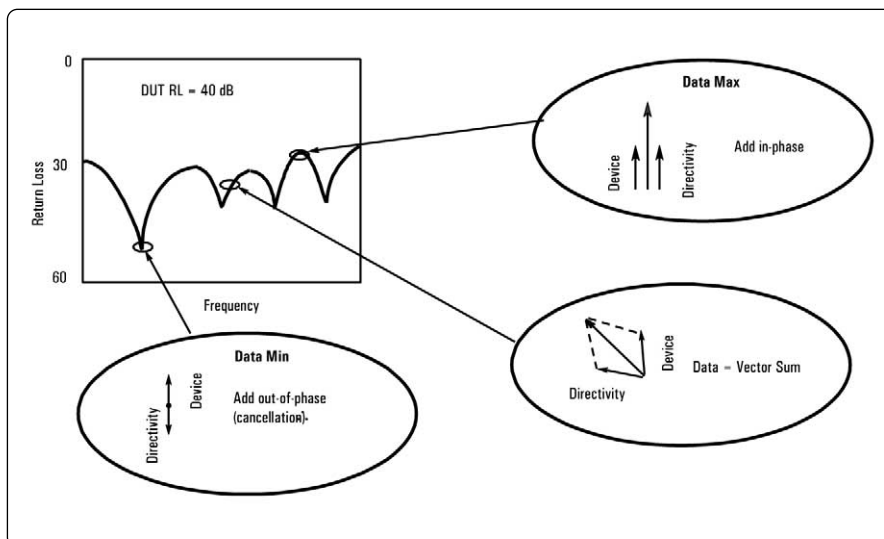
Ideally, a signal traveling in the coupler's reverse direction will not appear at all at the coupled port. In reality, however, some energy does leak through to the coupled arm, as a result of finite isolation.

One of the most important parameter for couplers is their directivity. Directivity is a measure of a coupler's ability to separate signals flowing in opposite directions within the coupler. It can be thought of as the dynamic range available for reflection measurements. Directivity can be defined as:

Directivity (dB) = Isolation (dB) - Forward Coupling Factor (dB) - Loss (through-arm) (dB)

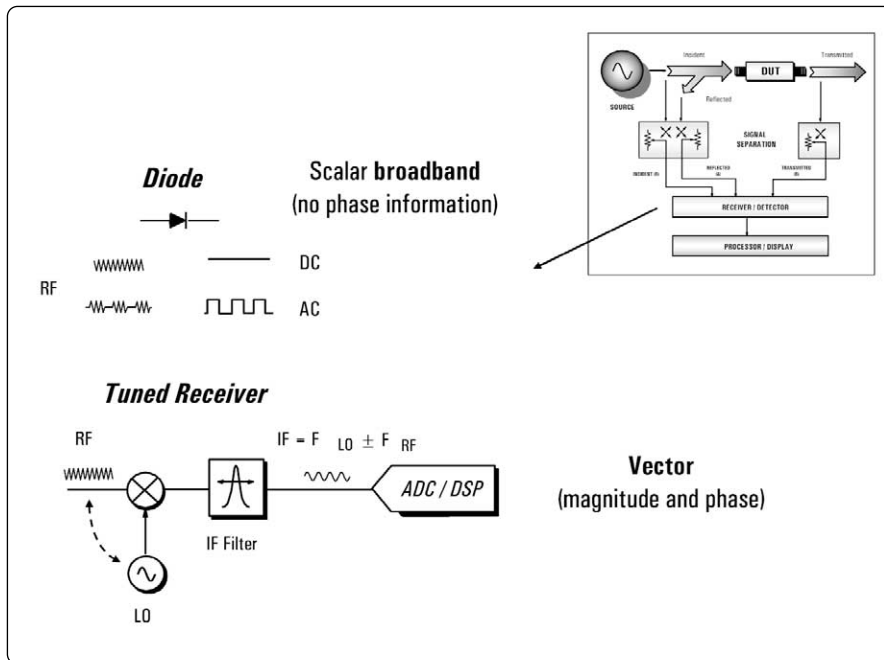
The appendix contains a slide showing how adding attenuation to the ports of a coupler can affect the effective directivity of a system (such as a network analyzer) that uses a directional coupler.

As we will see in the next slide, finite directivity adds error to our measured results.



Interaction of Directivity with the DUT (Without Error Correction)

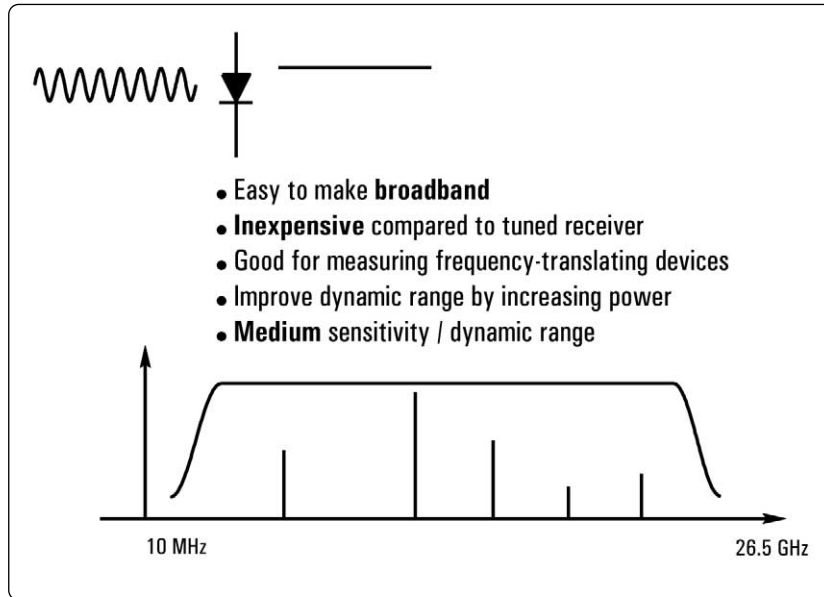
Directivity error is the main reason we see a large ripple pattern in many measurements of return loss. At the peaks of the ripple, directivity is adding in phase with the reflection from the DUT. In some cases, directivity will cancel the DUT's reflection, resulting in a sharp dip in the response.



Detector Types

The next portion of the network analyzer we'll look at is the signal-detection block. There are two basic ways of providing signal detection in network analyzers. Diode detectors convert the RF signal level to a proportional DC level. If the stimulus signal is amplitude modulated, the diode strips the RF carrier from the modulation (this is called AC detection). Diode detection is inherently scalar, as phase information of the RF carrier is lost.

The tuned receiver uses a local oscillator (LO) to mix the RF down to a lower "intermediate" frequency (IF). The LO is either locked to the RF or the IF signal so that the receivers in the network analyzer are always tuned to the RF signal present at the input. The IF signal is bandpass filtered, which narrows the receiver bandwidth and greatly improves sensitivity and dynamic range. Modern analyzers use an analog-to-digital converter (ADC) and digital-signal processing (DSP) to extract magnitude and phase information from the IF signal. The tuned-receiver approach is used in vector network analyzers and spectrum analyzers.

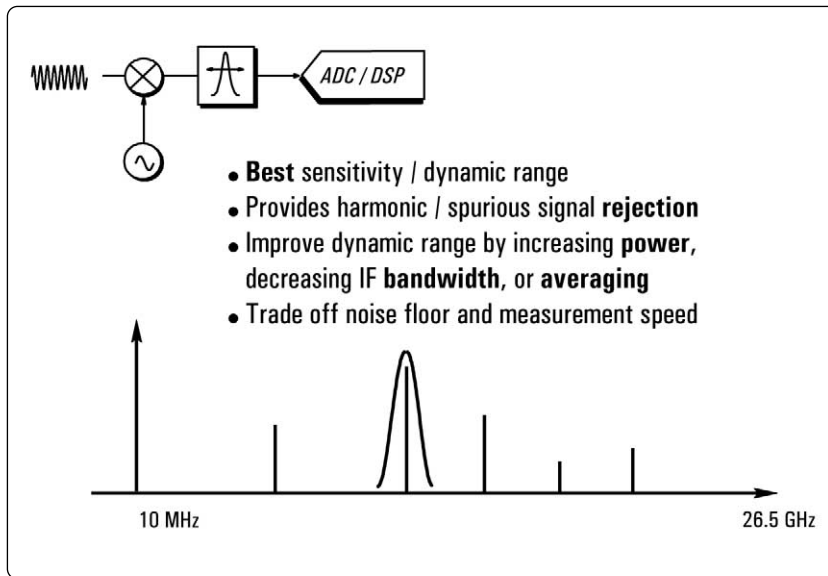


Broadband Diode Detection

The two main advantages of diode detectors are that they provide broadband frequency coverage (< 10 MHz on the low end to > 26.5 GHz at the high end) and they are inexpensive compared to a tuned receiver. Diode detectors provide medium sensitivity and dynamic range: they can measure signals to -60 dBm or so and have a dynamic range around 60 to 75 dB, depending on the detector type. Their broadband nature limits their sensitivity and makes them sensitive to source harmonics and other spurious signals. Dynamic range is improved in measurements by increasing input power.

AC detection eliminates the DC drift of the diode as an error source, resulting in more accurate measurements. This scheme also reduces noise and other unwanted signals. The major benefit of DC detection is that there is no modulation of the RF signal, which can have adverse effects on the measurement of some devices. Examples include amplifiers with AGC or large DC gain, and narrowband filters.

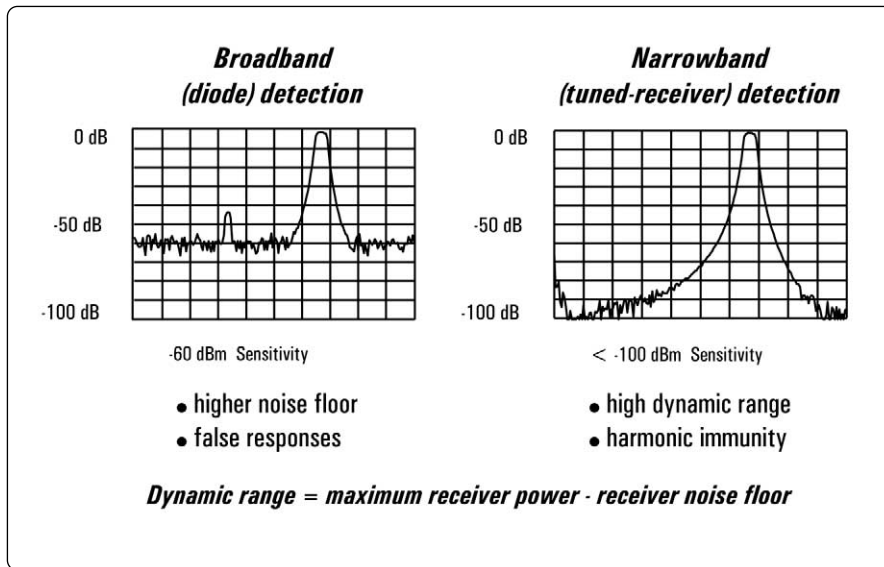
One application where broadband diode detectors are very useful is measuring frequency-translating devices, particularly those with internal LOs.



Narrowband Detection – Tuned Receiver

Tuned receivers provide the best sensitivity and dynamic range, and also provide harmonic and spurious-signal rejection. The narrow IF filter produces a considerably lower noise floor, resulting in a significant sensitivity improvement. For example, a microwave vector network analyzer (using a tuned receiver) might have a 3 kHz IF bandwidth, where a scalar analyzer's diode detector noise bandwidth might be 26.5 GHz. Measurement dynamic range is improved with tuned receivers by increasing input power, by decreasing IF bandwidth, or by averaging. The latter two techniques provide a trade off between noise floor and measurement speed. Averaging reduces the noise floor of the network analyzer (as opposed to just reducing the noise excursions as happens when averaging spectrum analyzer data) because we are averaging complex data. Without phase information, averaging does not improve analyzer sensitivity.

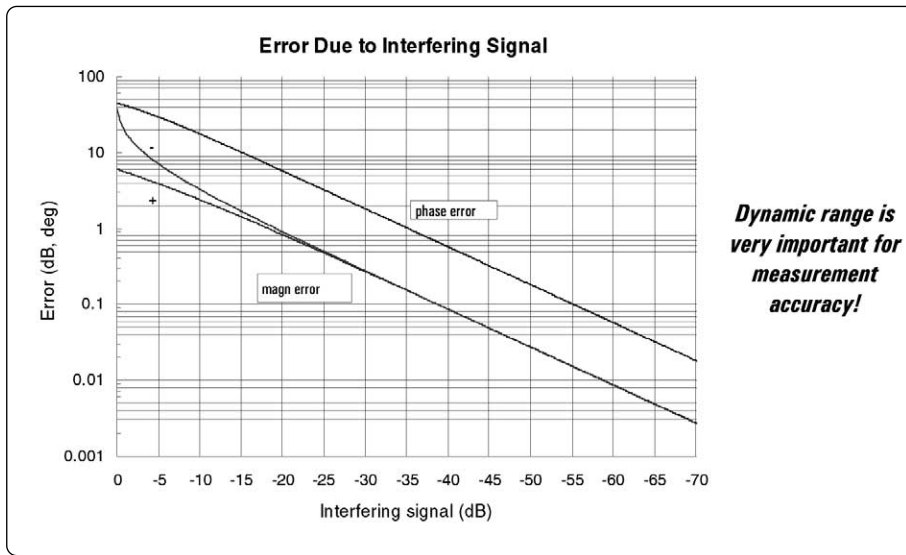
The same narrowband nature of tuned receivers that produces increased dynamic range also eliminates harmonic and spurious responses. As was mentioned earlier, the RF signal is downconverted and filtered before it is measured. The harmonics associated with the source are also downconverted, but they appear at frequencies outside the IF bandwidth and are therefore removed by filtering.



Comparison of Receiver Techniques

Dynamic range is generally defined as the maximum power the receiver can accurately measure minus the receiver noise floor. There are many applications requiring large dynamic range. One of the most common is measuring filter stopband performance. As you can see here, at least 80 dB dynamic range is needed to properly characterize the rejection characteristics of this filter. The plots show a typical narrowband filter measured on an 8757 scalar network analyzer and on an 8510 vector network analyzer. Notice that the filter exhibits 90 dB of rejection but the scalar analyzer is unable to measure it because of its higher noise floor.

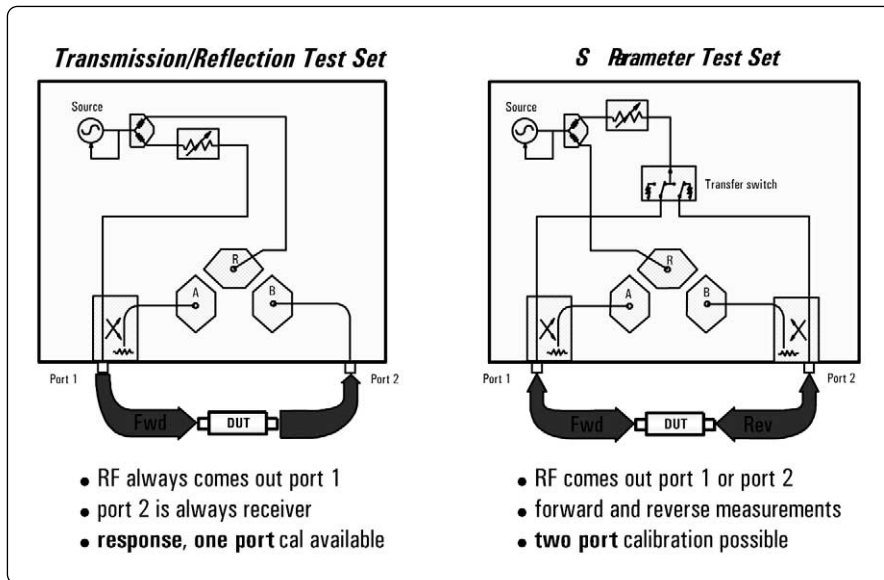
In the case where the scalar network analyzer was used with broadband diode detection, a harmonic from the source created a “false” response. For example, at some point on a broadband sweep, the second harmonic of the source might fall within the passband of the filter. If this occurs, the detector will register a response, even though the stopband of the filter is severely attenuating the frequency of the fundamental. This response from the second harmonic would show on the display at the frequency of the fundamental. On the tuned receiver, a false signal such as this would be filtered away and would not appear on the display. Note that source subharmonics and spurious outputs can also cause false display responses.



Dynamic Range and Accuracy

This plot shows the effect that interfering signals (sinusoids or noise) have on measurement accuracy. The magnitude error is calculated as $20 \cdot \log [1 \pm \text{interfering-signal}]$ and the phase error is calculated as $\text{arc-tangent} [\text{interfering-signal}]$, where the interfering signal is expressed in linear terms. Note that a 0 dB interfering signal results in (plus) 6 dB error when it adds in phase with the desired signal, and (negative) infinite error when it cancels the desired signal.

To get low measurement uncertainty, more dynamic range is needed than the device exhibits. For example, to get less than 0.1 dB magnitude error and less than 0.6 degree phase error, our noise floor needs to be more than 39 dB below our measured power levels (note that there are other sources of error besides noise that may limit measurement accuracy). To achieve that level of accuracy while measuring 80 dB of rejection would require 119 dB of dynamic range. One way to achieve this level is to average test data using a tuned-receiver based network analyzer.



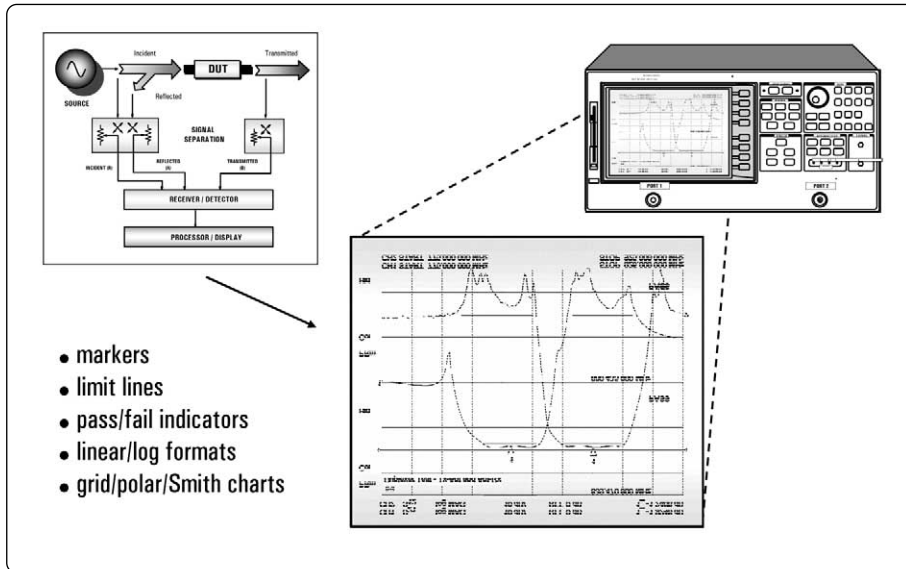
T/R Versus S-Parameter Test Sets

There are two basic types of test sets that are used with network analyzers. For transmission/reflection (T/R) test sets, the RF power always comes out of test port one and test port two is always connected to a receiver in the analyzer. To measure reverse transmission or output reflection of the DUT, we must disconnect it, turn it around, and re-connect it to the analyzer. T/R-based network analyzers offer only response and one-port calibrations, so measurement accuracy is not as good as that which can be achieved with S-parameter test sets. However, T/R-based analyzers are more economical. For the 8712, 8753 and 8720 families, Agilent uses the ET suffix to denote a T/R analyzer, and the ES suffix to denote an S-parameter analyzer.

S-parameter test sets allow both forward and reverse measurements on the DUT, which are needed to characterize all four S-parameters. RF power can come out of either test port one or two, and either test port can be connected to a receiver. S-parameter test sets also allow full two-port (12-term) error correction, which is the most accurate form available. S-parameter network analyzers provide more performance than T/R-based analyzers, but cost more due to extra RF components in the test set.

There are two different types of transfer switches that can be used in an S-parameter test set: solid-state and mechanical. Solid-state switches have the advantage of infinite lifetimes (assuming they are not damaged by too much power from the DUT). However, they are more lossy so they reduce the maximum output power of the network analyzer. Mechanical switches have very low loss and therefore allow higher output powers. Their main disadvantage is that eventually they wear out (after 5 million cycles or so). When using a network analyzer with mechanical switches, measurements are generally done in single-sweep mode, so the transfer switch is not continuously switching.

S-parameter test sets can have either a 3-receiver (shown on slide) or 4-receiver architecture. The 8753 series and standard 8720 series analyzers have a 3-receiver architecture. Option 400 adds a fourth receiver to 8720 series analyzers, to allow true TRL calibration. The 8510C family and the PNA Series uses a 4-receiver architecture. More detailed information of the two architecture is available in the appendix.



Processor/Display

The last major block of hardware in the network analyzer is the display/processor section. This is where the reflection and transmission data is formatted in ways that make it easy to interpret the measurement results. Most network analyzers have similar features such as linear and logarithmic sweeps, linear and log formats, polar plots, Smith charts, etc. Other common features are trace markers, limit lines, and pass/fail testing. Many of Agilent's network analyzers have specialized measurement features tailored to a particular market or application. One example is the E5100A/B, which has features specific to crystal-resonator manufacturers.

Simple: recall states

More powerful:

- **Test sequencing**
 - available on 8753/ 8720 families
 - keystroke recording
 - some advanced functions
- **IBASIC**
 - available on 8712 family
 - sophisticated programs
 - custom user interfaces
- **Windows compatible programs**
 - PNA Series:
 - use any Windows-compatible program
 - e.g. Visual Basic, VEE, LabView, C++, ...
 - ENA Series: VBA only

The screenshot shows a test sequencing program with a list of commands on the left and a graphical display on the right. The commands include:

- 1 ASSIGN @H8714 TO 800
- 2 OUTPUT @H8714:"SYST-PRES:" "WAI"
- 3 OUTPUT @H8714:"ABOR:INIT:CONT OFF:" "WAI"
- 4 OUTPUT @H8714:"DISP:ANN:FREQ1:MODE SSTOP"
- 5 OUTPUT @H8714:"DISP:ANN:FREQ1:MODE CSPAN"
- 6 OUTPUT @H8714:"SENS1:FREQ:CENT 175000000 HZ:" "WAI"
- 7 OUTPUT @H8714:"ABOR:INIT:CONT OFF:INIT1:" "WAI"
- 8 OUTPUT @H8714:"DISP:WIND1:TRAC:Y:AUTO ONCE"
- 9 OUTPUT @H8714:"CALC1:MARK1 ON"
- 10 OUTPUT @H8714:"CALC1:MARK1 OFF"
- 11 OUTPUT @H8714:"CALC1:MARK1 ON"
- 12 OUTPUT @H8714:"CALC1:MARK1 OFF"
- 13 OUTPUT @H8714:"CALC1:MARK1 ON"
- 14 OUTPUT @H8714:"CALC1:MARK1 OFF"
- 15 OUTPUT @H8714:"CALC1:MARK1 ON"
- 16 END

The graphical display shows a plot of the signal with various markers and a table of data points.

Internal Measurement Automation

All of Agilent's network analyzers offer some form of internal measurement automation. The most simple form is recall states. This is an easy way to set up the analyzer to a pre-configured measurement state, with all of the necessary instrument parameters.

More powerful automation can be achieved with test sequencing or Instrument BASIC (IBASIC). Test sequencing is available on the 8753/8720 families and provides keystroke recording and some advanced functions. IBASIC is available on the 8712ET/ES series and provides the user with sophisticated programs and custom user interfaces and measurement personalities.

The most powerful automation can be achieved with the PNA Series. Since these analyzers use Windows 2000, any PC-compatible programming language can be used to create an executable program. Popular programming languages such as Visual Basic, Visual C++, Agilent VEE, LabView or even Agilent BASIC for Windows can be used. The ENA series features Visual Basic Applications (VBA) as its automation language. VBA makes it easy to automate measurements and create intuitive Graphical User Interfaces (GUIs).

	PNA series • 3, 6, 9 GHz - 2, 3 ports • highest dynamic range • advanced LAN connectivity • internal Windows automation • program via SCPI or COM/DCOM		8753ET/ES series • 3, 6 GHz - 50/75 Ω • rich feature set • frequency offset and harmonic sweeps
	ENA series • 3, 8.5 GHz • 2, 3, 4, 7, 9 ports • excellent RF performance • balanced measurements • internal VBA automation		8712ET/ES series • 1.3, 3 GHz - 50/75 Ω • lowest cost • narrowband and broadband detection • IBASIC / LAN
	4395A series NA/SA combination • 500 MHz, 1.8 GHz • impedance-measuring option • fast, FFT-based spectrum analysis • IBASIC		E5100A/B • 180, 300 MHz • fast, small, economical • for crystals, resonators, filters

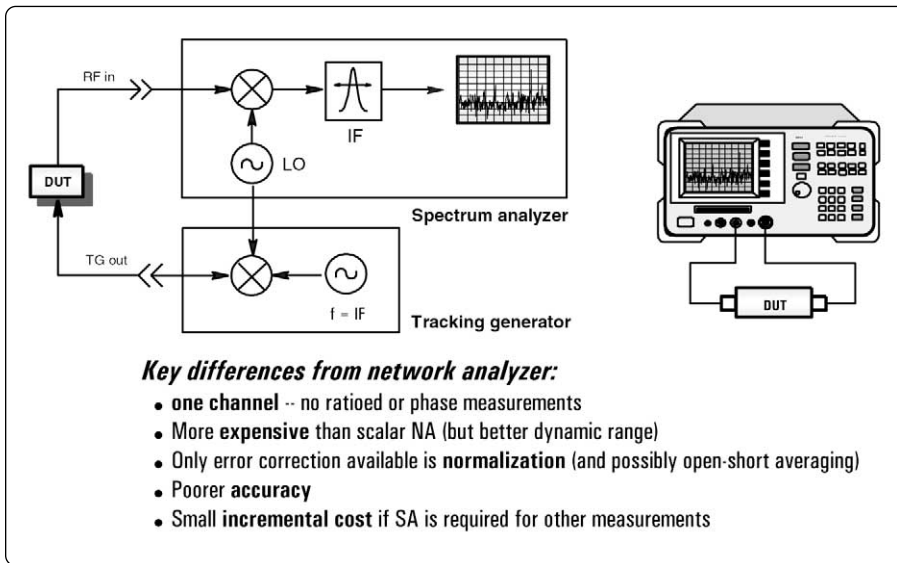
Agilent's Series of RF Vector Analyzers

Shown here is a summary of Agilent's RF families of vector network analyzers.

	8720ET/ES series • 13.5, 20, 40 GHz • economical • fast, small, integrated • test mixers, high-power amps		8510C series • 110 GHz in coax • modular, flexible • pulse systems • Tx/Rx module test
	PNA series • 20, 40, 50 GHz • highest dynamic range and speed • very low trace noise • advanced LAN connectivity • internal Windows automation • program via SCPI or COM/DCOM		

Agilent's Series of Microwave Vector Analyzers

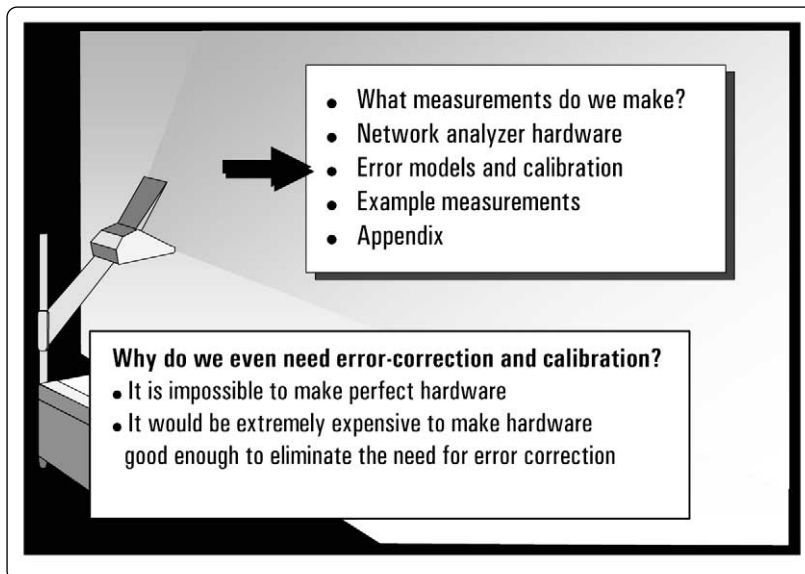
Shown here is a summary of Agilent's microwave families of vector network analyzers.



Spectrum Analyzer/Tracking Generator

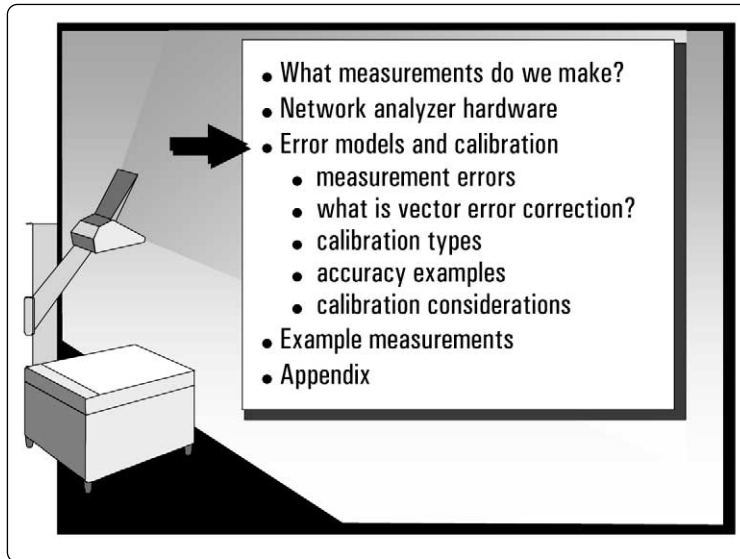
If the main difference between spectrum and network analyzers is a source, why don't we add a tracking generator (a source that tracks the tuned frequency of the spectrum analyzer) to our spectrum analyzer . . . then is it a network analyzer? Well, sort of.

A spectrum analyzer with a tracking generator can make swept scalar-magnitude measurements, but it is still a single-channel receiver. Therefore it cannot make ratioed or phase measurements. Also, the only error correction available is normalization (and possibly open-short averaging). The amplitude accuracy with a spectrum analyzer is roughly an order of magnitude worse than on a scalar network analyzer (dB vs. tenths of dB). Finally, a spectrum analyzer with a tracking generator costs more than a scalar network analyzer with similar frequency range, but it may be a small incremental cost to add a tracking generator if the spectrum analyzer is already needed for other spectrum-related measurements.



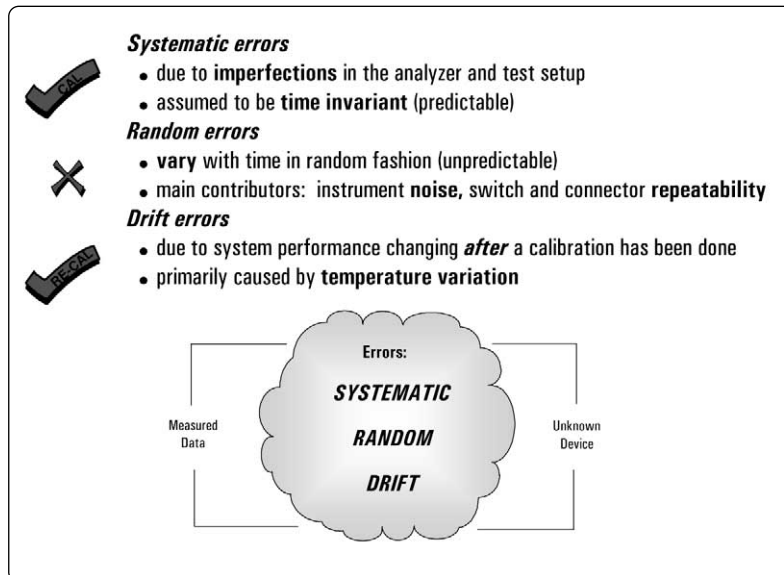
Agenda

In this next section, we will talk about the need for error correction and how it is accomplished. Why do we even need error-correction and calibration? It is impossible to make perfect hardware which obviously would not need any form of error correction. Even making the hardware good enough to eliminate the need for error correction for most devices would be extremely expensive. The best balance is to make the hardware as good as practically possible, balancing performance and cost. Error correction is then a very useful tool to improve measurement accuracy.



Calibration Topics

We will explain the sources of measurement error, how it can be corrected with calibration, and give accuracy examples using different calibration types.



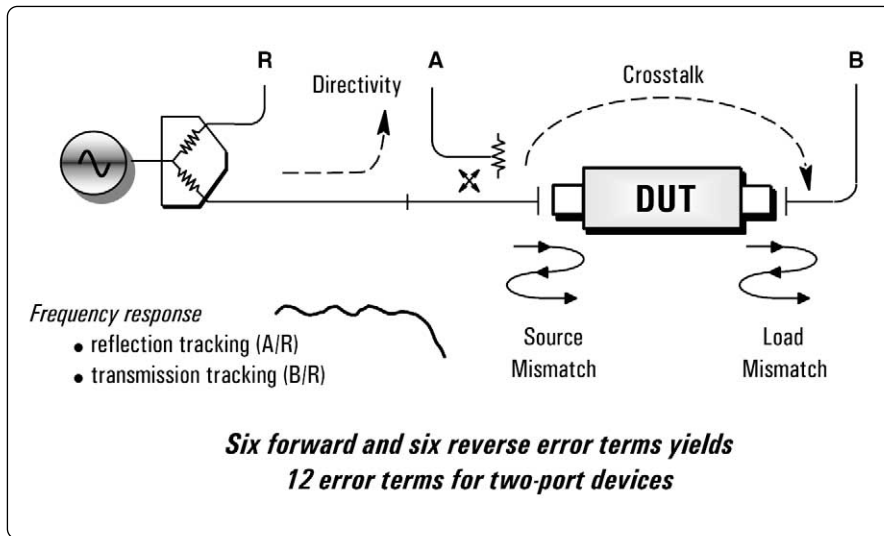
Measurement Error Modeling

Let's look at the three basic sources of measurement error: systematic, random and drift.

Systematic errors are due to imperfections in the analyzer and test setup. They are repeatable (and therefore predictable), and are assumed to be time invariant. Systematic errors are characterized during the calibration process and mathematically removed during measurements.

Random errors are unpredictable since they vary with time in a random fashion. Therefore, they cannot be removed by calibration. The main contributors to random error are instrument noise (source phase noise, sampler noise, IF noise).

Drift errors are due to the instrument or test-system performance changing *after* a calibration has been done. Drift is primarily caused by temperature variation and it can be removed by further calibration(s). The timeframe over which a calibration remains accurate is dependent on the rate of drift that the test system undergoes in the user's test environment. Providing a stable ambient temperature usually goes a long way towards minimizing drift.

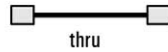


Systematic Measurement Errors

Shown here are the major systematic errors associated with network measurements. The errors relating to signal leakage are directivity and crosstalk. Errors related to signal reflections are source and load match. The final class of errors are related to frequency response of the receivers, and are called reflection and transmission tracking. The full two-port error model includes all six of these terms for the forward direction and the same six (with different data) in the reverse direction, for a total of twelve error terms. This is why we often refer to two-port calibration as twelve-term error correction

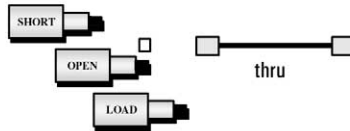
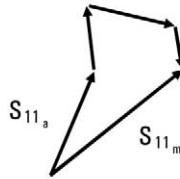
- **response (normalization)**

- simple to perform
- only corrects for tracking errors
- stores reference trace in memory, then does data divided by memory



- **vector**

- requires more standards
- requires an analyzer that can measure phase
- accounts for all major sources of systematic error



Types of Error Correction

The two main types of error correction that can be done are response (normalization) corrections and vector corrections. Response calibration is simple to perform, but only corrects for a few of the twelve possible systematic error terms (the tracking terms). Response calibration is essentially a normalized measurement where a reference trace is stored in memory, and subsequent measurement data is divided by this memory trace. A more advanced form of response calibration is open/short averaging for reflection measurements using broadband diode detectors. In this case, two traces are averaged together to derive the reference trace.

Vector-error correction requires an analyzer that can measure both magnitude and phase. It also requires measurements of more calibration standards. Vector-error correction can account for all the major sources of systematic error and can give very accurate measurements.

Note that a response calibration can be performed on a vector network analyzer, in which case we store a complex (vector) reference trace in memory, so that we can display normalized magnitude or phase data. This is not the same as vector-error correction however (and not as accurate), because we are not measuring and removing the individual systematic errors, all of which are complex or vector quantities.

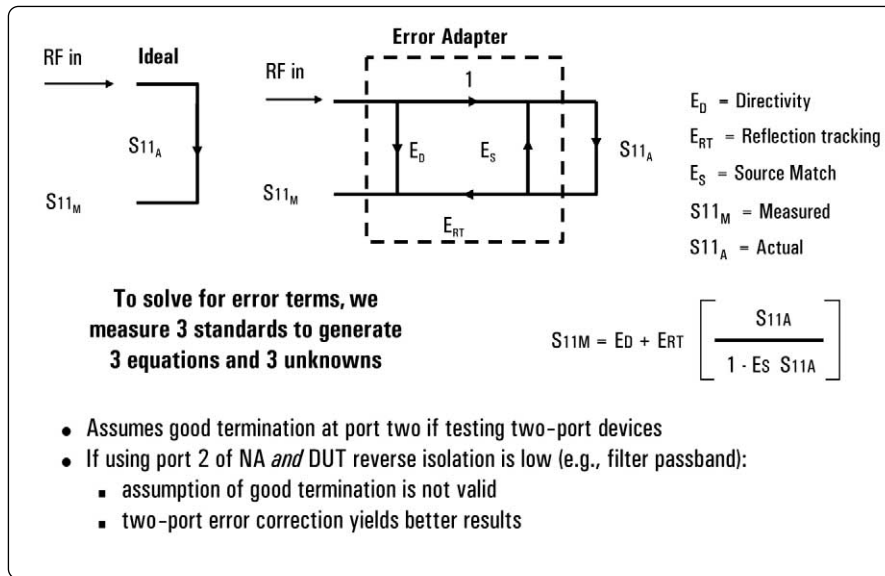
- Process of characterizing systematic error terms
 - measure **known standards**
 - remove effects from subsequent measurements
- **1-port calibration** (*reflection measurements*)
 - only 3 systematic error terms measured
 - directivity, source match, and reflection tracking
- **Full 2-port calibration** (*reflection and transmission measurements*)
 - 12 systematic error terms measured
 - usually requires 12 measurements on four known standards (SOLT)
- Standards defined in **cal kit definition** file
 - network analyzer contains standard cal kit definitions
 - **CAL KIT DEFINITION MUST MATCH ACTUAL CAL KIT USED!**
 - User-built standards must be characterized and entered into user cal-kit



What is Vector-Error Correction?

Vector-error correction is the process of characterizing systematic error terms by measuring known calibration standards, and then removing the effects of these errors from subsequent measurements.

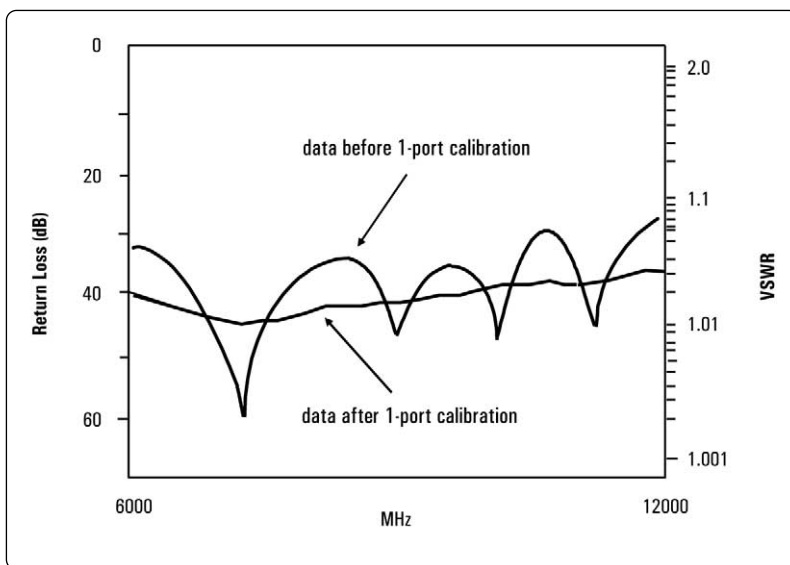
One-port calibration is used for reflection measurements and can measure and remove three systematic error terms (directivity, source match, and reflection tracking). Full two-port calibration can be used for both reflection and transmission measurements, and all twelve systematic error terms are measured and removed. Two-port calibration usually requires twelve measurements on four known standards (short-open-load-through or SOLT). Some standards are measured multiple times (e.g., the through standard is usually measured four times). The standards themselves are defined in a cal-kit definition file, which is stored in the network analyzer. Agilent network analyzers contain all of the cal-kit definitions for our standard calibration kits. In order to make accurate measurements, the cal-kit definition **MUST MATCH THE ACTUAL CALIBRATION KIT USED!** If userbuilt calibration standards are used (during fixtured measurements for example), then the user must characterize the calibration standards and enter the information into a user cal-kit file. Sources of more information about this topic can be found in the appendix.



Reflection: One-Port Model

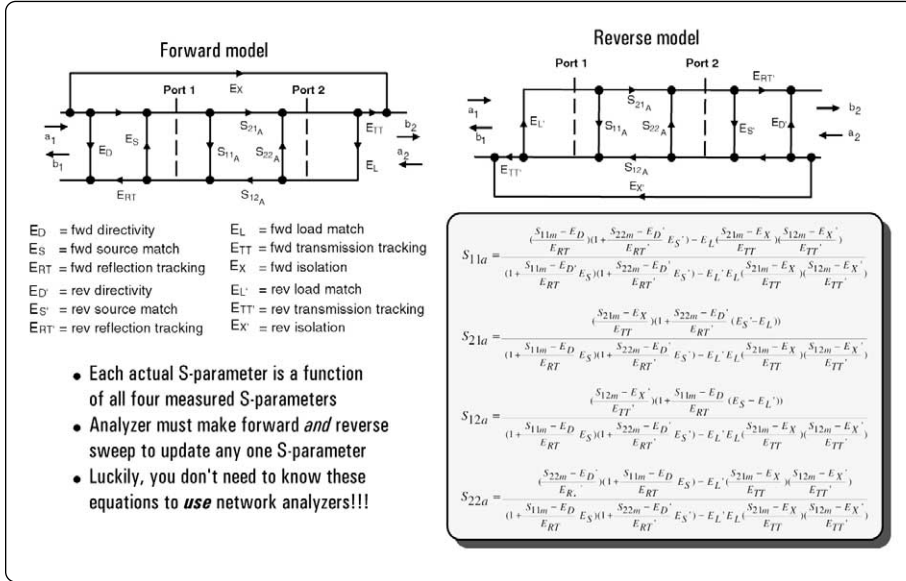
Taking the simplest case of a one-port reflection measurement, we have three systematic errors and one equation to solve in order to calculate the actual reflection coefficient from the measured value. In order to do this, we must first calculate the individual error terms contained in this equation. We do this by creating three more equations with three unknowns each, and solving them simultaneously. The three equations come from measuring three known calibration standards, for example, a short, an open, and a Z_0 load. Solving the equations will yield the systematic error terms and allow us to derive the actual reflection S-parameter of the device from our measurements.

When measuring reflection two-port devices, a one-port calibration assumes a good termination at port two of the device. If this condition is met (by connecting a load calibration standard for example), the one-port calibration is quite accurate. If port two of the device is connected to the network analyzer and the reverse isolation of the DUT is low (for example, filter passbands or cables), the assumption of a good load termination is not valid. In these cases, two-port error correction provides more accurate measurements. An example of a twoport device where load match is not important is an amplifier. The reverse isolation of the amplifier allows oneport calibration to be used effectively. An example of the measurement error that can occur when measuring a two-port filter using a one-port calibration will be shown shortly.



Before and After One-Port Calibration

Shown here is a plot of reflection with and without one-port calibration. Without error correction, we see the classic ripple pattern caused by the systematic errors interfering with the measured signal. The error-corrected trace is much smoother and better represents the device's actual reflection performance.



Two-Port Error Correction

Two-port error correction is the most accurate form of error correction since it accounts for all of the majorsources of systematic error. The error model for a two-port device is shown above. Shown below are the equations to derive the actual device S-parameters from the measured S-parameters, once the systematic error terms have been characterized. Notice that each actual S-parameter is a function of all four measured S-parameters. The network analyzer must make a forward and reverse sweep to update any one S-parameter. Luckily, you don't need to know these equations to use network analyzers!!!

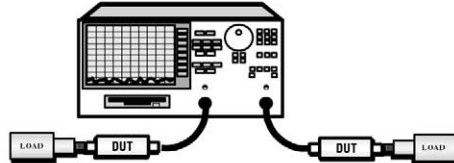
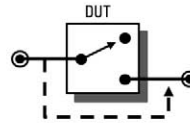
$$S_{11a} = \frac{\left(\frac{S_{11m} - E_D}{E_{RT}} \right) (1 + \frac{S_{22m} - E_{D'}}{E_{RT'}} E_{S'}) - E_L \left(\frac{S_{21m} - E_X}{E_{TT}} \right) \left(\frac{S_{12m} - E_{X'}}{E_{TT'}} \right)}{\left(1 + \frac{S_{11m} - E_D}{E_{RT}} E_{S'} \right) \left(1 + \frac{S_{22m} - E_{D'}}{E_{RT'}} E_{S'} \right) - E_L' E_L \left(\frac{S_{21m} - E_X}{E_{TT}} \right) \left(\frac{S_{12m} - E_{X'}}{E_{TT'}} \right)}$$

$$S_{21a} = \frac{\left(\frac{S_{21m} - E_X}{E_{TT}} \right) \left(1 + \frac{S_{22m} - E_{D'}}{E_{RT'}} (E_{S'} - E_L) \right)}{\left(1 + \frac{S_{11m} - E_D}{E_{RT}} E_{S'} \right) \left(1 + \frac{S_{22m} - E_{D'}}{E_{RT'}} E_{S'} \right) - E_L' E_L \left(\frac{S_{21m} - E_X}{E_{TT}} \right) \left(\frac{S_{12m} - E_{X'}}{E_{TT'}} \right)}$$

$$S_{12a} = \frac{\left(\frac{S_{12m} - E_{X'}}{E_{TT'}} \right) \left(1 + \frac{S_{11m} - E_D}{E_{RT}} (E_S - E_L') \right)}{\left(1 + \frac{S_{11m} - E_D}{E_{RT}} E_{S'} \right) \left(1 + \frac{S_{22m} - E_{D'}}{E_{RT'}} E_{S'} \right) - E_L' E_L \left(\frac{S_{21m} - E_X}{E_{TT}} \right) \left(\frac{S_{12m} - E_{X'}}{E_{TT'}} \right)}$$

$$S_{22a} = \frac{\left(\frac{S_{22m} - E_{D'}}{E_{RT'}} \right) \left(1 + \frac{S_{11m} - E_D}{E_{RT}} E_{S'} \right) - E_L' \left(\frac{S_{21m} - E_X}{E_{TT}} \right) \left(\frac{S_{12m} - E_{X'}}{E_{TT'}} \right)}{\left(1 + \frac{S_{11m} - E_D}{E_{RT}} E_{S'} \right) \left(1 + \frac{S_{22m} - E_{D'}}{E_{RT'}} E_{S'} \right) - E_L' E_L \left(\frac{S_{21m} - E_X}{E_{TT}} \right) \left(\frac{S_{12m} - E_{X'}}{E_{TT'}} \right)}$$

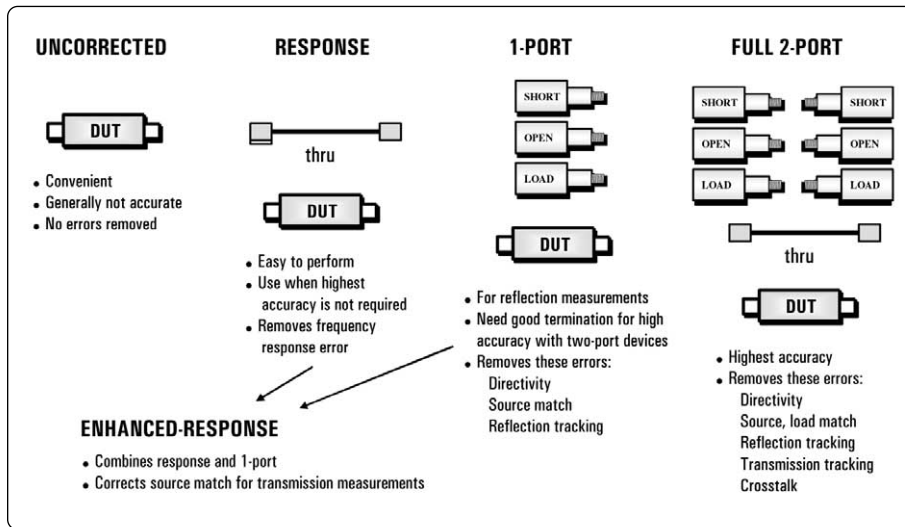
- Can be a problem with:
 - high isolation devices (e.g., switch in open position)
 - high dynamic range devices (some filter stopbands)
- Isolation calibration
 - adds noise to error model (measuring near noise floor of system)
 - only perform if really needed (use averaging if necessary)
 - if crosstalk is **independent** of DUT match, use two terminations
 - if **dependent** on DUT match, use DUT with termination on output



Crosstalk: Signal Leakage Between Test Ports During Transmission

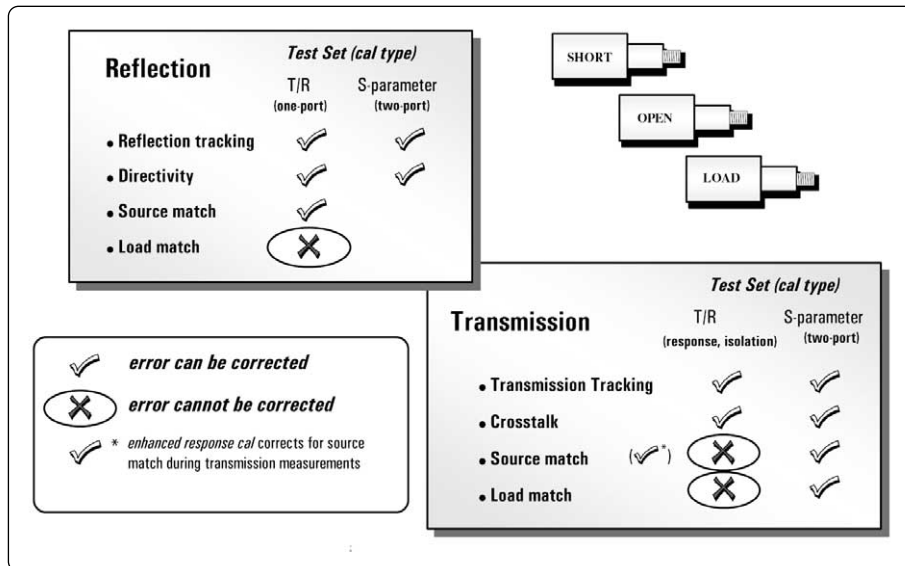
When performing a two-port calibration, the user has the option of omitting the part of the calibration that characterizes crosstalk or isolation. The definition of crosstalk is the signal leakage between test ports when no device is present. Crosstalk can be a problem with high-isolation devices (e.g., switch in open position) and highdynamic range devices (some filter stopbands). The isolation calibration adds noise to the error model since we usually are measuring near the noise floor of the system. For this reason, one should only perform the isolation calibration if it is really needed. If the isolation portion of the calibration is done, trace averaging should be used to ensure that the system crosstalk is not obscured by noise. In some network analyzers, crosstalk can be minimized by using the alternate sweep mode instead of the chop mode (the chop mode makes measurements on both the reflection (A) and transmission (B) receivers at each frequency point, whereas the alternate mode turns off the reflection receiver during the transmission measurement).

The best way to perform an isolation calibration is by placing the devices that will be measured on each test port of the network analyzer, with terminations on the other two device ports. Using this technique, the network analyzer sees the same impedance versus frequency during the isolation calibration as it will during subsequent measurements of the DUT. If this method is impractical (in test fixtures, or if only one DUT is available, for example), then placing a terminated DUT on the source port and a termination on the load port of the network analyzer is the next best alternative (the DUT and termination must be swapped for the reverse measurement). If no DUT is available or if the DUT will be tuned (which will change its port matches), then terminations should be placed on each network analyzer test port for the isolation calibration.



Errors and Calibration Standards

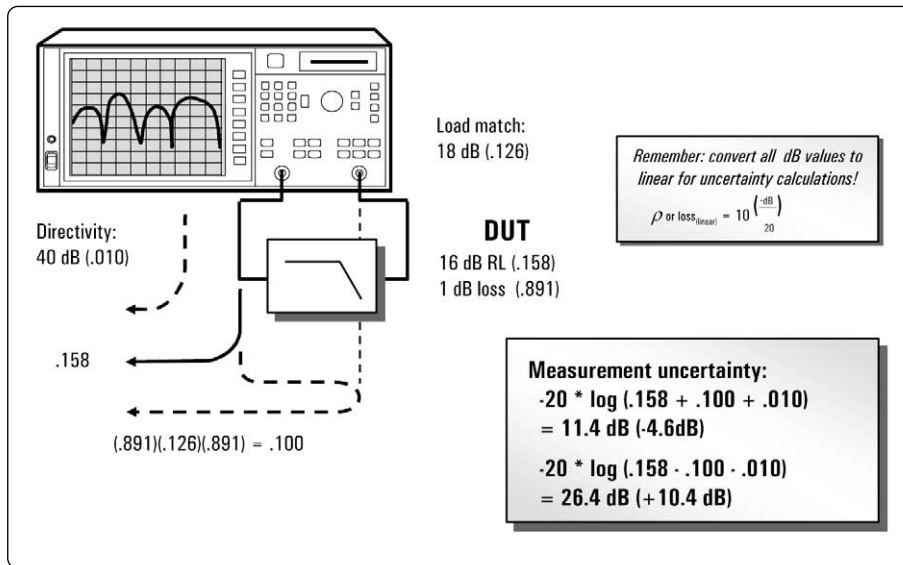
A network analyzer can be used for uncorrected measurements, or with any one of a number of calibration choices, including response calibrations and one- or two-port vector calibrations. A summary of these calibrations is shown above. We will explore the measurement uncertainties associated with the various calibration types in this section.



Calibration Summary

This summary shows which error terms are accounted for when using analyzers with T/R test sets (models ending with ET) and S-parameter test sets (models ending with ES). Notice that load match is the key error term that cannot be removed with a T/R-based network analyzer.

The following examples show how measurement uncertainty can be estimated when measuring two-port devices with a T/R-based network analyzer. We will also show how 2-port error correction provides the least measurement uncertainty.

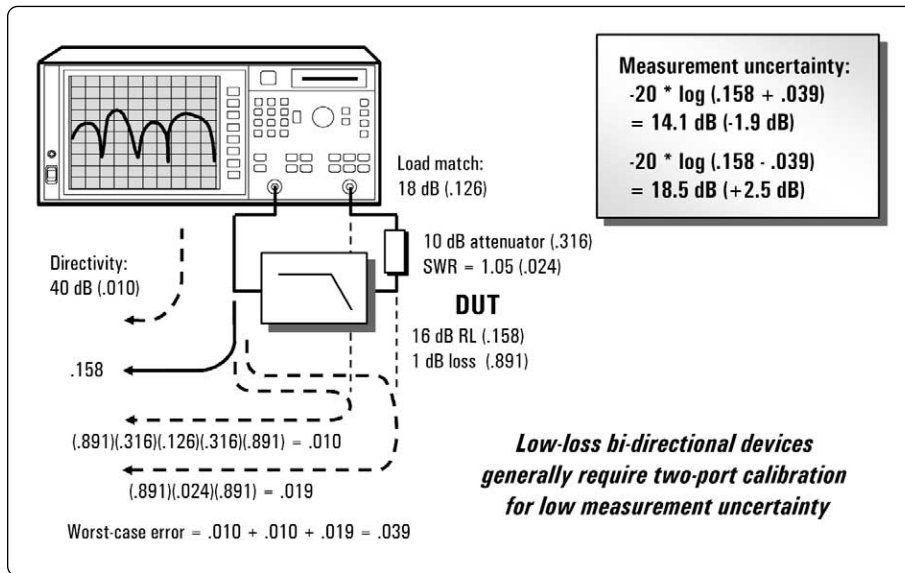


example, it is necessary to convert all dB values into linear values. Next, we must add and subtract the undesired reflection signal resulting from the load match (with a reflection coefficient of 0.100) with the signal reflecting from the DUT (0.158). To be consistent with the next example, we will also include the effect of the directivity error signal (0.010). As a worst case analysis, we will add this signal to the error signal resulting from the load match. The combined error signal is then $0.100 + 0.010 = 0.110$. When we add and subtract this error signal from the desired 0.158, we see the measured return loss of the 16-dB filter may appear to be anywhere from 11.4 dB to 26.4 dB, allowing too much margin for error. In production testing, these errors could easily cause filters which met specification to fail, while filters that actually did not meet specification might pass. In tuning applications, filters could be mistuned as operators try to compensate for the measurement error.

Reflection Example Using a One-Port Cal

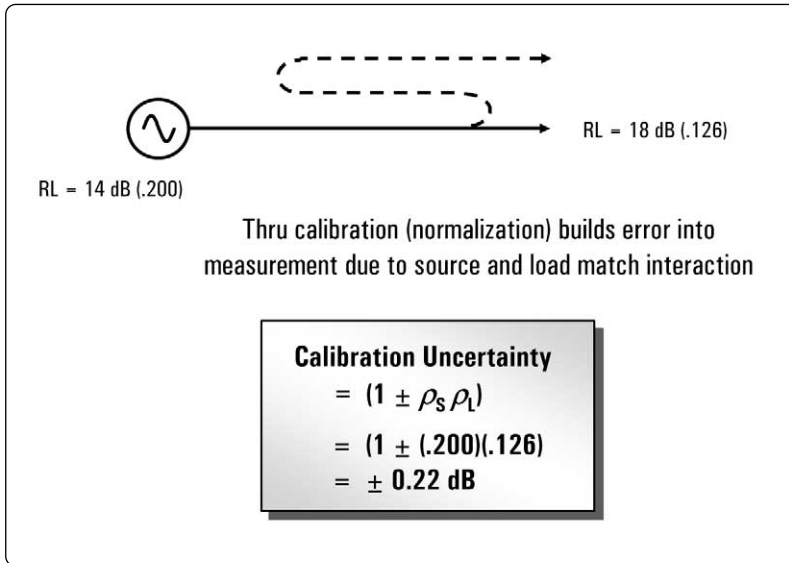
Here is an example of how much measurement uncertainty we might encounter when measuring the input match of a filter after a one-port calibration. In this example, our filter has a return loss of 16 dB, and 1 dB of insertion loss. Let's say the raw (uncorrected) load match of our network analyzer is specified to be 18 dB (generally, typical performance is significantly better than the specified performance). The reflection from the test port connected to the filter's output is attenuated by twice the filter loss, which is only 2 dB total in this case. This value is not adequate to sufficiently suppress the effects of this error signal, which illustrates why low-loss devices are difficult to measure accurately. To determine the measurement uncertainty of this

When measuring an amplifier with good isolation between output and input (i.e., where the isolation is much greater than the gain), there is much less measurement uncertainty. This is because the reflection caused by the load match is severely attenuated by the product of the amplifier's isolation and gain. To improve measurement uncertainty for a filter, the output of the filter must be disconnected from the analyzer and terminated with a high-quality load, or a high-quality attenuator can be inserted between the filter and port two of the network analyzer. Both techniques improve the analyzer's effective load match.



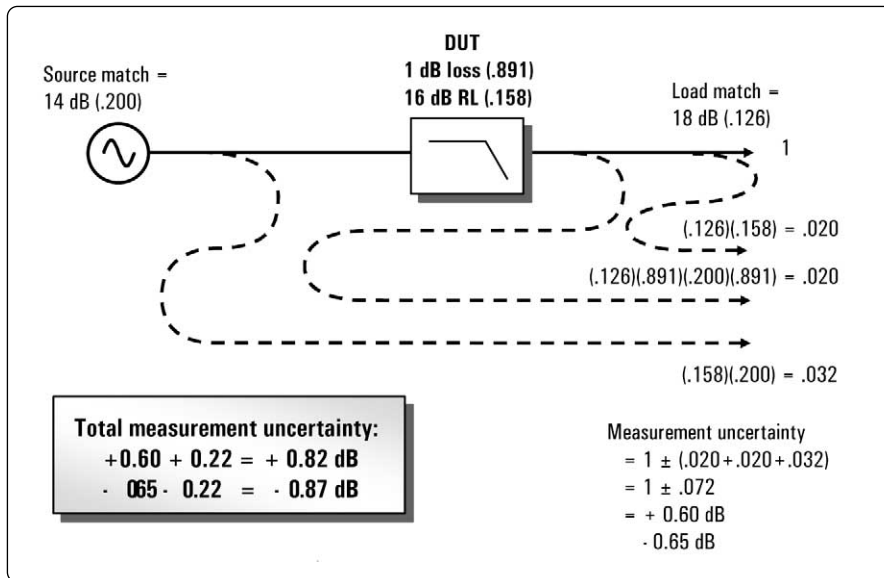
Using a One-Port + Attenuator

Let's see how much improvement we get by adding an attenuator between the output of the filter and our network analyzer. If we inserted a perfect 10 dB attenuator between port two of the network analyzer and the filter used in the previous example, we would expect the effective load match of our test system to improve by twice the value of the attenuator (since the error signal travels through the attenuator twice), which in this example, would be 20 dB. However, we must take into account the reflection introduced by the attenuator itself. For this example, we will assume the attenuator has a SWR of 1.05. Now, our effective load match is only 28.6 dB ($-20 * \log[10 \exp(-32.3/20) + 10 \exp(-38/20)]$), which is only about a 10 dB improvement. This value is the combination of a 32.3 dB match from the attenuator (SWR = 1.05) and the 38 dB effective match of the network analyzer with the 10 dB attenuator. Our worst-case uncertainty is now reduced to +2.5 dB, -1.9 dB, instead of the +10.4 dB, -4.6 dB we had without the 10 dB attenuator. While not as good as what could be achieved with two-port calibration, this level of accuracy may be sufficient for manufacturing applications. To minimize measurement uncertainty, it is important to use the best quality (lowest reflection) attenuators that your budget will allow.



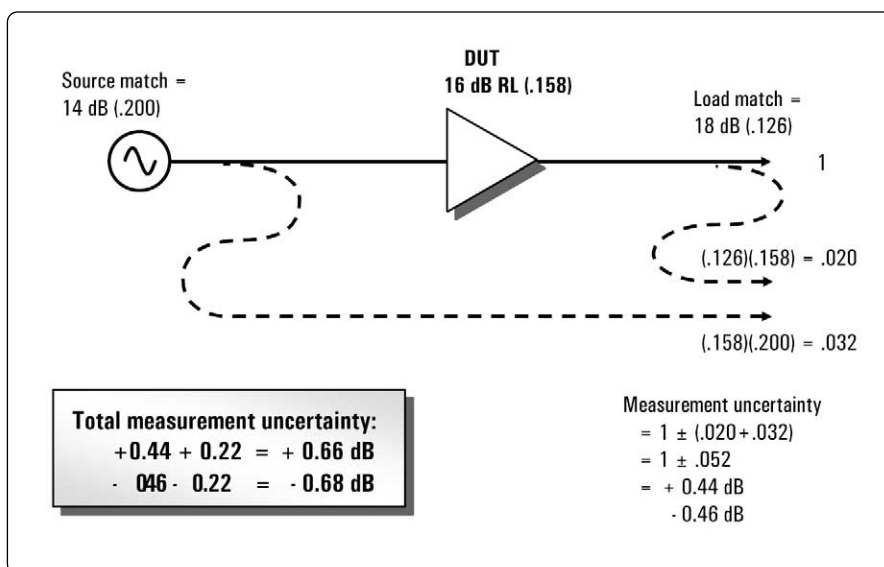
Transmission Example Using Response Cal

Let's do an example transmission measurement using only response calibration. Response calibrations offer simplicity, but with some compromise in measurement accuracy. In making a filter transmission measurement using only response calibration, the first step is to make a through connection between the two test port cables (with no DUT in place). For this example, some typical test port specifications will be used. The ripple caused by this amount of mismatch is calculated as ± 0.22 dB, and is now present in the reference data. Since we don't know the relative phase of this error signal once it passes through the DUT, it must be added to the uncertainty when the DUT is measured (see next slide) in order to compute the worst-case overall measurement uncertainty.



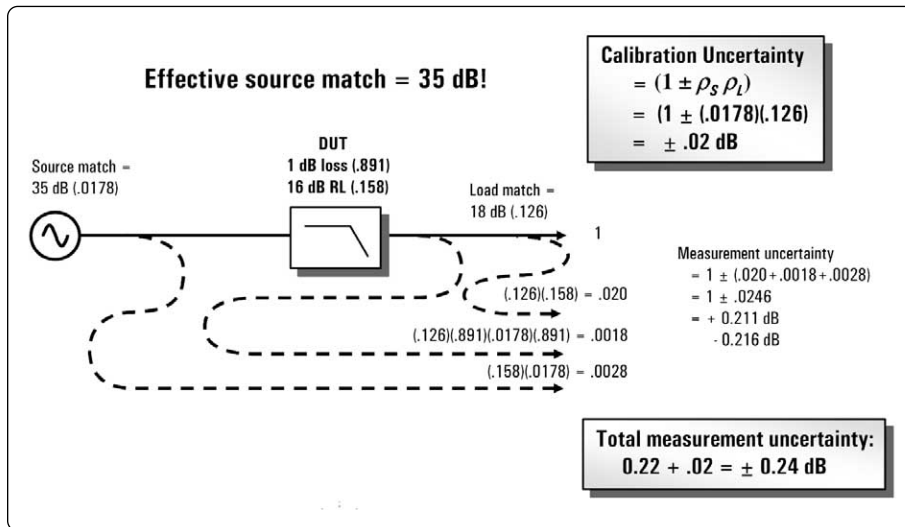
Filter Measurement with Response Cal

Now let's look at the measurement uncertainty when the DUT is inserted. We will use the same loss and mismatch specifications for the DUT and analyzer as before. We have three main error signals due to reflections between the ports of the analyzer and the DUT. Higher-order reflections are present as well, but they don't add any significant error since they are small compared to the three main terms. In this example, we will normalize the error terms to the desired signal that passes through the DUT once. The desired signal is therefore represented as 1, and error terms only show the additional transmission loss due to traveling more than once through the DUT. One of the signals passes through the DUT two extra times, so it is attenuated by twice the loss of the DUT. The worst case is when all of the reflected error signals add together in-phase $(.020 + .020 + .032 = .072)$. In that case, we get a measurement uncertainty of $+0.60 \text{ dB}$, -0.65 dB . The total measurement uncertainty, which must include the 0.22 dB of error incorporated into our calibration measurement, is about $\pm 0.85 \text{ dB}$.



Measuring Amplifiers with a Response Cal

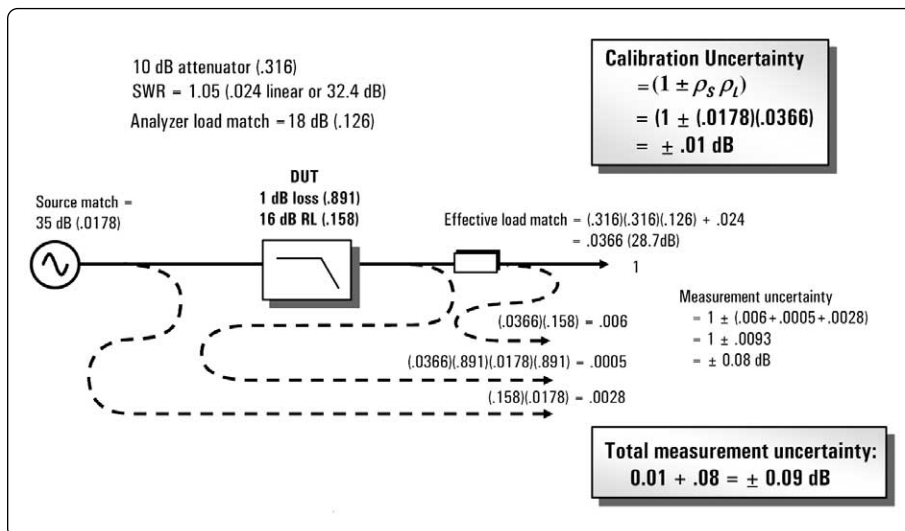
Now let's look at an example of measuring an amplifier that has port matches of 16 dB . The match of our test ports remains the same as our previous transmission response example. We see that the middle error term is no longer present, due to the reverse isolation of the amplifier. This fact has reduced our measurement uncertainty to about $\pm 0.45 \text{ dB}$. Our total measurement error now has been reduced to about $\pm 0.67 \text{ dB}$, versus the $\pm 0.85 \text{ dB}$ we had when measuring the filter.



Filter Measurements Using the Enhanced Response Calibration

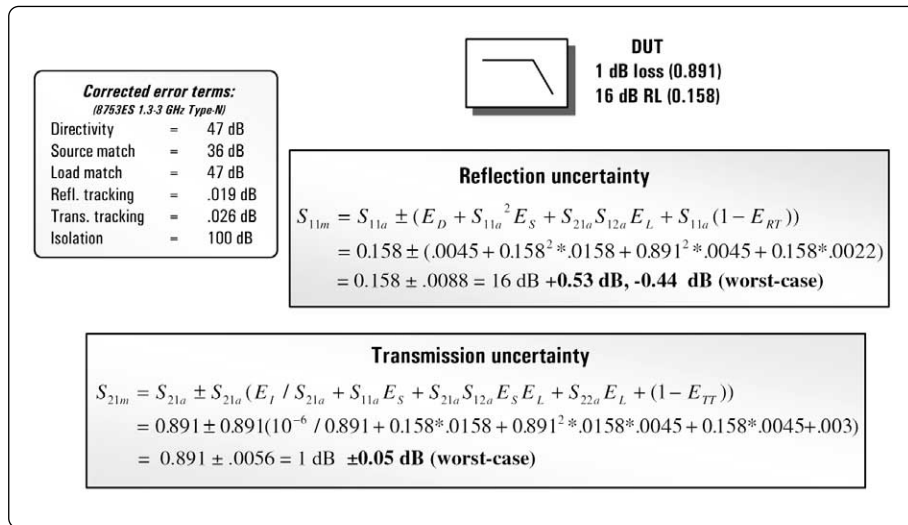
A feature contained in many of Agilent's T/R-based network analyzers is the enhanced response calibration. This calibration greatly reduces all the error terms involving a reflection from the source match. It requires the measurement of short, open, load, and through standards for transmission measurements. Essentially, it combines a one-port cal and a response cal to correct source match during transmission measurements. Recall that a standard response calibration cannot correct for the source and load match error terms.

Continuing with our filter example, we see the enhanced response calibration has improved the effective source match during transmission measurements to around 35 dB, instead of the 14 dB we used previously. This greatly reduces the calibration error ($\pm 0.02 \text{ dB}$ instead of $\pm 0.22 \text{ dB}$), as well as the two measurement error terms that involve interaction with the effective source match. Our total measurement error is now $\pm 0.24 \text{ dB}$, instead of the previously calculated $\pm 0.85 \text{ dB}$.



Using the Enhanced Response Calibration Plus an Attenuator

We can further improve transmission measurements by using the enhanced response calibration and by inserting a high-quality attenuator between the output port of the device and test port two of the network analyzer. In this example, we will use a 10 dB attenuator with a SWR of 1.05 (as we did with the reflection example). This makes the effective load match of the analyzer 28.7 dB, about a 10 dB improvement. Our calibration error is minuscule ($\pm 0.01 \text{ dB}$), and our total measurement uncertainty has been reduced to $\pm 0.09 \text{ dB}$. This is very close to what can be achieved with two-port error correction. As we have seen, adding a high-quality attenuator to port two of a T/R network analyzer can significantly improve measurement accuracy, with only a modest loss in dynamic range.



Calibrating Measurement Uncertainty After a Two-Port Calibration

Here is an example of calculating measurement error after a two-port calibration has been done. Agilent provides values on network analyzer data sheets for effective directivity, source and load match, tracking, and isolation, usually for several different calibration kits. The errors when measuring our example filter have been greatly reduced (± 0.5 dB reflection error, ± 0.05 dB transmission error). Phase errors would be similarly small.

Note that this is a worst-case analysis, since we assume that all of the errors would add in-phase. For many narrowband measurements, the error terms will not all align with one another. A less conservative approach to calculating measurement uncertainty would be to use a root-sum-squares (RSS) method. The best technique for estimating measurement uncertainty is to use a statistical approach (which requires knowing or estimating the probability-distribution function of the error terms) and calculating the $\pm 3\sigma$ (sigma) limits.

The terms used in the equations are forward terms only and are defined as:

- E_D = directivity error
- E_S = source match
- E_L = load match
- E_{RT} = reflection tracking
- E_{TT} = transmission tracking
- E_I = crosstalk (transmission isolation)
- a = actual
- m = measured

Reflection

Calibration type	Measurement uncertainty
One-port	-4.6/ 10.4 dB
One-port + attenuator	-1.9/ 2.5 dB
Two-port	-0.44/ 0.53 dB

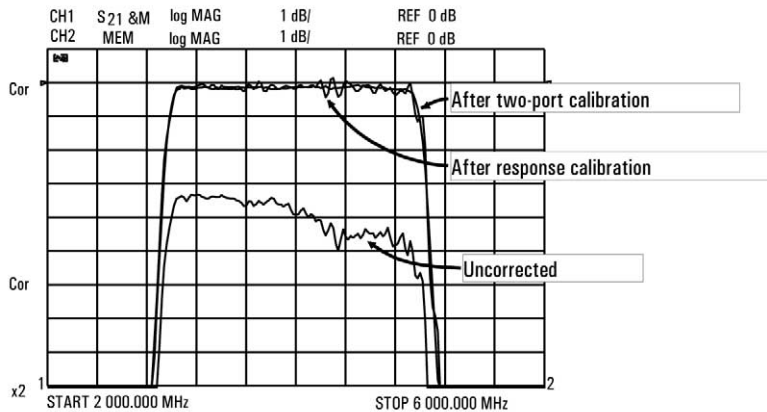
Transmission

Calibration type	Calibration uncertainty	Measurement uncertainty	Total uncertainty
Response	± 0.22 dB	0.60/ -0.65 dB	0.82/ -0.87 dB
Enhanced response	± 0.02 dB	± 0.22 dB	± 0.24 dB
Enh. response + attenuator	± 0.01 dB	± 0.08 dB	± 0.09 dB
Two port	----		± 0.05 dB

Comparison of Measurement Examples

Here is a summary of the measurement uncertainties we have discussed so far for different types of calibration.

Measuring filter insertion loss



Response Versus Two-Port Calibration

Let's look at some actual measurements done on a bandpass filter with different levels of error correction. The uncorrected trace shows considerable loss and ripple. In fact, the passband response varies about ± 1 dB around the filter's center frequency. Is the filter really this bad? No. What we are actually measuring is the sum of the filter's response and that of our test system.

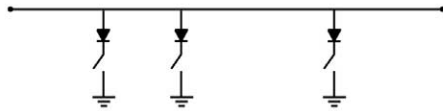
Performing a normalization prior to the measurement of the filter removes the frequency response of the system (transmission tracking error) from the measurement. The loss that was removed was most likely caused by the test cables. After normalization, the frequency response of the filter still contains ripple caused by an interaction between the system's source and load match. This ripple even goes above the 0 dB reference line, indicating gain! However, we know that a passive device cannot amplify signals. This apparent anomaly is due to mismatch error.

The measurement shown after a two-port calibration is the most accurate of the three measurements shown. Using vector-error correction, the filter's passband response shows variation of about ± 0.1 dB around its center frequency. This increased level of measurement flatness will ensure minimum amplitude distortion, increase confidence in the filter's design, and ultimately increase manufacturing yields due to lower test-failure rates.

- Variety of modules cover 300 kHz to 26.5 GHz
- 2 and 4-port versions available
- Choose from six connector types (50 Ω and 75 Ω)
- Mix and match connectors (3.5mm, Type-N, 7/16)
- Single-connection
 - reduces calibration time
 - makes calibrations easy to perform
 - minimizes wear on cables and standards
 - eliminates operator errors
- Highly repeatable temperature-compensated terminations provide excellent accuracy



Microwave modules use a transmission line shunted by PIN-diode switches in various combinations



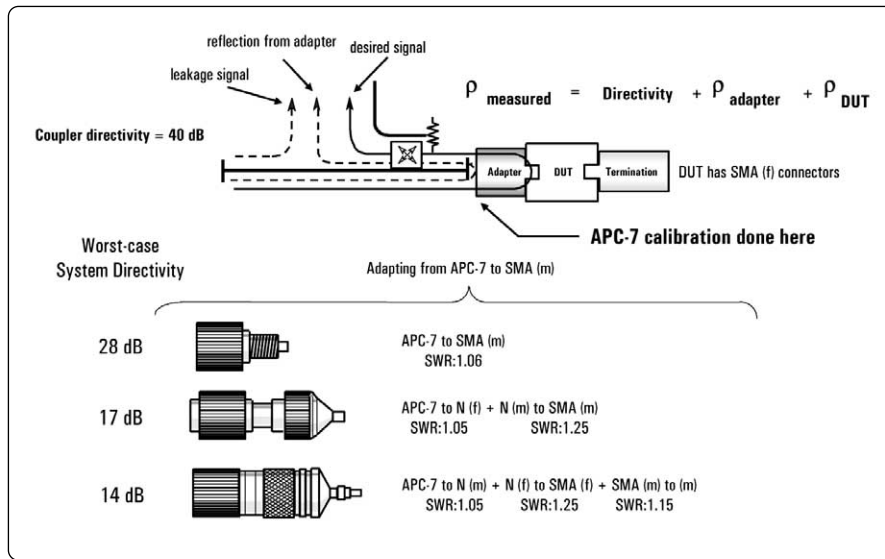
ECal: Electronic Calibration (85060/90 Series)

Although the previous slides have all shown mechanical calibration standards, Agilent offers a solid-state calibration solution which makes two, three, and four-port calibration fast, easy, and less prone to operator errors. A variety of calibration modules are available with different connector types and frequency ranges. You can configure a single module with different connector types or choose all the same type. The calibration modules are solid-state devices with programmable, repeatable impedance states. These states are characterized at the Agilent factory using a network analyzer calibrated with coaxial, airline-TRL standards (the best calibration available), making the ECal modules transfer standards (rather than direct standards).

For the microwave calibration modules, the various impedance states are achieved by PIN-diode switches which shunt the transmission line to ground. The number of diodes and their location vary depending upon the module's frequency range. A multitude of reflection coefficients can be generated by applying various combinations of the shunts. With no shunts, the network acts as a low loss transmission line. High isolation between the ports is obtained by driving several of the PIN shunts simultaneously. Four different states are used to compute the error terms at each frequency point. Four states are used because this gives the best trade-off between high accuracy and the time required for the calibration. With four reflection states, we have four equations but only three unknowns. To achieve the best accuracy from this over-determined set of equations, a least-squares-fit algorithm is used. Adding more impedance states at each frequency point would further improve accuracy but at the expense of more calibration time.

The RF module uses the more traditional short, open, and load terminations, and a through transmission line.

For more information about these products, please order Agilent literature number 5963-3743E.



Adapter Considerations

Whenever possible, reflection calibrations should be done with a cal kit that matches the connector type of the DUT. If adapters need to be used to mate the calibrated test system to the DUT, the effect of these adapters on measurement accuracy can be very large. This error is often ignored, which may or may not be acceptable. As the slide shows, the adapter causes an error signal which can add or subtract with the desired reflection signal from the DUT. Worst-case effective directivity (in dB) is now:

$$-20 \log (\text{Corrected-coupler-directivity} + \rho_{\text{adapters}})$$

If the adapter has a SWR of say 1.5 (the less-expensive variety), the effective directivity of the coupler drops to around 14 dB worst case, even if the coupler itself had infinite directivity! In other words, with a perfect Z_0 load ($\rho_L = 0$) on the output of the adapter; the reflected signal appearing at the coupled port would only be 14 dB less than the reflection from a short or open circuit. Stacking adapters compounds the problem, as is illustrated above. Consequently, it is very important to use quality adapters (or preferably, no adapters at all) in your measurement system, so system directivity is not excessively degraded. While error-correction can mitigate the effect of adapters on the test port, our system is more susceptible to drift with degraded raw (uncorrected) directivity.

When doing a through cal, normally test ports mate directly

- cables can be connected directly without an adapter
- result is a zero length through



What is an insertable device?

- has same type of connector, but different sex on each port
- has same type of sexless connector on each port (e.g. APC 7)

What is a non insertable device?

- one that cannot be inserted in place of a zero length through
- has same connectors on each port (type and sex)
- has different type of connector on each port
(e.g., waveguide on one port, coaxial on the other)



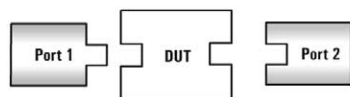
What calibration choices do I have for non insertable devices?

- use an *uncharacterized* through adapter
- use a *characterized* through adapter (modify cal kit definition)
- swap equal adapters
- adapter removal

Calibrating Non-Insertable Devices

When performing a through calibration, often the test ports mate directly. For example, two cables with the appropriate connectors can be joined without a through adapter, resulting in a zero-length through path. An insertable device is one that can be substituted for a zero-length through. This device has the same connector type on each port but of the opposite sex, or the same sexless connector on each port, either of which makes connection to the test ports quite simple. A noninsertable device is one that cannot be substituted for a zero-length through. It has the same type and sex connectors on each port or a different type of connector on each port, such as 7/16 at one end and SMA on the other end.

There are several calibration choices available for noninsertable devices. One choice is to use an uncharacterized through adapter. While not recommended, this might be acceptable at low frequencies where the electrical length of the adapter is relatively small. In general, it is preferable to use a characterized through adapter (where the electrical length and loss are specified), which requires modifying the calibration-kit definition. A high-quality through adapter (with good match) should be used since reflections from the adapter cannot be removed. The other two choices (swapping equal adapters and adapter removal) will be discussed next.



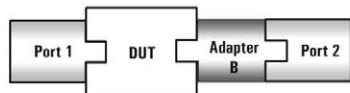
Accuracy depends on how well the adapters are matched - loss, electrical length, match and impedance should all be equal



1. Transmission cal using adapter A.



2. Reflection cal using adapter B.



3. Measure DUT using adapter B.

Swap Equal Adapters Method

The swap-equal-adapters method is very useful for devices with the same connector type and sex (female SMA on both ends for example). It requires the use of two precision matched adapters that are equal in performance but have connectors of different sexes. For example, for measuring a device with female SMA connectors on both ends using APC-7 mm test cables, the adapters could be 7-mm-to-male-3.5-mm and 7-mm-to-female-3.5-mm. To be equal, the adapters must have the same match, characteristic impedance, insertion loss, and electrical delay. Many of Agilent's calibration kits include matched adapters for this purpose.

The first step in the swap-equal-adapters method is to perform the transmission portion of a two-port calibration with the adapter needed to make the through connection. This adapter is then removed and the second adapter is used in its place during the reflection portion of the calibration, which is performed on both test ports. This swap changes the sex of one of the test ports so that the DUT can be inserted and measured (with the second adapter still in place) after the calibration procedure is finished. The errors remaining after calibration are equal to the difference between the two adapters. The technique provides a high level of accuracy, but not as high as the more complicated adapter-removal technique.

- Calibration is very accurate and traceable
- In firmware of 8753, 8720 and 8510 series
- Also accomplished with ECal modules (85060/90)
- Uses adapter with same connectors as DUT
- Must specify electrical length of adapter to within 1/4 wavelength of highest frequency (to avoid phase ambiguity)



1. Perform 2-port cal with adapter on port 2.
Save in cal set 1.



2. Perform 2-port cal with adapter on port 1.
Save in cal set 2.

[CAL] [MORE] [MODIFY CAL SET]
[ADAPTER REMOVAL]

3. Use ADAPTER REMOVAL
to generate new cal set.



4. Measure DUT without cal adapter.

Adapter Removal Calibration

Adapter-removal calibration provides the most complete and accurate calibration procedure for noninsertable devices. It is available in the 8753, 8720, and 8510 series of network analyzers. This method uses a through adapter that has the same connectors as the noninsertable DUT (this adapter is sometimes referred to as the calibration adapter). The electrical length of the adapter must be specified within one-quarter wavelength at each calibration frequency. Type N, 3.5-mm, and 2.4-mm calibration kits for the 8510 contain adapters specified for this purpose. For other adapters, the user can simply enter the electrical length.

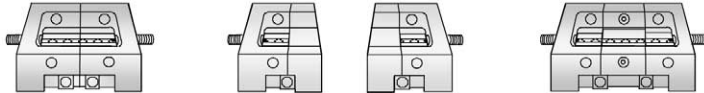
Two full two-port calibrations are needed for an adapter-removal calibration. In the first calibration, the through adapter is placed on test port two, and the results are saved into a calibration set. In the second calibration, the adapter is moved to test port one and the resulting data is saved into a second calibration set. Two different calibration kits may be used during this process to accommodate devices with different connector types. To complete the adapter-removal calibration, the network analyzer uses the two sets of calibration data to generate a new set of error coefficients that completely eliminate the effects of the calibration adapter. At this point, the adapter can be removed and measurements can be made directly on the DUT.

We know about Short-Open-Load-Thru (SOLT) calibration...

What is TRL?

- A two-port calibration technique
- Good for noncoaxial environments (waveguide, fixtures, wafer probing)
- Uses the same 12-term error model as the more common SOLT cal
- Uses practical calibration standards that are easily fabricated and characterized
- Two variations: TRL (requires 4 receivers) and TRL* (only three receivers needed)
- Other variations: Line-Reflect-Match (LRM), Thru-Reflect-Match (TRM), plus many others

TRL was developed for non-coaxial microwave measurements



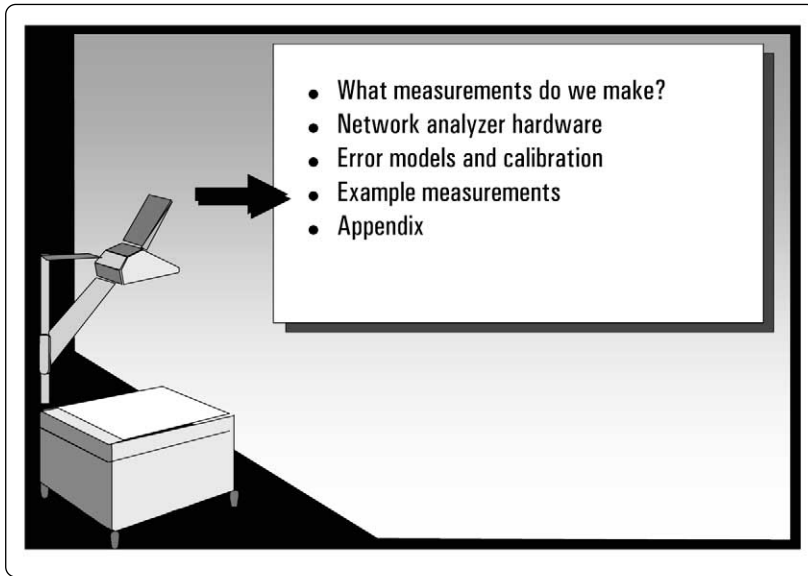
Thru-Reflect-Line (TRL) Calibration

When performing a two-port calibration, we have some choices based on the type of calibration standards we want to use. So far, we have only discussed coaxial calibration techniques. Let's briefly look at TRL (throughreflect-line), a calibration technique that is especially useful for microwave, noncoaxial environments such as fixture, wafer probing, or waveguide. It is the second-most common type of two-port calibration, after SOLT. TRL solves for the same 12 error terms as the more common SOLT calibration, but uses a slightly different error model.

The main advantage of TRL is that the calibration standards are relatively easy to make and define at microwave frequencies. This is a big benefit since it is difficult to build good, noncoaxial, open and load standards at microwave frequencies. TRL uses a transmission line of known length and impedance as one standard. The only restriction is that the line needs to be significantly longer in electrical length than the through line, which typically is of zero length. TRL also requires a high-reflection standard (usually, a short or open) whose impedance does not have to be well characterized, but it must be electrically the same for both test ports.

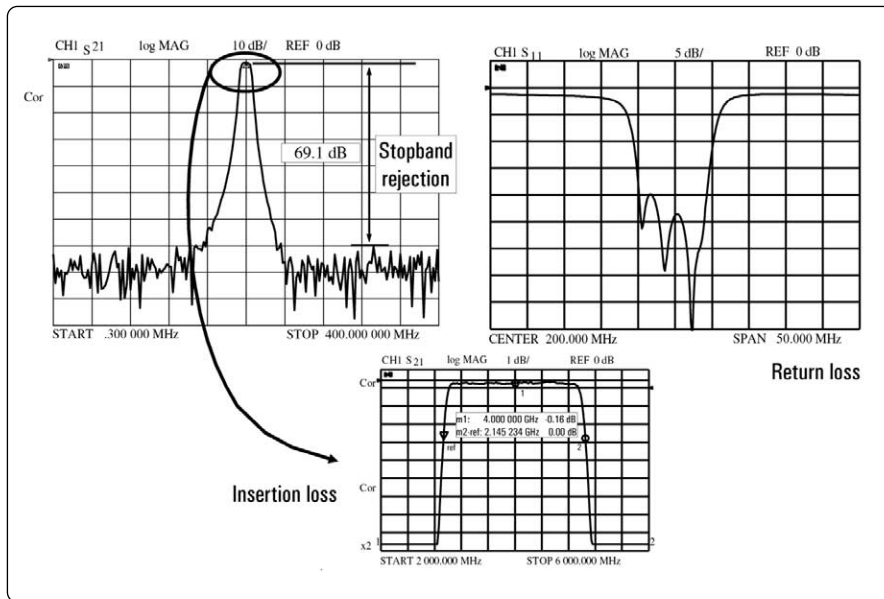
For RF applications, the lengths of the transmission lines needed to cover down to low frequencies become impractical (too long). It is also difficult to make good TRL standards on printed-circuit boards, due to dielectric, line-dimension, and board-thickness variations. And, the typical TRL fixture tends to be more complicated and expensive, due to the need to accommodate throughs of two different physical lengths.

There are two variations of TRL. True TRL calibration requires a 4-receiver network analyzer. The version for three-receiver analyzers is called TRL* ("TRL-star"). Other variations of this type of calibration (that share a common error model) are Line-Reflect-Line (LRL), Line-Reflect-Match (LRM), Thru-Reflect-Match (TRM), plus many others.



Agenda

This section will cover some typical measurements. We will look at swept-frequency testing of filters and swept-power testing of amplifiers.



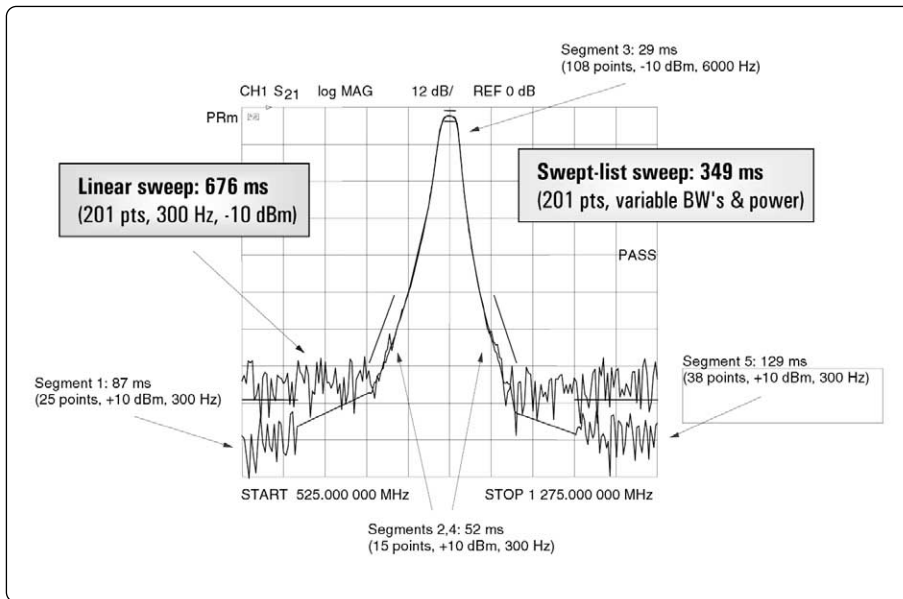
Frequency Sweep – Filter Test

Shown above are the frequency responses of a filter. On the left and bottom we see the transmission response in log magnitude format, and on the right we see the reflection response (return loss).

The most commonly measured filter characteristics are insertion loss and bandwidth, shown on the lower plot with an expanded vertical scale. Another common parameter we might measure is out-of-band rejection. This is a measure of how well a filter passes signals within its bandwidth while simultaneously rejecting all other signals outside of that same bandwidth. The ability of a test system to measure out-of-band rejection is directly dependent on its system dynamic-range specification.

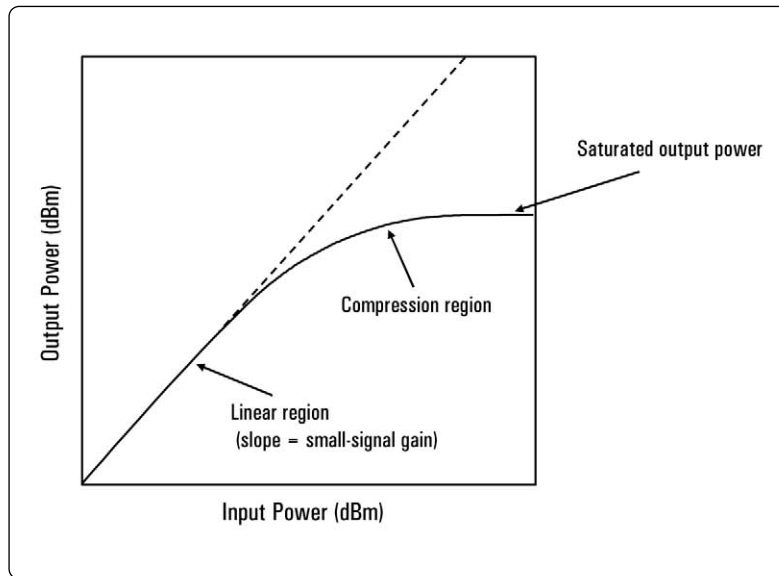
The return loss plot is very typical, showing high reflection (near 0 dB) in the stopbands, and reasonable match in the passband. Most passive filters work in this manner. A special class of filters exist that are absorptive in both the passband and stopband. These filters exhibit a good match over a broad frequency range.

For very narrowband devices, such as crystal filters, the network analyzer must sweep slow enough to allow the filter to respond properly. If the default sweep speed is too fast for the device, significant measurement errors can occur. This can also happen with devices that are electrically very long. The large time delay of the device can result in the receiver being tuned to frequencies that are higher than those coming out of the device, which also can cause significant measurement errors.



Optimize Filter Measurements With Swept-List Mode

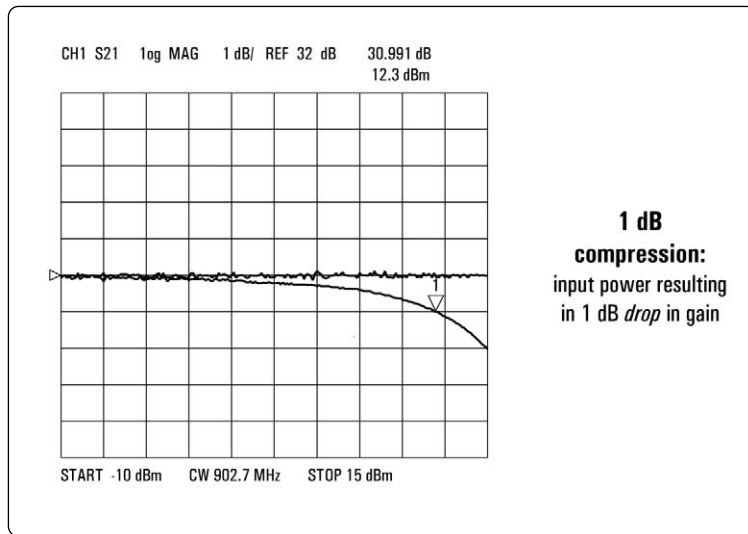
Many network analyzers have the ability to define a sweep consisting of several individual segments (called swept-list mode in the 8753 and 8720 series, and segment sweep in the PNA Series). These segments can have their own stop and start frequency, number of data points, IF bandwidth, and power level. Using a segmented sweep, the sweep can be optimized for speed and dynamic range. Data resolution can be made high where needed (more data points) and low where not needed (less data points); frequency ranges can be skipped where data is not needed at all; the IF bandwidth can be large when high dynamic range is not necessary (in filter passbands, for example), which decreases the sweep time, and small when high dynamic range is required (in filter stopbands, for example); the power level can be decreased in the passband and increased in the stopband for DUTs that contain a filter followed by an amplifier (for example, a cellular-telephone base-station receive filter/LNA combination). The slide shows an example of a filter/amplifier combination where the sweep time and dynamic range using a segmented sweep are considerably better compared to using a linear sweep, where the IF bandwidth and power level are fixed.



Power Sweeps – Compression

Many network analyzers have the ability to do power sweeps as well as frequency sweeps. Power sweeps help characterize the nonlinear performance of an amplifier. Shown above is a plot of an amplifier's output power versus input power at a single frequency. Amplifier gain at any particular power level is the slope of this curve. Notice that the amplifier has a linear region of operation where gain is constant and independent of power level. The gain in this region is commonly referred to as "small-signal gain". At some point as the input power is increased, the amplifier gain appears to decrease, and the amplifier is said to be in compression. Under this nonlinear condition, the amplifier output is no longer sinusoidal – some of the output power is present in harmonics, rather than occurring only at the fundamental frequency. As input power is increased even more, the amplifier becomes saturated, and output power remains constant. At this point, the amplifier gain is essentially zero, since further increases in input power result in no change in output power. In some cases (such as with TWT amplifiers), output power actually decreases with further increases in input power after saturation, which means the amplifier has negative gain.

Saturated output power can be read directly from the above plot. In order to measure the saturated output power of an amplifier, the network analyzer must be able to provide a power sweep with sufficient output power to drive the amplifier from its linear region into saturation. A preamp at the input of the amplifier under test may be necessary to achieve this.



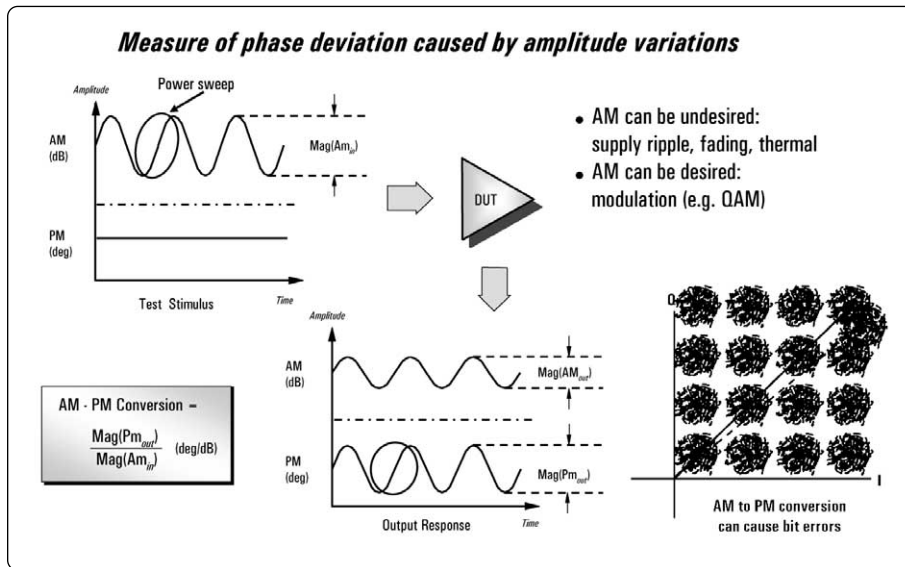
Power Sweep – Gain Compression

The most common measurement of amplifier compression is the 1-dB-compression point, defined here as the input power¹ which results in a 1-dB decrease in amplifier gain (referenced to the amplifier's small-signal gain). The easiest way to measure the 1-dB-compression point is to directly display normalized gain (B/R) from a power sweep. The flat part of the trace is the linear, small-signal region, and the curved part on the right side corresponds to compression caused by higher input power. As shown above, the 1-dB-compression point of the amplifier-under-test is 12.3 dBm, at a CW frequency of 902.7 MHz.

It is often helpful to also know the output power corresponding to the 1-dB-compression point. Using the dualchannel feature found on most modern network analyzers, absolute power and normalized gain can be displayed simultaneously. Display markers can read out both the output power and the input power where 1-dB-compression occurs. Alternatively, the gain of the amplifier at the 1-dB-compression point can simply be added to the 1-dB-compression power to compute the corresponding output power. As seen above, the output power at the 1-dB-compression point is $12.3 \text{ dBm} + 31.0 \text{ dB} = 43.3 \text{ dBm}$.

It should be noted that the power-sweep range needs to be large enough to ensure that the amplifier under test is driven from its linear region into compression. Modern network analyzers typically provide power sweeps with 15 to 25 dB of range, which is more than adequate for most amplifiers. It is also very important to sufficiently attenuate the output of high-power amplifiers to prevent damage to the network analyzer's receiver.

1. The 1-dB-compression point is sometimes defined as the output power resulting in a 1-dB decrease in amplifier gain (as opposed to the input power).



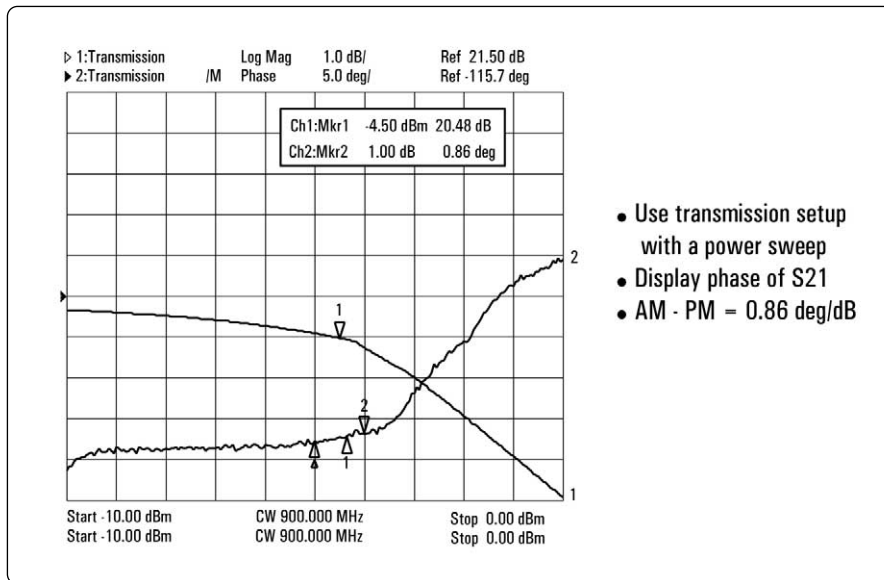
AM to PM Conversion

Another common measurement which helps characterize the nonlinear behavior of amplifiers is AM-to-PM conversion, which is a measure of the amount of undesired phase deviation (the PM) which is induced by amplitude variations inherent in the system (the AM). In a communications system, this unwanted PM is caused by unintentional amplitude variations such as power supply ripple, thermal drift, or multipath fading, or by intentional amplitude change that is a result of the type of modulation used, such as the case with QAM or burst modulation.

AM-to-PM conversion is a particularly critical parameter in systems where phase (angular) modulation is employed, because undesired phase distortion causes analog signal degradation, or increased bit-error rates (BER) in digital systems. Examples of common modulation types that use phase modulation are FM, QPSK, and 16QAM. While it is easy to measure the BER of a digital communication system, this measurement alone does not provide any insight into the underlying phenomena which cause bit errors. AM-to-PM conversion is one of the fundamental contributors to BER, and therefore it is important to quantify this parameter in communication systems.

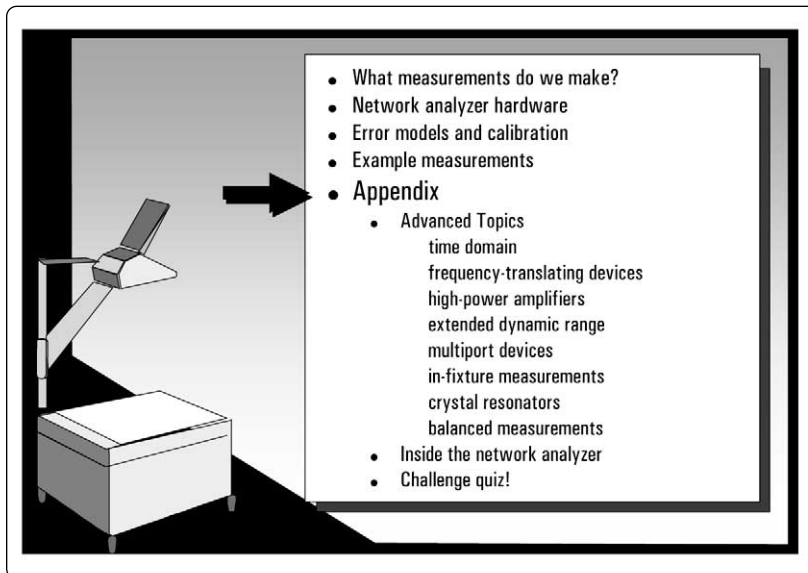
The I/Q diagram shown above shows how AM-to-PM conversion can cause bit errors. Let's say the desirable state change is from the small solid vector to the large solid vector. With AM-PM conversion, the large vector may actually end up as shown with the dotted line. This is due to phase shift that results from a change in power level. For a 64QAM signal as shown (only one quadrant is drawn), we see that the noise circles that surround each state would actually overlap, which means that statistically, some bit errors would occur.

AM-to-PM conversion is usually defined as the change in output phase for a 1-dB increment in the input power to an amplifier, expressed in degrees-per-dB ($^{\circ}/\text{dB}$). An ideal amplifier would have no interaction between its phase response and the level of the input signal. To measure AM to PM conversion, we can use the power sweep capability of the network analyzer. As the slide shows, a power sweep is equivalent to a quarter cycle of sinusoidal modulation.



Measuring AM to PM Conversion

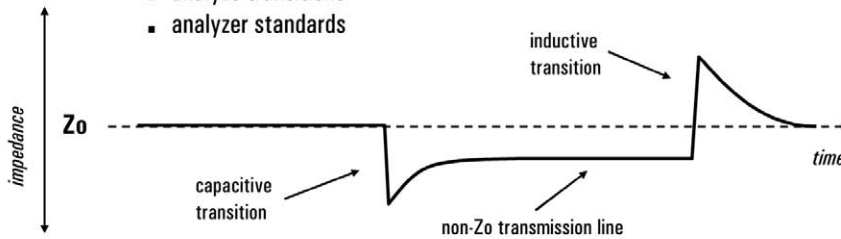
AM-PM conversion can be measured by performing a power sweep with a vector network analyzer, using the same transmission setup that we used for gain compression. The displayed data is formatted as the phase of S21 (transmission) versus power. AM-PM conversion can be computed by choosing a small amplitude increment (typically 1 dB) centered at a particular RF power level, and noting the resultant change in phase. The easiest way to read out the amplitude and phase deltas is to use trace markers. Dividing the phase change by the amplitude change yields AM-PM conversion. The plot above shows AM-PM conversion of 0.86°/dB, centered at an input power of -4.5 dBm and an output power of 16.0 dBm. Had we chosen to measure AM-PM conversion at a higher power level, we would have seen a much larger value (around 7°/dB).



Agenda

The appendix is intended to provide more detail on selected topics, such as time domain and balanced measurements, and to give pointers to reference material which covers some of these topics in more detail.

- What is TDR?
 - time domain reflectometry
 - analyze impedance versus time
 - distinguish between inductive and capacitive transitions
- With gating:
 - analyze transitions
 - analyzer standards



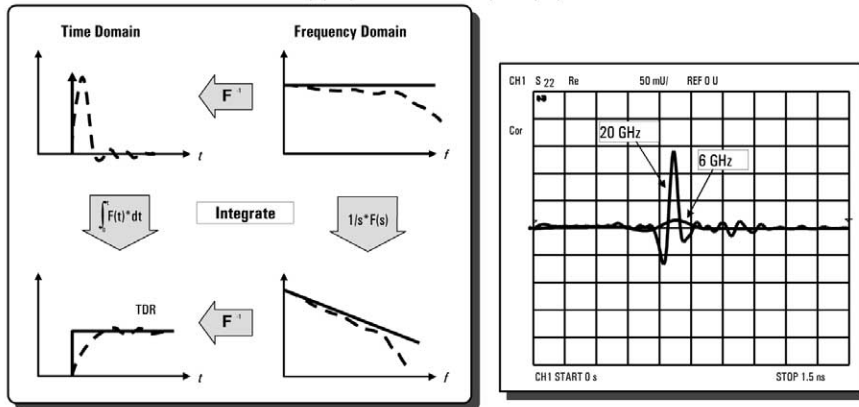
Time-Domain Reflectometry (TDR)

Time-domain reflectometry (TDR) is a very useful tool that allows us to measure impedance versus distance. One good application for TDR is fixture design and the design of corresponding in-fixture calibration standards. We can distinguish between capacitive and inductive transitions, and see non- Z_0 transmission lines. TDR can help us determine the magnitude and position of reflections from transitions within the fixture, and we can measure the quality of the calibration standards. As long as we have enough spatial resolution, we can see the reflections of the connector launches independently from the reflections of the calibration standards. It is very easy to determine which transition is which, as the designer can place a probe on a transition and look for a large spike on the TDR trace.

With time-domain gating, we can isolate various sections of the fixture and see the effects in the frequency domain. For example, we can choose to look at just the connector launches (without interference from the reflections of the calibration standards), or just the calibration standards by themselves.

Another application for TDR is fault-location for coaxial cables in cellular and CATV installations. We can use TDR in these cases to precisely determine the location of cable faults such as crimps, poor connections, shorts, opens — anything that causes a portion of the incident signal to be reflected.

- start with broadband frequency sweep (often requires microwave VNA)
- use inverse Fourier transform to compute time domain
- resolution inversely proportionate to frequency span



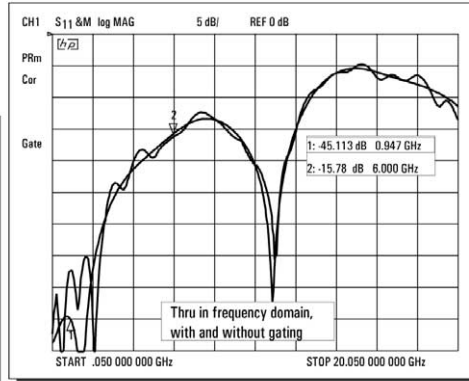
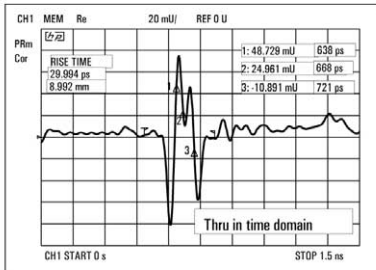
perform an inverse-Fourier transform to get from the frequency domain to the time domain, and voilà, we have the step response. Note that we could also perform the inverse-Fourier transform first, and then integrate the time-domain data. The result would be the same. The actual math used in the network analyzer is somewhat more complicated than described above, in order to take care of other effects (one example is extrapolating a value for the DC term, since the analyzer doesn't measure all the way down to 0 Hz).

To get more resolution in the time domain (to separate transitions), we need a faster effective rise time for our step response. This translates to a sharper (narrower) effective impulse, which means a broader input-frequency range must be applied to our DUT. In other words, the higher the stop frequency, the smaller the distance that can be resolved. For this reason, it is generally necessary to make microwave measurements on the fixture to get sufficient resolution to analyze the various transitions. Providing sufficient spacing between transitions may eliminate the need for microwave characterization, but can result in very large fixtures. The plot above of a fixtured-load standard shows the extra resolution obtained with a 20 GHz sweep versus only a 6 GHz sweep.

TDR Basics Using a Network Analyzer

TDR measurements using a vector network analyzer start with a broadband sweep in the frequency domain. The inverse-Fourier transform is used to transform frequency-domain data to the time domain. The figure on the left of the slide shows a simplified conceptual model of how a network analyzer derives time-domain traces. For step-response TDR, we want to end up at the lower left-hand plot. The network analyzer gathers data in the frequency domain (upper right) from a broadband sweep (note: all the data is collected from a reflection measurement). In effect, we are stimulating the DUT with a flat frequency input, which is equivalent to an impulse in the time domain. The output response of our DUT is therefore the frequency response of its impulse response. Since a step in the time domain is the integral of an impulse, if we integrate the frequency-response data of our DUT, we will have frequency-domain data corresponding to the step response in the time domain. Now, we simply

- TDR and gating can **remove** undesired reflections (a form of error **correction**)
- Only useful for **broadband** devices (a load or thru for example)
- Define **gate** to only include DUT
- Use two-port calibration



Time-Domain Gating

Gating can be used in conjunction with time-domain measurements to separate and remove undesirable reflections from those of interest. For example, gating can isolate the reflections of a DUT in a fixture from those of the fixture itself. This is a form of error correction. For time-domain gating to work effectively, the time domain responses need to be well-separated in time (and therefore distance). The gate itself looks like a filter in time, and has a finite transition range between passing and rejecting a reflection (similar to the skirts of a filter in the frequency domain).

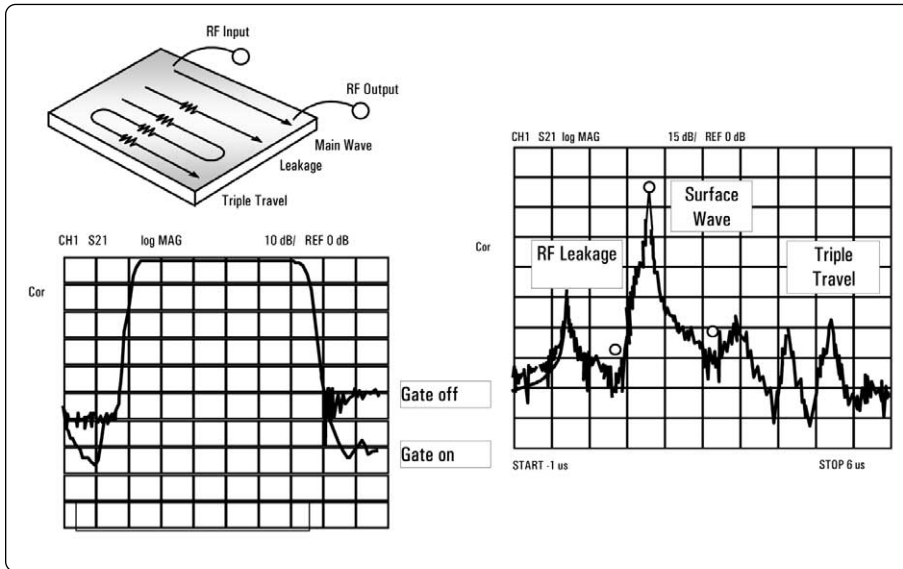
The plots above show the performance of an in-fixture thru standard (without normalization). We see about a 7 dB improvement in the measurement of return loss at 947 MHz using time-domain gating - in this case, the through standard is quite good, having a return loss of 45 dB. The gating effectively removes the effects of the SMA connectors at either end of the test fixture.

1. Set up desired frequency range (need wide span for good spatial resolution)
2. Under SYSTEM, transform menu, press "set freq low pass"
3. Perform one- or two-port calibration
4. Select S11 measurement *
5. Turn on transform (low pass step) *
6. Set format to real *
7. Adjust transform window to trade off rise time with ringing and overshoot *
8. Adjust start and stop times if desired
9. For gating:
 - set start and stop frequencies for gate
 - turn gating on *
 - adjust gate shape to trade off resolution with ripple *
10. To display gated response in frequency domain
 - turn transform off (leave gating on) *
 - change format to log-magnitude *

* If using two channels (even if coupled), these parameters must be set independently for second channel

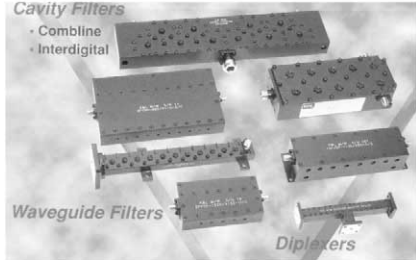
Ten Steps for Performing TDR

Here is a summary of how to perform TDR measurements. Without such a checklist, it is easy to overlook some of the more subtle steps, resulting in confusing or misleading measurements. A one-port calibration is all that is needed when characterizing connectors and the open, short and load standards. A two-port calibration is needed to characterize the reflection or line impedance of the thru standard.



Time-Domain Transmission

Time-domain transmission (TDT) is a similar tool which uses the transmission response instead of the reflection response. It is useful in analyzing signal timing in devices such as SAW filters. Gating is also useful in TDT applications. In the above example, a designer could look at the frequency response of the main surface wave without the effect of the leakage and triple-travel error signals.



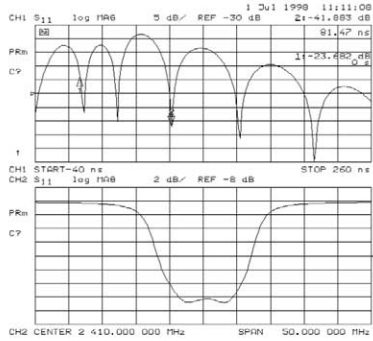
- Deterministic method used for tuning cavity-resonator filters
- Traditional frequency-domain tuning is very difficult:
 - lots of training needed
 - may take 20 to 90 minutes to tune a single filter
- Need VNA with fast sweep speeds and fast time-domain processing

Time-Domain Filter Tuning

Many communications systems use coupled-cavity-resonator bandpass filters, since they can handle high power levels and provide high rejection and sharp skirts. These filters are typically tuned to achieve the desired frequency response, which can be a tedious and time-consuming job. The difficulty of tuning these filters quickly and accurately often limits manufacturers from increasing their production volumes and reducing manufacturing costs.

Tuning these filters using the time domain response of S_{11} or S_{22} is vastly easier and faster. It is possible to tune each resonator and coupling aperture individually, since time-domain measurements can distinguish the individual responses of each filter element. Such clear identification of responses is extremely difficult (or impossible) in the frequency domain.

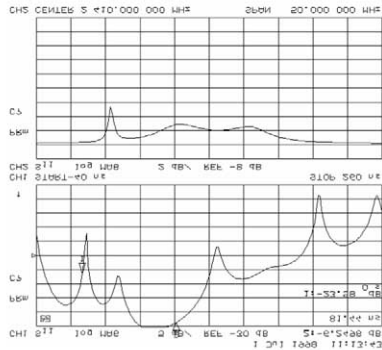
This technique requires a vector network analyzer with sufficient sweep and processing speed to allow real-time tuning while the time-domain transform is applied. Agilent's 8753, 8720, and PNA Series analyzers are all well suited for this application.



- Set analyzer's center frequency
= center frequency of the filter
- Measure S_{11} or S_{22} in the time domain
- Nulls in the time-domain response correspond to individual resonators in filter

Filter Reflection in Time Domain

The slide shows the reflection response of a well-tuned filter in both the frequency and time domains. The nulls in the time-domain response occur when the resonators are exactly tuned. The peaks between the nulls relate to the coupling factors of the filter's apertures. To set up the measurement for time-domain tuning, the frequency sweep MUST be centered at the desired center frequency of the bandpass filter. This is critical, since the tuning method will tune the filter to exactly that center frequency. The frequency span should be set to approximately two to five times the expected bandwidth.



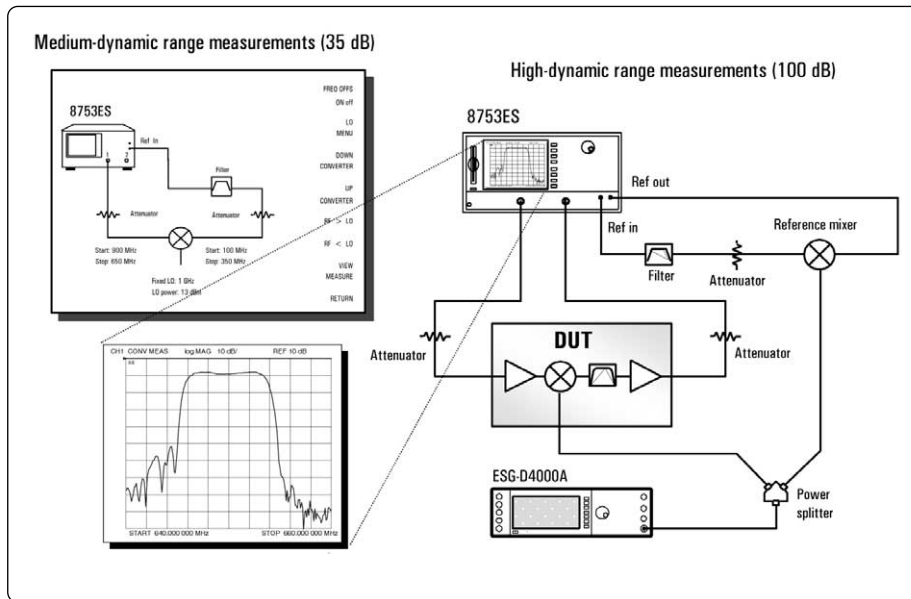
- Easier to identify mistuned resonator in time-domain: null #3 is missing
- Hard to tell which resonator is mistuned from frequency-domain response
- Adjust resonators by minimizing null
- Adjust coupling apertures using the peaks in-between the dips

Tuning Resonator #3

This slide shows that resonator #3 is mistuned, since the null for that cavity is missing. Tuning can be made even easier by overlaying the desired time-domain response in a memory trace, making it easy to see where the resonator nulls and peaks should occur. The memory trace can either be obtained from a "golden" filter or from a circuit simulation of the desired filter response.

More information about this application can be obtained from the following source:

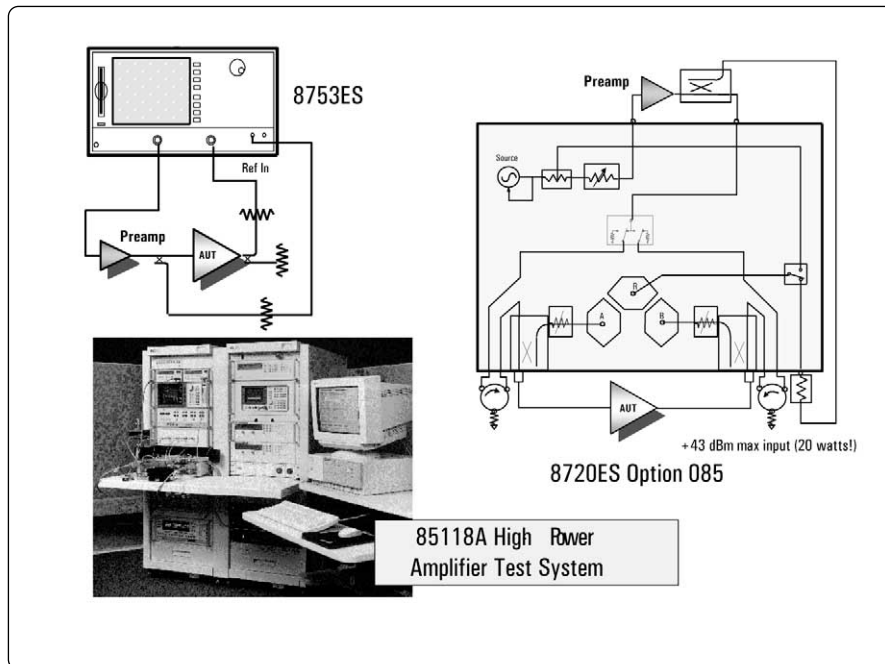
"Simplified Filter Tuning Using Time Domain", Application Note 1287-8, 5968-5328E (6/99)



Frequency-Translating Devices

Measuring frequency-translating devices requires a network analyzer that has frequency-offset capability. This means the network analyzer's source can be tuned independently from its receivers. The 8753 and 8720 series of network analyzers provided this function. The slide shows two ways to measure a frequency-translating device. In the upper left corner, a device with limited dynamic range, such as a mixer, can be measured as shown, by putting the output of the mixer directly into the reference input of the network analyzer (instead of into port 2 as is normally done). This configuration allows measurements of up to 30 dB of dynamic range, which is generally sufficient to measure the conversion loss of a mixer.

The drawing on the right half of the slide shows a more complicated setup involving a reference mixer. The reference mixer is needed to establish phase-lock of the source, and can be used as a phase reference for measurements of insertion phase or group delay. This setup utilizes the full dynamic range of the network analyzer, and is particularly useful for measuring mixer/filter combinations.



High-Power Amplifiers

More information about measuring high-power amplifiers can be obtained from the following sources:

"Using a Network Analyzer to Characterize High-Power Devices," Application Note 1287-6, 5966-3319E (4/98)

"Measurement Solutions for Test Base-Station Amplifiers," 1996 Device Test Seminar handout, 5964-9803E (4/96)

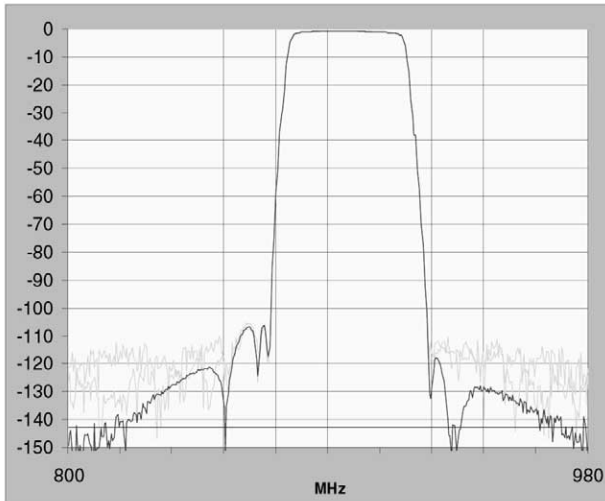
"Modern Solutions for Testing RF Communications Amplifiers," 1995 Device Test Seminar handout, 5963-5191 E (12/94)

"Amplifier Measurements using the 8753 Network Analyzer," Product Note 8753-1, 5956-4361 (5/88)

"Testing Amplifiers and Active Devices with the 8720 Network Analyzer," Product Note 8720-1, 5091-1942E (8/91)

"85108 Series Network Analyzer Systems for Isothermal, High-Power, and Pulsed Applications," Product Overview, 5091-8965E ('94)

"85118 Series High Power Amplifier Test System," Product Overview, 5963-9930E (5/95)



Take advantage of
extended dynamic
range with direct-
receiver access

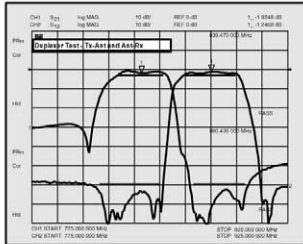
High-Dynamic Range Measurements

Extended dynamic range filter measurements can be achieved by either bypassing or reversing the coupler at test port 2 (for forward transmission measurements). By reversing the port 2 coupler, the transmitted signal travels to the “B” receiver via the main arm of the coupler, instead of the coupled arm. This increases the effective sensitivity of the analyzer by around 12 dB. To take advantage of this increased sensitivity, the power level must be decreased in the passband, to prevent the receiver from compressing. This is easily done using a segmented sweep, where the power is set high in the stopbands (+10 dBm typically), and low in the passband (-6 dBm typically). The IF bandwidth can be widened for the passband segment to speed up the overall sweep. When the port 2 coupler is reversed, 2-port error correction can still be used, but the available power at port 2 is 12 dB lower than the normal configuration.

Direct access to the receivers (which also extends dynamic range) can be achieved on PNA Series analyzers (either standard or with option 014), or by ordering Options 014 or H85 on the 8753, or Options 012 or 085 on the 8720 series. Agilent also offers a special option (H16) for the 8753 that adds a switch that can reverse the port-two coupler. The coupler can be switched to the usual configuration for normal operation.



8753 H39



Multiport analyzers and test sets:

- improve **throughput** by reducing the number of connections to DUTs with more than two ports
- allow **simultaneous** viewing of two paths (good for tuning duplexers)
- include **mechanical** or **solid-state** switches, **50** or **75** ohms
- degrade raw performance so **calibration** is a **must** (use two-port cals whenever possible)
- Agilent offers a variety of standard and custom multiport analyzers and test sets

Multiport Device Test

High-volume tuning and testing of multiport devices (devices with more than two ports) can be greatly simplified by using a multiport test set between the DUT and the network analyzer. A single connection to each port of the DUT allows complete testing of all transmission paths and port reflection characteristics. Agilent multiport test systems eliminate time-consuming reconnections to the DUT, keeping production costs down and throughput up. By reducing the number of RF connections, the risk of misconnections is lowered, operator fatigue is reduced, and the wear on cables, fixtures, connectors, and the DUT is minimized.

Agilent offers a variety of multiport test systems, both standard and custom. Some, like the 8753 H39, which is targeted to duplexer manufacturers, have built-in multiport test sets. Custom test sets can be created in 50 and 75 ohms with a variety of connector types and switching configurations, to exactly suit a user's application.

More information about measuring multiport devices can be obtained from the following source:

"Improve Test Throughput for Duplexers and Other Multiport Devices", 1996 Device Test Seminar handout, 5964-9803E (4/96)

"Novel Techniques Simplify Calibration of New Multiport Device Test System", proceedings of the 1999 European Microwave Conference.

- Single connection measurements for up to 9-port devices
- Exceptionally fast measurement speed
- Excellent dynamic range
- Built-in balanced measurements to interpret mixed-mode S-parameters
- Solid-state switches for fast and reliable measurements
- Fast and easy multiport calibration using the N4431A 4-port ECal module
- Easy to use operation



E5091A Test Set For ENA Series Analyzers

The Agilent E5091 A multiport test set combined with the 4-port ENA Series network analyzer is a complete solution for multiport device measurements. The multiport test set is available in 7- and 9-port configurations. The system is tailored for testing antenna-switch modules for mobile handsets, particularly those modules with balanced ports, although it can be used in a wide range of multiport measurement applications. The system is well-suited for both manufacturing and R&D. It has exceptionally fast measurement speed, and various features that facilitate test automation. The N4431A 4-port ECal module is available for efficient multiport calibration. Easy to use operation of the multiport system minimizes measurement setup time.



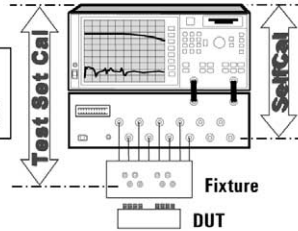
- Two and three-port calibration
- Differential measurements (mixed-mode S-parameters) for single-ended to balanced paths
- Fixture-compensation module
 - impedance transformations
 - embedding/de-embedding
 - port extensions (loss and delay) for each port
- Pass/Fail by channel
- Hardware interface to control DUT or other test gear

9-Port PNA Series Solution for High-Volume Manufacturing of LTCC/SAW Front-End Modules

Agilent has test sets that are specifically meant to work with the PNA Series of network analyzers, for the ultimate in measurement accuracy, speed, and convenience. Both solid-state and mechanical switching is available. The test sets are controlled by a Visual Basic program that runs internally in the network analyzer to coordinate the test set and VNA. The VB program also provides additional measurement capability beyond the native firmware of the network analyzer. The solution shown on the slide is targeted towards high-volume manufacturing of front-end modules for cellular telephones and other wireless appliances.



Once a month:
perform a **Test Set Cal** with external standards to remove systematic errors in the analyzer, test set, cables, and fixture



Once an hour:
automatically perform a **SelfCal** using internal standards to remove systematic errors in the analyzer and test set

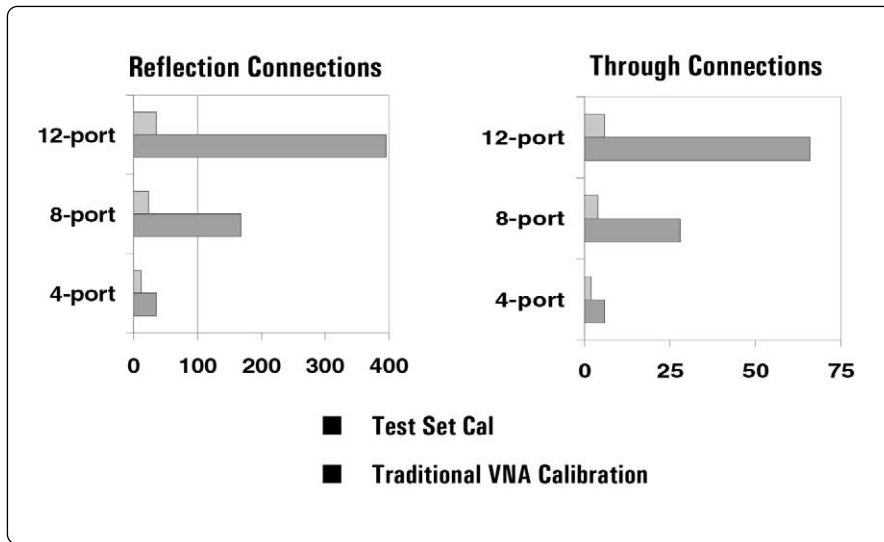
- For use with 8712E family
- 50 Ω : 3 MHz to 2.2 GHz, 4, 8, or 12 ports
- 75 Ω : 3 MHz to 1.3 GHz, 6 or 12 ports
- Test Set Cal and SelfCal dramatically improve calibration times
- Systems offer fully-specified performance at test ports

87050E/87075C Standard Multiport Test Sets

Agilent offers a line of standard multiport test sets that are designed to work with the 8712E series of network analyzers to provide a complete, low-cost multiport test system. The 87075C features specified performance to 1.3 GHz with 6 or 12 test ports (75 ohm), and the 87050E features specified performance to 2.2 GHz with 4, 8, or 12 test ports (50 ohm). These test sets contain solid-state switches for fast, repeatable, and reliable switching between measurement paths.

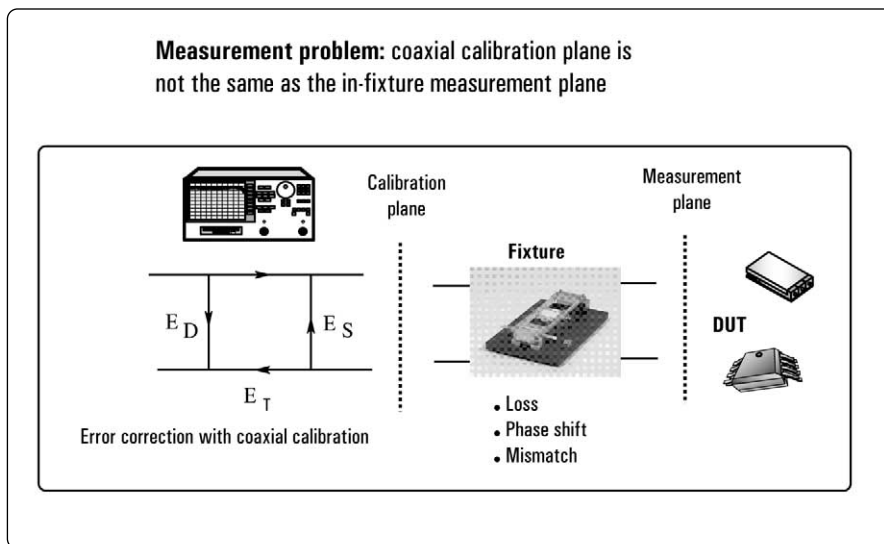
New calibration techniques can dramatically reduce the time needed to calibrate the test system. Test Set Cal is a mechanical-standards based calibration that eliminates redundant connections of reflection standards and minimizes the number of through standards needed to test all possible measurement paths. SelfCal is an internally automated calibration technique that uses solid-state switches to measure calibration standards located inside the test set. SelfCal executes automatically in just a few seconds (at a user-defined interval), restoring the measurement accuracy of the Test Set Cal. This effectively eliminates test-system drift, and greatly increase the interval between Test Set Cals. With SelfCal, a Test Set Cal needs to be performed only about once per month, unlike other test systems that typically require calibration once or twice a day. This combination can easily reduce overall calibration times by a factor of twenty or more, increasing the amount of time a test station can be used to measure components.

For more information about these multiport test systems, please order Agilent literature number 5968-4763E (50 ohm), or 5968-4766E (75 ohm).



Test Set Cal Eliminates Redundant Connections of Calibration Standards

This slide shows the decrease in the number of connections needed to fully calibrate the multiport test systems using Test Set Cal.



In-Fixture Measurements

More information about in-fixture measurements can be obtained from the following sources:

"In-Fixture Measurements Using Vector Network Analyzers", Application Note 1287-9, 5968-5329E (5/99)

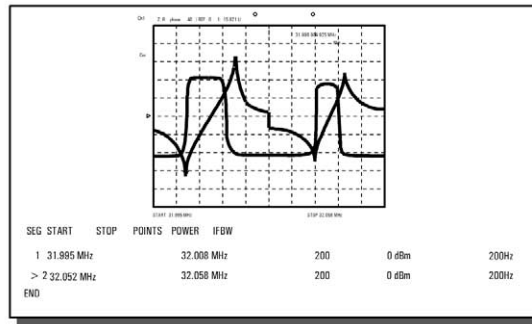
"Designing and Calibrating RF Fixtures for Surface-Mount Devices", 1996 Device Test Seminar handout, 5964-9803E (4/96)

"Accurate Measurements of Packaged RF Devices", 1995 Device Test Seminar handout, 5963-5191 E (12/94)

"Specifying calibration standards for the 8510 network analyzer", Product Note 8510-5A, 08510-90352 (1/88)

"Applying the 8510 TRL calibration for non-coaxial measurements", Product Note 8510-8A, 5091-3645E (2/92)

E5100A/B Network Analyzer



Example of crystal resonator measurement

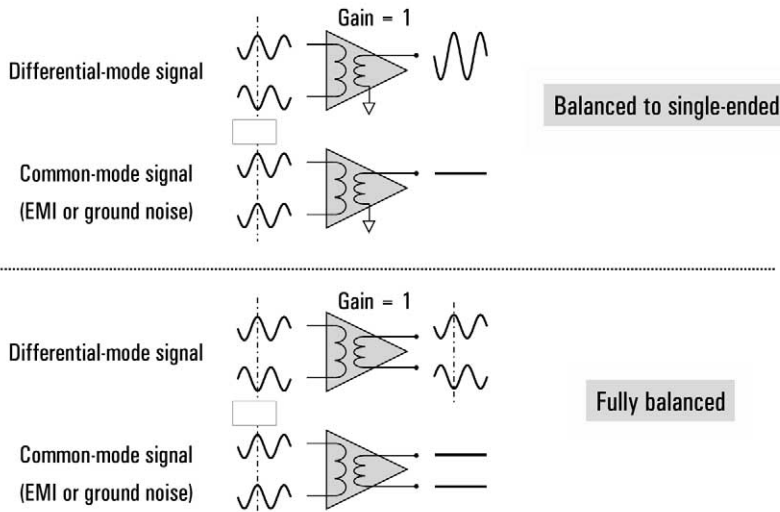
Characterizing Crystal Resonators/Filters

More information about measuring crystal resonators and filters can be obtained from the following sources:

“Crystal Resonators Measuring Functions of E5100A/B Network Analyzer”, Product Note, (5965-4972E)

“E5100A/B Network Analyzer”, Product Overview, 5963-3991E

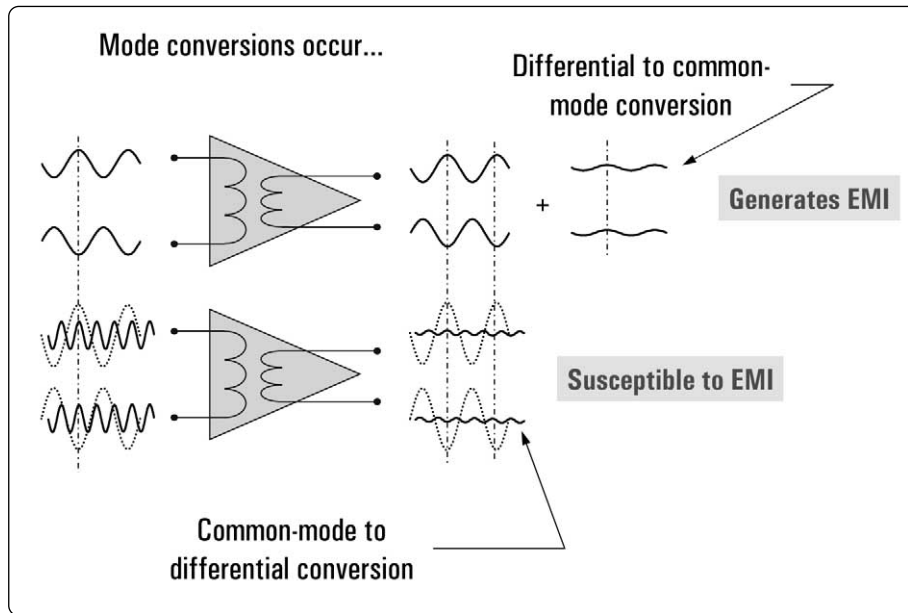
Ideally, respond to differential and reject common-mode signals



What are Balanced Devices?

Let's briefly review how balanced devices work. Ideally, a balanced device only responds to or generates differential-mode signals, which are defined as two signals that are 180° out of phase with one another. These devices do not respond to or generate in-phase signals, which are called common-mode signals. In the top example of a balanced-to-single-ended amplifier, we see that the amplifier is responding the differential input, but there is no output when common-mode or in-phase signals are present at the input of the amplifier. The lower example shows a fully balanced amplifier, which is both differential inputs and outputs. Again, the amplifier only responds to the differential input signals, and does not produce an output in response to the common-mode input.

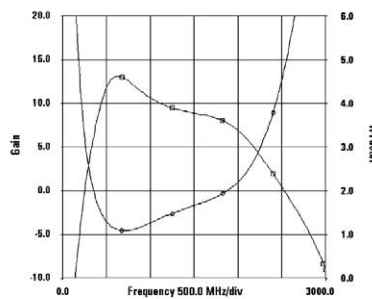
One of the main reasons that balanced circuits are desirable is because external signals that are radiated from an RF emitter show up at the terminals of the device as common mode, and are therefore rejected by the device. These interfering signals may be from other RF circuitry or from the harmonics of digital clocks or data. Balanced circuits also reject noise on the electrical ground, since the noise appears in phase to both input terminals, making it a common-mode signal.



What About Non-Ideal Devices?

We just discussed ideal balanced circuits. Real world devices on the other hand, don't completely reject common mode noise or generate only differential signals. The example on the top shows that a balanced device will produce a small amount of common mode signal that rides on top of the differential signal output. This common mode signal is the result of differential to common mode conversion, and it is a source of electromagnetic interference. The bottom example shows what can happen when a common-mode signal is present at the input to the device. Common to differential mode conversion results in a differential signal at the output of the device, which will interfere with the desired differential output, which is not shown on the slide for simplicity. This mode conversion makes a circuit susceptible to electromagnetic interference.

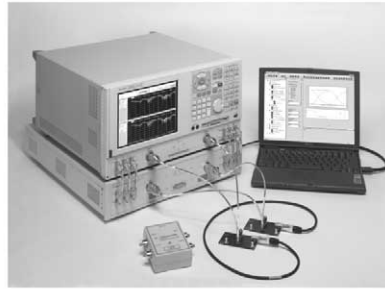
- RF and digital designers need to characterize:
- Differential to differential mode (desired operation)
- Mode conversions (undesired operation)
- Operation in non-50-ohm environments
- Other differential parameters:
 - common-mode rejection ratio
 - K-factor
 - phase/amplitude balance
 - conjugate matches



So What?

Now that you see how these balanced circuits operate, you can begin to understand why characterizing this behavior is very important to RF and digital designers. They need to characterize the differential to differential mode, which represents the desired mode of operation, as well as the undesired mode conversions. Another challenge is that differential devices and circuits often have input and output impedances other than 50 ohms. Characterizing these devices with a standard two-port network analyzer is difficult. And finally, there are other differential parameters that are also not measured with standard two-port network analyzers, such as commonmode rejection ratio or conjugate matching.

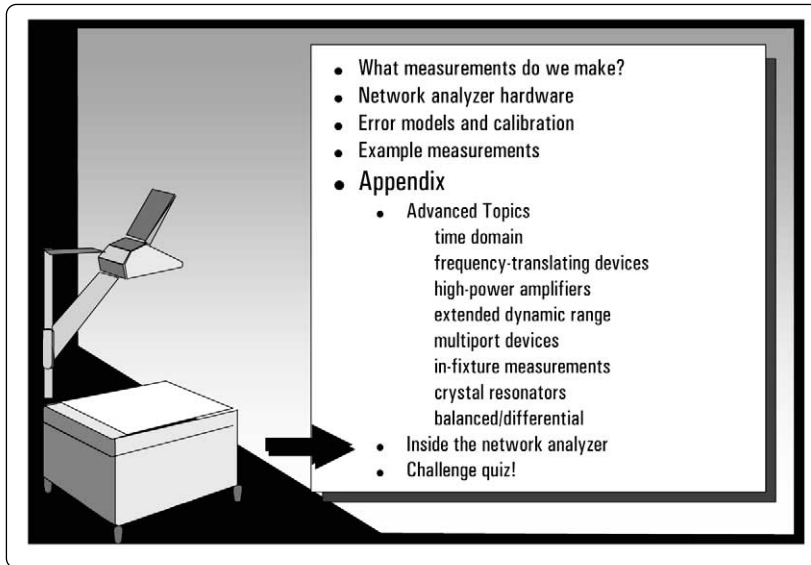
- For RF wireless and general-purpose devices:
 - ENA series is recommended choice (3, 8.5 GHz)
 - measure standard or mixed-mode S-parameters
 - fixture simulator to embed/de-embed and transform test port impedances
 - 4-port ECal for fast and easy calibration
- For microwave or signal-integrity applications:
 - N1947/48/51A systems
 - consist of network analyzer, test set, external software
 - time domain, eye diagrams, RLCG extraction
 - ECal available up to 9 GHz



Agilent Solution for Balanced Measurements

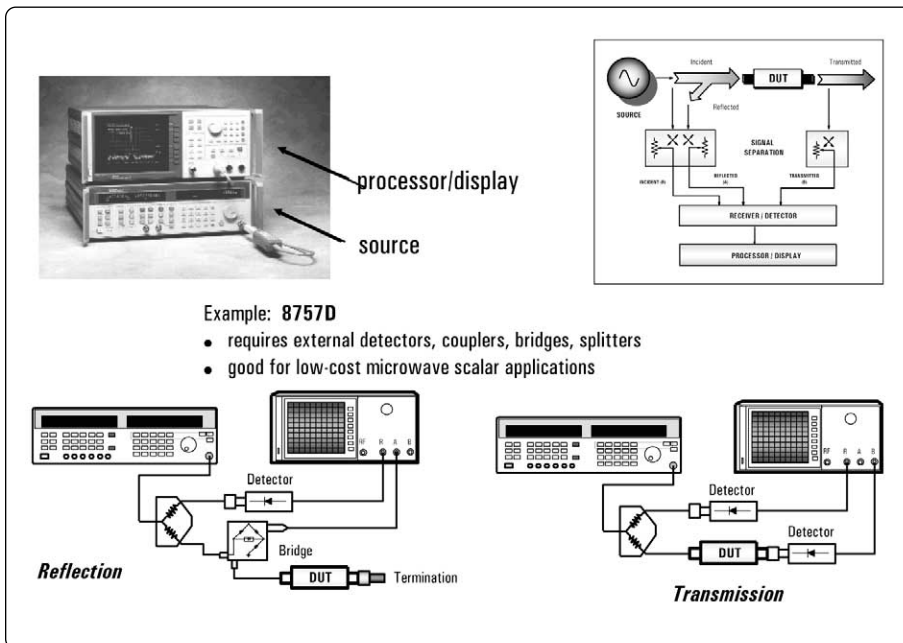
Agilent's solutions for balanced measurements address the challenges that wireless and digital designers face when characterizing differential devices. The ENA series is the best choice for wireless and general RF applications, as it offers a superior value for 4-port single-ended and balanced measurements. Data can be presented as mixed-mode S-parameters, which allows designers to clearly see the desired and undesired modes of operation.

For microwave or signal integrity applications, the N1947/48/51A physical-layer test systems are the best choice. These systems include software (N1930A) that offers many features specifically for digital designers. Agilent's physical-layer test systems provide you with a signal-integrity "expert in a box" that displays data in ways that make sense to digital-design and signal-integrity engineers. You can easily analyze your device-under-test in the time or frequency domain with exceptional accuracy. Use eye diagrams to evaluate eye closure and deterministic jitter due to your device's performance. Easily extract RLCG parameters to accurately verify transmission line designs. You can also perform "what if" analysis to test the limits of your design. Perform skew analysis to help understand mode conversions and their impact on your system's electromagnetic-interference performance. Gate out unwanted responses in the time domain, giving you a clear picture of your device's performance versus frequency. Optimize signal quality by experimenting with different compensation networks and their effect on eye diagrams.



Agenda

This next section goes into more detail about the inside workings of vector network analyzers.

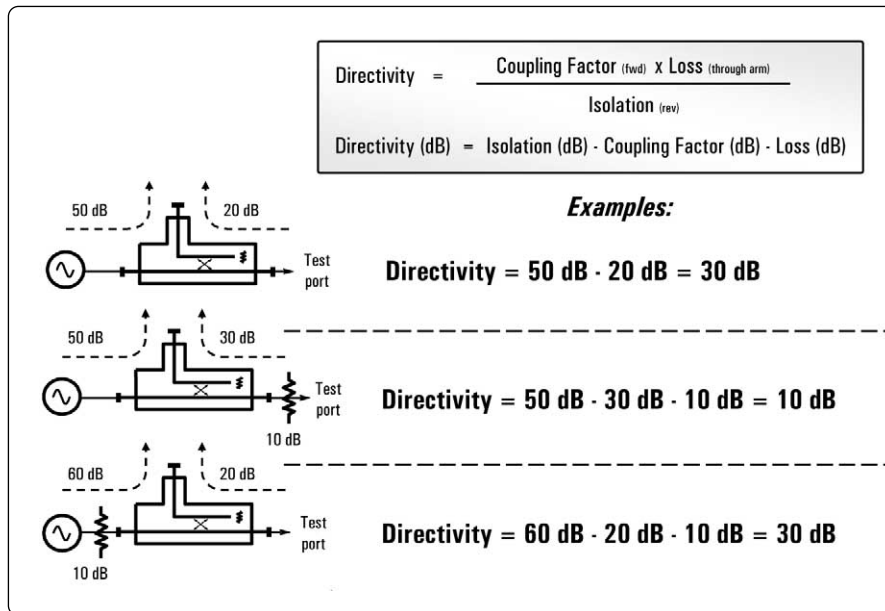


Traditional Scalar Analyzer

The following slides give more detailed information about scalar network analyzers, signal separation devices, and the receiver section within a vector network analyzer.

Here is a picture of a traditional scalar system consisting of a processor/display unit (8757D) and a stand-alone source (8750B series). This type of system requires external splitters, couplers, detectors, and bridges. While not as common as they used to be, scalar systems such as this are good for low-cost microwave scalar applications.

The configuration shown for reflection provides the best measurement accuracy (assuming the termination is a high-quality Z_0 load), especially for low-loss, bidirectional devices (i.e., devices that have low loss in both the forward and reverse directions). Alternately, the transmission detector can be connected to port two of the DUT, allowing both reflection and transmission measurements with a single setup. The drawback to this approach is that the detector match (which is considerably worse than a good load) will cause mismatch errors during reflection measurements.



Directional Coupler Directivity

One of the most important parameter for couplers is their directivity. Directivity is a measure of a coupler's ability to separate signals flowing in opposite directions within the coupler. It can be thought of as the dynamic range available for reflection measurements. Directivity can be defined as:

$$\text{Directivity (dB)} = \text{Isolation (dB)} - \text{Forward Coupling Factor (dB)} - \text{Loss (through-arm) (dB)}$$

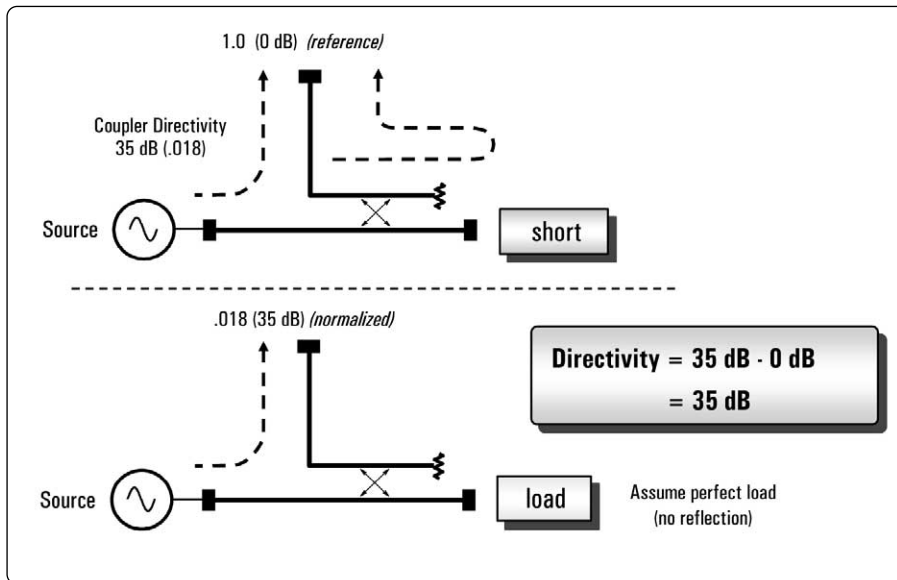
(Note: this definition deviates somewhat from the traditional definition of directivity for a dual directional coupler, which is simply the forward-coupling factor divided by the reverse-coupling factor).

In the upper example in the above slide, our coupler exhibits a directivity of 30 dB. This means that during a reflection measurement, the directivity error signal is 30 dB below the desired signal (when measuring a device with full reflection or $\rho = 1$). The better the match of the device under test, the more measurement error the directivity error term will cause.

The slide also shows the effect of adding attenuators to the various ports of the coupler. The middle example shows that adding attenuation to the test port of a network analyzer reduces the raw (uncorrected) directivity by twice the value of the attenuator. While vector-error correction can correct for this, the stability of the calibration will be greatly reduced due to the degraded raw performance.

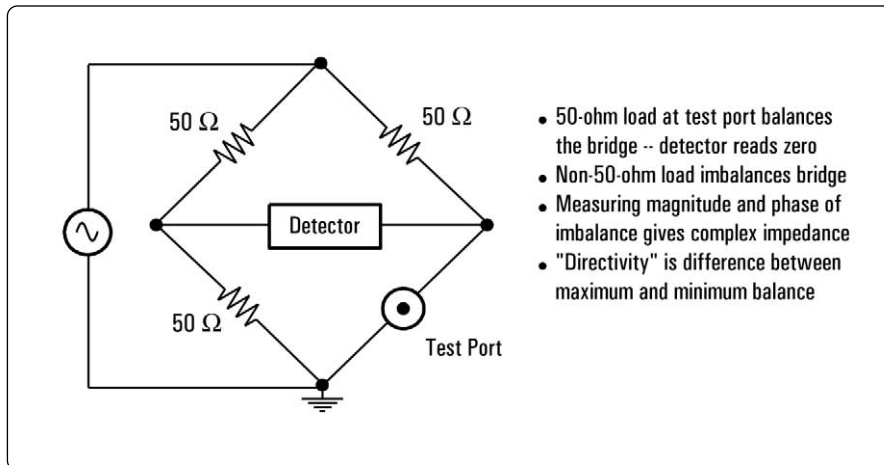
The lower example shows that adding an attenuator to the source side of the coupler has no effect on directivity. This makes sense since directivity is not a function of input-power level.

Adding an attenuator to the coupled port (not shown) affects both the isolation and forward-coupling factor by the same amount, so directivity is also unaffected.



One Method of Measuring Coupler Directivity

This is one method of measuring directivity in couplers (or in a network analyzer) that doesn't require forward and reverse measurements. First we place a short at the output port of the main arm (the coupler is in the reverse direction). We normalize our power measurement to this value, giving a 0 dB reference. This step accounts for the coupling factor and loss. Next, we place a (perfect) load at the coupler's main port. Now, the only signal we measure at the coupled port is due to leakage. Since we have already normalized the measurement, the measured value is the coupler's directivity.



Directional Bridge

Another device used for measuring reflected signals is the directional bridge. Its operation is similar to the simple Wheatstone bridge. If all four arms are equal in resistance (50 Ω connected to the test port) a voltage null is measured (the bridge is balanced). If the test-port load is not 50 Ω, then the voltage across the bridge is proportional to the mismatch presented by the DUT's input. The bridge is unbalanced in this case. If we measure both magnitude and phase across the bridge, we can measure the complex impedance at the test port.

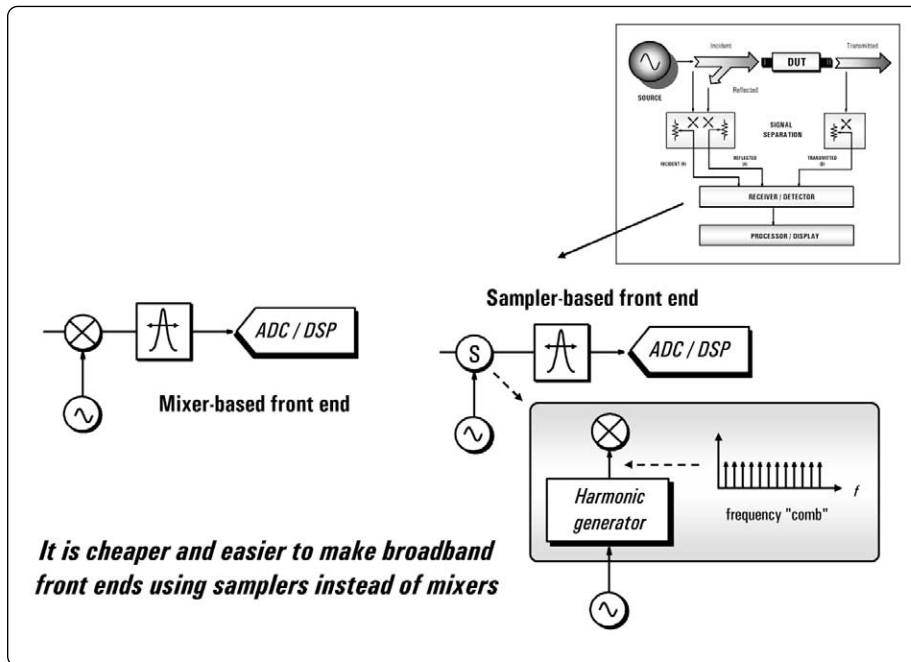
A bridge's equivalent directivity is the ratio (or difference in dB) between maximum balance (measuring a perfect Z_0 load) and minimum balance (measuring a short or open). The effect of bridge directivity on measurement uncertainty is exactly the same as we discussed for couplers.

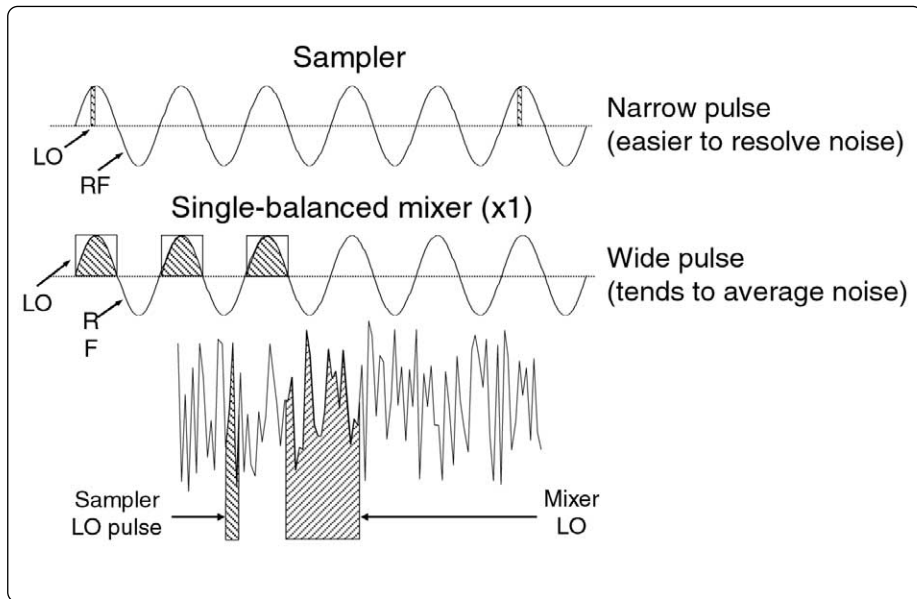
NA Hardware: Front Ends, Mixers Versus Samplers

Tuned receivers can be implemented with mixer- or sampler-based front ends. It is often cheaper and easier to make wideband front ends using samplers instead of mixers, especially for microwave frequency coverage. Samplers are used with many of Agilent's network analyzers, such as the 8753 series of RF analyzers, and the 8720 series of microwave network analyzers. The PNA Series uses mixers.

The sampler uses diodes to sample very short time slices of the incoming RF signal. Conceptually, the sampler can be thought of as a mixer with an internal pulse generator driven by the LO signal. The pulse generator creates a broadband frequency spectrum (often referred to as a "comb") composed of harmonics of the LO. The RF signal mixes with one of the spectral lines (or "comb tooth") to produce the desired IF. Compared to a mixer-based network analyzer, the LO in a sampler-based front end covers a much smaller frequency range, and a broadband mixer is no longer needed. The tradeoff is that the phase-lock algorithms for locking to the various comb teeth are more complex and time consuming.

Sampler-based front ends also have somewhat less dynamic range than those based on mixers and fundamental LOs. This is due to the fact that additional noise is converted into the IF from all of the comb teeth. Network analyzers with narrowband detection based on samplers still have far greater dynamic range than analyzers that use diode detection.

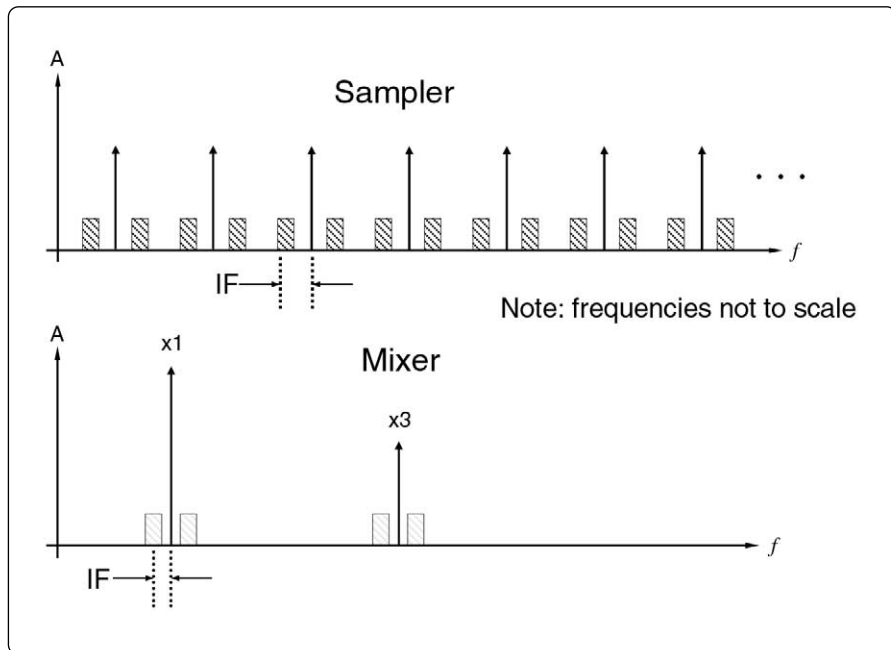




Mixers Versus Samplers: Time Domain

Let's look at the difference between samplers and mixers in the time domain first. Samplers use very narrow pulses to sample the RF input, compared to fundamental or third-order mixing. The narrow pulse is what makes a harmonic-rich LO in the frequency domain. This narrow pulse also gives more time-domain resolution, making it easier to follow the peaks and valleys of the noise. The result is that there is more noise on the IF signal.

In contrast, the mixer's LO is on for roughly half of the RF cycle, assuming a single-balanced mixer, which is typically the case for RF front ends. This longer period provides much more noise averaging. The result is less noise on the IF signal.

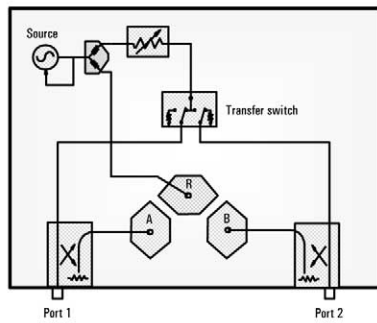


Mixers Versus Samplers: Frequency Domain

Now let's use a frequency-domain approach to explain why there is more noise conversion using samplers.

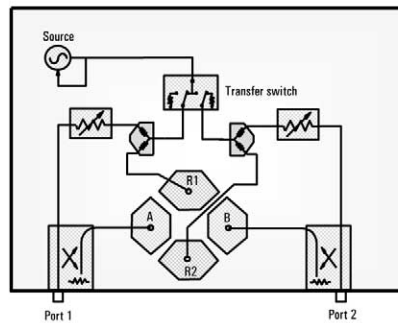
As was mentioned earlier, there are many harmonics of the LO in the frequency domain when using a sampler. Any noise present one IF away from every comb tooth, on either side, will be down-converted and detected in the IF. Since there are so many more harmonics, much more noise conversion takes place compared to using mixers, where noise is converted only around the fundamental and third harmonic of the LO. The noise multiplication effect from all of the sampler LO harmonics result in the sampler having a worse noise figure than the mixer. Typically, the difference is around 20 to 30 dB, depending on the frequencies involved.

Both the time domain and frequency domain approaches are valid ways at looking at the down-conversion process. They are just two different ways of explaining the same phenomenon.



3 receivers

- more economical
- TRL*, LRM* cals only
- includes:
 - 8753ES
 - 8720ES (standard)



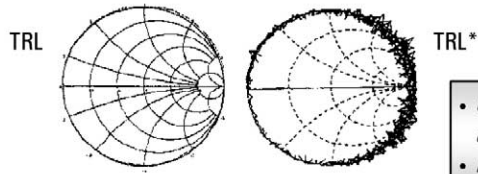
4 receivers

- more expensive
- true TRL, LRM cals
- includes:
 - PNA Series
 - 8720ES (option 400)
 - 8510C

Three Versus Four-Receiver Analyzers

As already discussed there are two main types of test sets, transmission/reflection and S-parameter test sets. The S-parameter test set has two basic test set architectures: one employing three receivers (either samplers or mixers) and one employing four receivers. The three-receiver architecture is simpler and less expensive, but the calibration choices are not as good. This type of network analyzer can do TRL* and LRM* calibrations, but not true TRL or LRM.

Four-receiver analyzers employ a second reference receiver, so forward and reverse sweeps each have their own reference receiver. This eliminates any nonsymmetrical effects of the transfer switch. Four-receiver analyzers are more expensive, but provide better accuracy for noncoaxial measurements. With a four-receiver architecture, true TRL calibrations can be performed.



- **8720ES Option 400 adds fourth sampler, allowing full TRL calibration**
- **PNA Series has four receivers standard**

- **TRL***
 - assumes the source and load match of a test port are equal (port symmetry between forward and reverse measurements)
 - this is only a fair assumption for three-receiver network analyzers
- **TRL**
 - four receivers are necessary to make the required measurements
 - TRL and TRL* use identical calibration standards
- **In noncoaxial applications, TRL achieves better source and load match correction than TRL***
- **What about coaxial applications?**
 - SOLT is usually the preferred calibration method
 - coaxial TRL can be more accurate than SOLT, but not commonly used

Why Are Four Receivers Better Than Three?

Just what is the difference between TRL and TRL*? TRL* assumes the source and load match of a test port are equal (i.e., there is port-impedance symmetry between forward and reverse measurements). This is only a fair assumption for a three-receiver network analyzer. TRL* requires ten measurements to quantify eight unknowns. True TRL calibration requires four receivers (two reference receivers plus one each for reflection and transmission) and fourteen measurements to quantify ten unknowns. TRL and TRL* use identical calibration standards. The isolation portion of a TRL calibration is the same as for SOLT.

In noncoaxial applications, TRL achieves better source match and load match correction than TRL*, resulting in less measurement error. For coaxial applications, SOLT calibration is almost always the preferred method. Agilent can provide coaxial calibration kits all the way up to 110 GHz, with a variety of connector types. While not commonly done, coaxial TRL calibration can be more accurate than SOLT calibration, but only if very-high quality coaxial transmission lines (such as beadless airlines) are used.

Option 400 for the 8720 series adds a fourth sampler, allowing these analyzers to do a full TRL calibration. The PNA Series feature four measurement receivers in the standard product.

Challenge Quiz

- 1. Can filters cause distortion in communications systems?**
 - A. Yes, due to impairment of phase and magnitude response
 - B. Yes, due to nonlinear components such as ferrite inductors
 - C. No, only active devices can cause distortion
 - D. No, filters only cause linear phase shifts
 - E. Both A and B above
- 2. Which statement about transmission lines is false?**
 - A. Useful for efficient transmission of RF power
 - B. Requires termination in characteristic impedance for low VSWR
 - C. Envelope voltage of RF signal is independent of position along line
 - D. Used when wavelength of signal is small compared to length of line
 - E. Can be realized in a variety of forms such as coaxial, waveguide, microstrip
- 3. Which statement about narrowband detection is false?**
 - A. Is generally the cheapest way to detect microwave signals
 - B. Provides much greater dynamic range than diode detection
 - C. Uses variable-bandwidth IF filters to set analyzer noise floor
 - D. Provides rejection of harmonic and spurious signals
 - E. Uses mixers or samplers as downconverters
- 4. Maximum dynamic range with narrowband detection is defined as:**
 - A. Maximum receiver input power minus the stopband of the device under test
 - B. Maximum receiver input power minus the receiver's noise floor
 - C. Detector 1-dB-compression point minus the harmonic level of the source
 - D. Receiver damage level plus the maximum source output power
 - E. Maximum source output power minus the receiver's noise floor
- 5. With a T/R analyzer, the following error terms can be corrected:**
 - A. Source match, load match, transmission tracking
 - B. Load match, reflection tracking, transmission tracking
 - C. Source match, reflection tracking, transmission tracking
 - D. Directivity, source match, load match
 - E. Directivity, reflection tracking, load match
- 6. Calibration(s) can remove which of the following types of measurement error?**
 - A. Systematic and drift
 - B. Systematic and random
 - C. Random and drift
 - D. Repeatability and systematic
 - E. Repeatability and drift
- 7. Which statement about TRL calibration is false?**
 - A. Is a type of two-port error correction
 - B. Uses easily fabricated and characterized standards
 - C. Most commonly used in noncoaxial environments
 - D. Is not available on the 8720ES family of microwave network analyzers
 - E. Has a special version for three-sampler network analyzers
- 8. For which component is it hardest to get accurate transmission and reflection measurements when using a T/R network analyzer?**
 - A. Amplifiers because output power causes receiver compression
 - B. Cables because load match cannot be corrected
 - C. Filter stopbands because of lack of dynamic range
 - D. Mixers because of lack of broadband detectors
 - E. Attenuators because source match cannot be corrected
- 9. Power sweeps are good for which measurements?**
 - A. Gain compression
 - B. AM to PM conversion
 - C. Saturated output power
 - D. Power linearity
 - E. All of the above

1. E
2. C
3. A
4. B
5. C
6. A
7. D
8. B
9. E

Answers to Challenge Quiz

The correct answers to the challenge quiz are:

1. E
2. C
3. A
4. B
5. C
6. A
7. D
8. B

Web Sources

For additional information regarding Agilent Network Analyzers, visit:
www.agilent.com/find/na

For Agilent "Back To Basics" eSeminar, series information, visit:
www.agilent.com/find/backtobasics

Agilent Technologies' Test and Measurement Support, Services, and Assistance

Agilent Technologies aims to maximize the value you receive, while minimizing your risk and problems. We strive to ensure that you get the test and measurement capabilities you paid for and obtain the support you need. Our extensive support resources and services can help you choose the right Agilent products for your applications and apply them successfully. Every instrument and system we sell has a global warranty. Support is available for at least five years beyond the production life of the product. Two concepts underlie Agilent's overall support policy: "Our Promise" and "Your Advantage."

Our Promise

Our Promise means your Agilent test and measurement equipment will meet its advertised performance and functionality. When you are choosing new equipment, we will help you with product information, including realistic performance specifications and practical recommendations from experienced test engineers. When you receive your new Agilent equipment, we can help verify that it works properly and help with initial product operation.

Your Advantage

Your Advantage means that Agilent offers a wide range of additional expert test and measurement services, which you can purchase according to your unique technical and business needs. Solve problems efficiently and gain a competitive edge by contracting with us for calibration, extra-cost upgrades, out-of-warranty repairs, and onsite education and training, as well as design, system integration, project management, and other professional engineering services. Experienced Agilent engineers and technicians worldwide can help you maximize your productivity, optimize the return on investment of your Agilent instruments and systems, and obtain dependable measurement accuracy for the life of those products.



Agilent Email Updates

www.agilent.com/find/emailupdates

Get the latest information on the products and applications you select.

Agilent T&M Software and Connectivity

Agilent's Test and Measurement software and connectivity products, solutions and developer network allows you to take time out of connecting your instruments to your computer with tools based on PC standards, so you can focus on your tasks, not on your connections. Visit
www.agilent.com/find/connectivity
for more information.

For more information on Agilent Technologies' products, applications or services, please contact your local Agilent office. The complete list is available at:
www.agilent.com/find/contactus

Phone or Fax

United States:

(tel) 800 829 4444
(fax) 800 829 4433

Canada:

(tel) 877 894 4414
(fax) 905 282 6495

China:

(tel) 800 810 0189
(fax) 800 820 2816

Europe:

(tel) 31 20 547 2111

Japan:

(tel) (81) 426 56 7832
(fax) (81) 426 56 7840

Korea:

(tel) (080) 769 0800
(fax) (080) 769 0900

Latin America:

(tel) (305) 269 7500

Taiwan:

(tel) 0800 047 866
(fax) 0800 286 331

Other Asia Pacific Countries:

(tel) (65) 6375 8100
(fax) (65) 6755 0042
Email: tm_ap@agilent.com

Product specifications and descriptions in this document subject to change without notice.

© Agilent Technologies, Inc. 2004
Printed in USA, August 31, 2004
5965-7917E



Agilent Technologies

SUNSTAR 商斯达实业集团是集研发、生产、工程、销售、代理经销、技术咨询、信息服务等为一体的高科技企业，是专业高科技电子产品生产厂家，是具有10多年历史的专业电子元器件供应商，是中国最早和最大的仓储式连锁规模经营大型综合电子零部件代理分销商之一，是一家专业代理和分销世界各大品牌IC芯片和电子元器件的连锁经营综合性国际公司，专业经营进口、国产名厂名牌电子元件，型号、种类齐全。在香港、北京、深圳、上海、西安、成都等全国主要电子市场设有直属分公司和产品展示展销窗口门市部专卖店及代理分销商，已在全国范围内建成强大统一的供货和代理分销网络。我们专业代理经销、开发生产电子元器件、集成电路、传感器、微波光电元器件、工控机/DOC/DOM电子盘、专用电路、单片机开发、MCU/DSP/ARM/FPGA软件硬件、二极管、三极管、模块等，是您可靠的一站式现货配套供应商、方案提供商、部件功能模块开发配套商。商斯达实业公司拥有庞大的资料库，有数位毕业于著名高校——有中国电子工业摇篮之称的西安电子科技大学（西军电）并长期从事国防尖端科技研究的高级工程师为您精挑细选、量身订做各种高科技电子元器件，并解决各种技术问题。

微波光电部专业研制、代理经销高频、微波、光纤、光电元器件、组件、部件、模块、整机；电磁兼容元器件、材料、设备；微波CAD、EDA软件、开发测试仿真工具；微波、光纤仪器仪表。欢迎国外高科技微波、光纤厂商将优秀产品介绍到中国、共同开拓市场。长期大量现货专业批发高频、微波、卫星、光纤、电视、CATV器件：晶振、VCO、连接器、PIN开关、变容二极管、开关二极管、低噪晶体管、功率电阻及电容、放大器、功率管、MMIC、混频器、耦合器、功分器、振荡器、合成器、衰减器、滤波器、隔离器、环行器、移相器、调制解调器；光电子元件和组件：红外发射管、红外接收管、光电开关、光敏管、发光二极管和发光二极管组件、半导体激光二极管和激光器组件、光电探测器和光接收组件、光发射接收模块、光纤激光器和光放大器、光调制器、光开关、DWDM用光发射和接收器件、用户接入系统光收发器件与模块、光纤连接器、光纤跳线/尾纤、光衰减器、光纤适配器、光隔离器、光耦合器、光环行器、光复用器/转换器；无线收发芯片和模组、蓝牙芯片和模组。

更多产品请看本公司产品专用销售网站：[欢迎索取免费详细资料、设计指南和光盘](#)；产品凡多，未能尽录，欢迎来电查询

商斯达中国传感器科技信息网：<http://www.sensor-ic.com/>

商斯达工控安防网：<http://www.pc-ps.net/>

商斯达电子元器件网：<http://www.sunstare.com/>

商斯达微波光电产品网：[HTTP://www.rfoe.net/](http://www.rfoe.net/)

商斯达消费电子产品网：<http://www.icasic.com/>

商斯达实业科技产品网：<http://www.sunstars.cn/> 微波元器件销售热线：

地址：深圳市福田区福华路福庆街鸿图大厦1602室

电话：0755-82884100 83397033 83396822 83398585

传真：0755-83376182 (0) 13823648918 MSN: SUNS8888@hotmail.com

邮编：518033 E-mail:szss20@163.com QQ: 195847376

技术支持: 0755-83394033 13501568376