

Chapter 2

Specification-Governed Telecommunication and High-Frequency-Electronics Aspects of Low-Noise Amplifier Research

Abstract The first part of this chapter focuses on placing millimeter-wave research in the context of telecommunication. The second part focuses more strongly on some high-frequency amplifier electronics that were neglected in Chap. 1. In essence, this chapter aims to illustrate the convergence of communications, circuits and antennas, which is necessary in millimeter-wave LNA research. These apparently unrelated aspects of LNA research can furthermore be treated in a single chapter because they set the research constraints (i.e., result in design specifications). For example, there may be a requirement for a 60 GHz communication network deploying a certain type of modulation, but with a particular gain and noise figure; this, once again, illustrates the multidisciplinary nature of LNA research.

This chapter will open Part I of the book, the aim of which is to contextualize the research presented in Part II. It will combine two aspects, both extremely important in research into LNAs. The first part of this chapter will focus on placing millimeter-wave research in the context of telecommunication. The second part of the chapter will focus more strongly on some high-frequency amplifier electronics that were neglected in Chap. 1. In essence, this chapter aims to illustrate the convergence of communications, circuits and antennas, which is necessary in millimeter-wave LNA research [1]. These apparently unrelated aspects of LNA research can furthermore be treated in a single chapter because they set the research constraints (i.e., result in design specifications). For example, there may be a requirement for a 60 GHz communication network deploying a certain type of modulation, but with a particular gain and noise figure; this, once again, illustrates the multidisciplinary nature of LNA research.

This chapter is set to complement the fundamental theory of LNAs introduced in Chap. 1. The first part of the chapter starts with the introduction of the concept of wavelength, followed by the analysis of the frequency spectrum and various transmission bands and their implications for transceiver system research. The millimeter-wave portion of the spectrum itself is quite wide and there are a few sub-frequency allocations that are the focus of active research and warrant further discussion. The feasibility of passive component implementations, especially

transmission lines, in each range is investigated further, seeing that the component size is dependent on wavelength. Various digital modulation schemes, commonly used in transceiver systems, are also presented in some detail, since the modulation schemes often steer the direction of research. The first part of the chapter also examines some of the antenna and propagation theory, specifically as applicable to millimeter-waves. A typical LNA connects either directly to the antenna or via a bandpass filter, and needs to amplify extremely weak and noisy signals and an effort thus needs to be made to research the antennas as well to avoid introducing additional attenuation on the receiving antenna.

Secondly, a large part of this chapter investigates high-frequency electronics. In general, electronics at high frequencies are treated by means of a two-port analysis. This part of the chapter therefore mostly focuses on the two-port theory as applicable to LNA research. This includes two-port modeling and applicable parameters, which are important for practical LNA analysis. This is then followed by a discussion of the amplifier and LNA fundamentals, but mostly structured from the two-port perspective. This includes various two-port gain equations and the stability of amplifiers. Furthermore, concepts such as two-port impedance matching and biasing are examined to some extent. Another two very important concepts introduced in Chap. 1 will be investigated here in much detail and from the two-port perspective: noise and linearity.

2.1 Frequency and Wavelength

Frequency of operation has a major influence on the behavior of passive and active devices and thus on any transceiver component that is built to operate at a certain frequency. Inherently, building LNAs for increased frequencies becomes progressively more challenging.

The active device responsible for the gain of the amplifier is the transistor. If one looks at the physical structure of the transistor, it is not difficult to identify a number of parasitic components, primarily capacitors. If these are worked into amplifier gain equations, they result in a frequency response of the amplifier that is not flat over all frequencies, but rather starts decreasing with an increase in frequency. A transistor's transitional frequency is typically used as the measure of the maximum frequency at which an amplifier can be designed to operate. This chapter will assume that a transistor or a technology is chosen such that an active device has a transitional frequency, which is going to allow operation well into the millimeter-wave range (discussed in the following section); thus the performance of the amplifier is only limited by the performance of the passive components, unless otherwise stated. Of course this is often untrue, but the issue will be deferred until later because active devices and their modeling will be researched in Chap. 3 in great detail.

When discussing passive components, especially transmission lines, it is somewhat more important to relate the feasibility of passives to a wavelength. The size of the antenna also depends on the wavelength, requiring wavelength to be defined early in this chapter.

The frequency is related to wavelength according to a well-known relation

$$\lambda = \frac{v}{f} \quad (2.1)$$

where v is the phase speed of the wave and f is the wave frequency. The phase speed of the electromagnetic wave in free space is the speed of light, which is about 3×10^8 m/s, but decreases in semiconductors or substrates that have relative permittivity of more than 1.0 (slow waves) [1]. At lower frequencies, the wavelengths of signals are quite large, so the size of the passive electrical components has little impact on these signals. For example, at 2.4 GHz, a frequency where most commercial WiFi systems operate, the wavelength is 12.5 cm. This means that any component or a connection should not be greater than a tenth of the wavelength (12.5 mm) for a system to behave with minimal loss, i.e. so that the wave propagation does not need to be taken into account and transmission lines can be avoided. This can still be accomplished even on a PCB where tracks are generally longer. At 60 GHz, the wavelength is 5 mm, which means that any connection greater than one tenth of 5 mm, or 500 μm , has to be treated as a transmission line. This motivates the on-chip transmission line implementation of passives and matching of all circuitry.

2.2 Frequency Spectrum and Transmission Bands

Based on the discussion in the previous section, at a frequency of 30 GHz, the wavelength of a transmitted signal has a value of 1 cm. This value has made it convenient for this point in the spectrum to be taken as the transition between microwave and millimeter-wave ranges. Similarly, the frequency of 3 GHz ($\lambda = 10$ cm) was taken as the transition between RF and microwave. Coincidentally, transceivers constructed by lumped elements can be more compact than designs based on transmission lines up to the higher end of the RF range, whether built on chip or off chip (discrete implementation). Above 30 GHz, or in millimeter-wave range, transmission lines and waveguides are more practical (albeit with some exceptions generally for good lumped component designs, as discussed later in Chap. 4), even on chip. Furthermore, transceivers and their elements require accurate modeling and high-precision manufacturing in this range. The microwave range, squeezed between 3 and 30 GHz, presents almost a “gray area” of amplifier design, where both lumped and transmission-line implementations are possible, depending on the application and methodology and whether the system is implemented on chip, on package or on a PCB.

Research into LNA and other transceiver components regularly leads to articles that use alternative, more detailed nomenclature of the frequency spectrum division (that is other than RF, microwave and millimeter-wave), and all these definitions may be confusing to the reader. Thus the classification of the different bands in the frequency spectrum is beneficial for understanding LNA design and applications and is included in this section. Frequency bands are defined by the international telecommunication union (ITU) [2]. The frequency spectrum is illustrated in Table 2.1. The same table also summarizes the feasibility of passives in each frequency range.

The lower end of the spectrum, that is, the ELF, VF and VLF ranges, spans from 30 Hz to 30 kHz, contains audible frequencies and is thus unsuitable for radio transmission. Low frequencies (LF) span from 30 to 300 kHz and are used for long-range navigation, submarine communication and telegraphy. Medium frequencies (MF) or medium waves span from 300 kHz to 3 MHz and are used for commercial radio. The high-frequency (HF) range with frequencies from 3 to 30 MHz is used for military tactical radios and by amateur radio operators because of the long-distance propagation properties of the 30-m-long waves.

The very-high-frequency (VHF) range with frequencies from 30 to 300 MHz and the ultra-high-frequency (UHF) range with frequencies from 300 MHz to 3 GHz are used for television and radio broadcast, cordless and cellular telephone transmission, as well as for other wireless applications, such as wireless local area networks (WLANs) and Bluetooth®. These bands are also suitable for industrial heating and microwave ovens.

The higher end of the spectrum comprises of super-high-frequency (SHF) and extra-high frequency (EHF) bands. The SHF range includes frequencies from 3 to 30 GHz and the EHF range includes frequencies from 30 to 300 GHz. These two ranges are mostly used for satellite communication and radar applications, but as mentioned in Chap. 1, millimeter-waves have great potential for other commercial applications. The terahertz range is currently used mostly for fiber communications, and infrared remote controls, such as the remote controls for television.

More precise frequency division based on the letters of the alphabet is sometimes used for the UHF, SHF and EHF frequency bands. The *L*-band spans from 1 to 2 GHz, the *S*-band spans from 2 to 4 GHz, the *C*-band spans from 4 to 8 GHz and the *X*-band spans from 8 to 12 GHz. The next band is typically divided into three bands, the *Ku*, *K* and *Ka* bands, where “u” stands for “under” and “a” stands for “above”. The *Ku* band spans from 12 to 18 GHz, the *K* band from 18 to 26.5 GHz and the *Ka*-band from 26.5 to 40 GHz. The top part of the *Ka*-band is already in the millimeter range, together with the *V*-band (40–75 GHz) and *W*-band (75–110 GHz). The *D*-band starts at 110 GHz.

Electronic circuitry and the required IC technologies for applications above 110 GHz are still in the early phases of research, and it remains customary to refer to any range above 110 GHz as simply the millimeter-wave range, or if the application is above 300 GHz, the sub-terahertz range. The sub-terahertz range coincides with the far-infrared range, which is one decade lower than the infrared range and two decades lower than visible light.

Table 2.1 The frequency spectrum and feasibility of passives

Frequency range	Wavelength range	Range name	UHF/SHF/EHF band name	Other names	Feasibility of passives
30–300 Hz	10,000–1000 km	Extremely low frequency (ELF)	–	No transmission	–
300–3000 Hz	1000–100 km	Voice frequency (VF)		RF	Lumped passives
3–30 kHz	100–10 km	Very low frequency (VLF)			
30–300 kHz	10–1 km	Low frequency (LF)			
300–3000 kHz	1000–100 m	Medium frequency (MF)			
3–30 MHz	100–10 m	High frequency (HF)			
30–300 MHz	10–1 m	Very high frequency (VHF)			
300–1000 MHz	100–30 cm	Ultra-high frequency (UHF)			
1–2 GHz	30–15 cm				
2–3 GHz	15–10 cm				
3–4 GHz	10–7.5 cm				
4–8 GHz	7.5–3.75 cm	Super-high frequency (SHF)	C-band	Microwave	Lumped passives, distributed passives (transmission lines) off-chip
8–12.4 GHz	37.5–24 cm		X-band		
12.4–18 GHz	2.4–1.7 cm		Ku-band		
18–26.5 GHz	1.7–1.1 cm		K-band		
26.5–30 GHz	1.1–1 cm		Ka-band		
30–40 GHz	10–1 mm		Extremely high frequency (EHF)		
40–75 GHz		W-band			
75–110 GHz		D-band			
110–170 GHz		–			
170–300 GHz		–			
300 + GHz	< 1 mm	Sub-terahertz waves	–	THz frequencies Far infrared	

As discussed in the opening sections of this book, the main focus is on millimeter-wave frequencies. However, this frequency range is so wide that it is necessary to focus more specifically on the particular parts of the millimeter-wave spectrum. This will be handled in the following section.

2.3 The Millimeter-Wave Frequency Range

As seen in the previous section, the millimeter-wave part of the electromagnetic spectrum, or EHF, according to the definition by the ITU, spans the frequency range of 30–300 GHz. Being one of the least explored frequency bands, its use across the world is not standardized, and therefore in different countries, different amounts of bandwidth are available in different frequency groups, which is typically regulated on government level. Nevertheless, the amount of bandwidth available is still much more than in UHF and SHF, where the communication networks of today (WiFi, global positioning system, cellular) primarily operate.

One of the main considerations in millimeter-wave research regarding the abundance of bandwidth allocation is millimeter-wave propagation. The amount of atmospheric attenuation increases with frequencies, albeit with some exceptions.

2.3.1 Millimeter-Wave Bandwidth Allocations

Frequency allocation remains one of the major bottlenecks of millimeter-wave research. Even though, by definition, the millimeter-wave range spans a total bandwidth of 270 GHz, only a small percentage of the total frequency allocation is typically available without the need for licensing, which is an incentive for research into commercial applications. If one considers the frequency distribution as the one shown in Fig. 2.1 [3], it appears that only two bands are available, one centered at 60 GHz and the second one at 180 GHz. This leaves about 225 GHz of bandwidth that is unused.

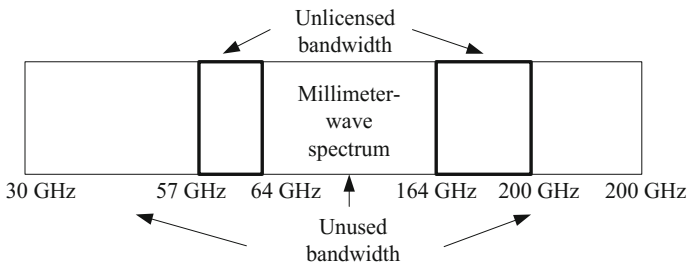


Fig. 2.1 Millimeter-wave bandwidth distribution

The largest research effort into millimeter-wave communication went into the unlicensed 60-GHz band, even though this band suffers from large amounts of oxygen attenuation, an aspect that will be discussed later. Europe allows transmissions between 57 and 66 GHz, which translates to a 9-GHz range. This, as mentioned before, is not standardized across the world, and the United States of America (USA), for example, allows 2 GHz less, that is, transmission between 57 and 64 GHz. Channels are typically just over 2 GHz, with top and bottom guards of 240 and 120 MHz respectively [4], as illustrated in Fig. 2.2. The first standard governing bandwidth usage around 60 GHz was the IEEE 802.11ad standard [5].

The 60-GHz band is also used for satellite applications. Beyond 60 GHz, however, there are some reserved frequency allocations. For example, it was already mentioned in Chap. 1 that the 77 GHz (76–77 GHz) band is utilized for automotive radar and safety applications. Automotive radar application around 77 GHz quickly gained popularity for several reasons (previously, automotive radar used the 24 GHz band) [6]. Firstly, the size of the radar antenna allowed for seamless integration into, for example, the bumpers of vehicles [7]. Secondly, the millimeter-wave propagation characteristics at these frequencies allowed for practical narrow beams. Finally, the amount of radiated power allowed in this band is also greater.

In Europe, the 77–81 GHz window is allocated for UWB short-range radar. 77-GHz bands are sometimes also used for millimeter-wave imaging. Imaging is sometimes also done around 94 GHz, where oxygen attenuation exhibits the local minimum.

Furthermore, the band around 40 GHz is used for licensed high-speed microwave data links in the USA. The 71–76, 81–86 and 92–95 GHz bands are also used for point-to-point communication links because the oxygen attenuation is not as prominent as in the 60 GHz band, but licensing is needed. Other millimeter-wave frequencies are typically used for radio astronomy.

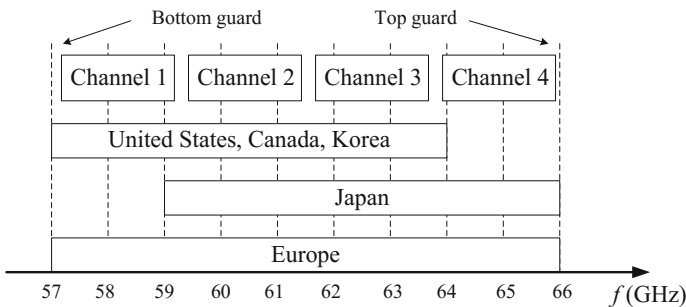


Fig. 2.2 60-GHz millimeter-wave band frequency allocation

2.3.2 Propagation of Millimeter Waves

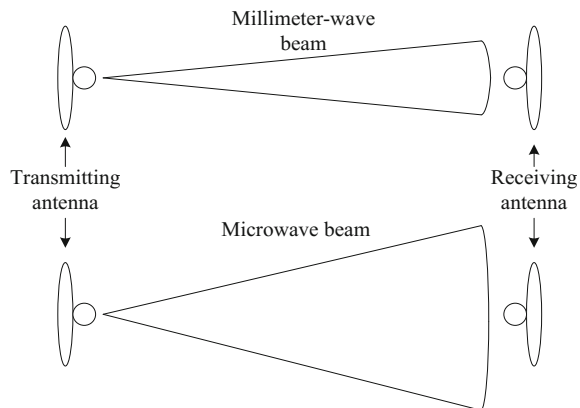
For transmission above 30 MHz, propagation is possible only by line-of sight waves [8]. There are two types of these waves: direct waves and ground-reflected waves. Most of the transmission at millimeter-wave frequencies are accomplished by direct waves. Propagation of the waves is influenced by different types of losses, including space loss, atmospheric loss, polarization mismatch loss, impedance mismatch loss and pointing loss.

It is a common misconception that propagation of the waves is influenced by frequency, in other words that as the frequency increases, the waves cannot propagate as well [3]. One counter-argument to that is the size of the antenna; as will be seen later in this chapter, the required size of the antenna is proportional to the wavelength and therefore at millimeter-wave frequencies, as waveforms decrease, so does the size of the antenna, and more antennas can be packed into a smaller space.

The concept of the increase of the number of antennas in the millimeter waves is closely related to the directivity of the millimeter-wave beam. It is possible to build highly directional antennas, resulting in narrow transmitted beams. This allows for more channels to be present in a small geographical area, thus increasing overall data throughput. This is best understood if the size of the millimeter-wave beam is compared to the size of, for example, a microwave beam, as illustrated in Fig. 2.3 [9].

Despite these advantages, the atmospheric loss does tend to limit the range of millimeter-wave communication. This is mostly due to the absorption of atmospheric gases. Oxygen (O_2) absorption is coincidentally most prominent at 60 GHz (around the range of the unlicensed bands both in Europe and the USA) with an attenuation of over 10 dB/km. Water vapor (H_2O) absorption also influences the propagation, but it is more prominent above 100 GHz. Atmospheric attenuation of electromagnetic waves is well researched and is usually represented by the attenuation curves 8 depicted in Fig. 2.4.

Fig. 2.3 Directivity of the millimeter-wave beam compared to the microwave beam



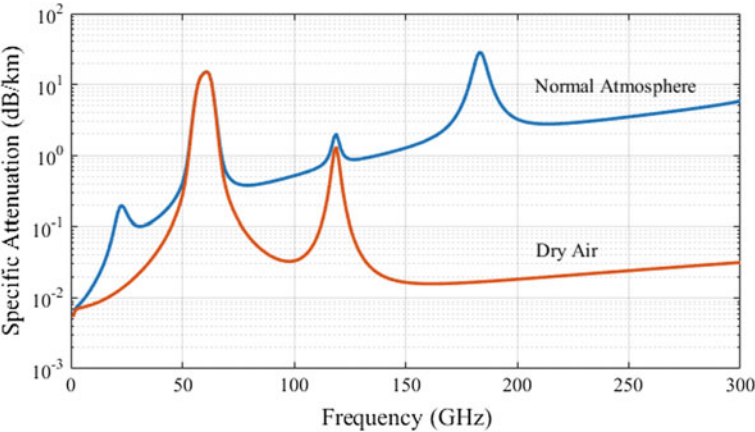


Fig. 2.4 Atmospheric attenuation of electromagnetic waves (sea level) for normal atmospheric air and dry air

Precipitation is also something that causes attenuation. This is due to the fact that the electromagnetic waves are scattered by raindrops. Consider, for example, applications above 70 GHz where oxygen attenuation is negligible. Here, communication would be seriously disrupted by precipitation, with the attenuation figures as shown in Table 2.2 [9].

Signal loss due to atmospheric conditions therefore allows reliable outdoor communication to ranges of only a few kilometers.

Another major cause of loss in millimeter-waves is multipath interference. This results in reflections and scattering. Scattering arises in response to sizes of objects that are physically similar to the wavelength, while reflection occurs if waves reach objects that are larger in size than the wavelength. With millimeter waves, where wavelengths are smaller than 10 mm, most objects act as reflectors. This is why the line of sight, mentioned at the beginning of this section, is so important. Hence this also influences the environmental use of this frequency band i.e. for indoor or outdoor propagation. Careful selection of both transmit and receive antennas can help alleviate the multipath issue to some extent [10].

Table 2.2 Atmospheric loss between 70 and 80 GHz in various conditions

Effect	Conditions	Signal loss (dB/km)
No precipitation	Sea level	0.22
Light rain	1 mm/h	0.9
Humidity	100% at 30 °C	1.8
Moderate rain	4 mm/h	2.6
Heavy fog	10 °C, 50 m visibility	3.2
Heavy rain	22 mm/h	10.7
Intense rain	50 mm/h	18.4

2.4 Digital Modulation Schemes for Millimeter-Wave Applications

In Chap. 1, two block diagrams were presented: one of the complete transmitter/receiver, and one of the zero-IF direct conversion transmitter. One aspect common to both diagrams was that both introduced the concept of modulation. Modulation implies that properties of the carrier signal, i.e. the signal that can physically be amplified and transmitted, are varied so that the information of interest is superimposed [11, 12]. Thus, before a signal is transmitted, a certain modulation scheme needs to be deployed. On the receiver's side, the signal is demodulated, therefore the LNA is required to amplify the received signal, while retaining the properties of the signal achieved by modulation. Modulation is accomplished by means of a modulator. On the receiver side, a demodulator is used to recover the same information.

Modern telecommunication systems are moving from employing analog modulation towards employing digital modulation, and this is especially true for the millimeter-wave range. In a digital modulation scheme, the carrier signal is modulated by a discrete signal. To fully understand the requirements of the design of an LNA for a particular application that may be of interest to the reader, at least minimal understanding of the modulation scheme that is deployed is required, and therefore, various digital modulation schemes are discussed in this review. The modulation schemes listed below all have different bandwidth utilization efficiency:

- On-off keying (OOK);
- Phase-shift keying (PSK);
- Frequency-shift keying (FSK);
- Pulse-amplitude modulation (PAM);
- Quadrature amplitude modulation (QAM);
- Orthogonal frequency-division multiplexing (OFDM); and
- Various direct-sequence spread spectrum (DSSS) techniques.

OOK, PSK, PAM and QAM are referred to as single-carrier modulation schemes. FSK and OFDM are examples of multi-carrier modulation schemes. DSSS refers to methods by which a specific bandwidth signal is intentionally spread in the frequency domain by using a spreading sequence, resulting in a signal with a wider bandwidth. This is done for various reasons, most often to make the signal less prone to noise or to prevent unintended detection [13]. Because of its complexity, the discussion on the DSSS will be omitted.

If the mapping of the digital sequence is performed without requiring the information on the previously transmitted signals, then the modulation is memoryless. By definition then, all the modulation schemes discussed here are memoryless.

2.4.1 On-Off Keying

OOK is the simplest digital modulation technique. In this scheme, the presence of the carrier indicates a digital one (1), and the absence of the signal indicates a digital zero (0), as illustrated in Fig. 2.5. Because of its simplicity, the data transfer rates that can be achieved with OOK are the same as BPSK, and thus, in millimeter-wave applications, more complex modulation techniques are typically used, such as the ones discussed in subsequent sections.

2.4.2 Phase Shift-Keying

If the digital signal is modulated onto the carrier by changing its phase, then PSK is accomplished. A finite number of phases (M) is used, usually two (0 and 180°) for bits 0 and 1 (BPSK), four (0, 90°, 180°, 270°) for bit combinations 00, 01, 10 and 11 (QPSK) or eight for eight three-bit combinations (octal PSK). If the signal pulse has the shape defined by $g(t)$, then M signals waveforms (symbols) are represented by

$$s_m(t) = g(t) \cos \frac{2\pi}{M}(m-1) \cos 2\pi f_c t - g(t) \sin \frac{2\pi}{M}(m-1) \sin 2\pi f_c t \quad (2.2)$$

for $m = 1, 2, \dots, M$, $0 \leq t \leq T$, where f_c is the carrier and T is the period of the signal. A typical waveform of the QPSK-modulated signal is shown in Fig. 2.6. Complex modulation schemes are often also represented by signal space diagrams or constellations. Constellations of BPSK, QPSK and octal PSK are shown in Fig. 2.7.

The bandwidth utilization efficiency of BPSK is 1 bps/Hz, but the bandwidth increases as the complexity of modulation is increased. For example, in the 60-GHz band, a data rate of 3.5 Gb/s can be achieved by QPSK in a 2.16-GHz channel [14].

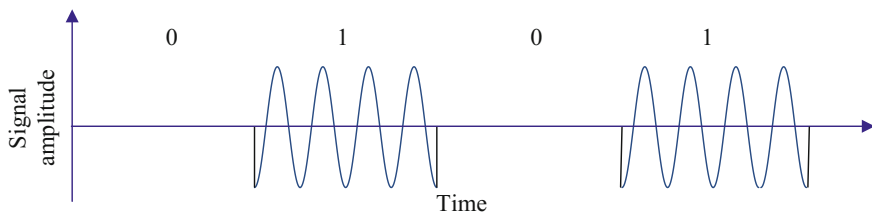


Fig. 2.5 OOK-modulated signal waveform

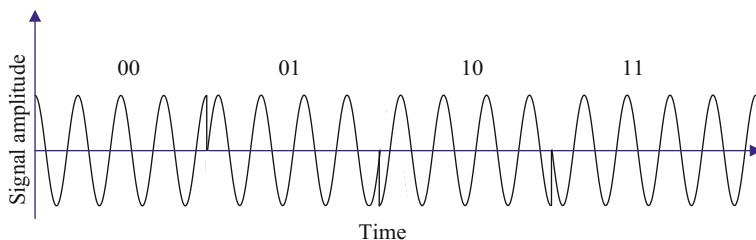


Fig. 2.6 The QPSK-modulated signal waveform

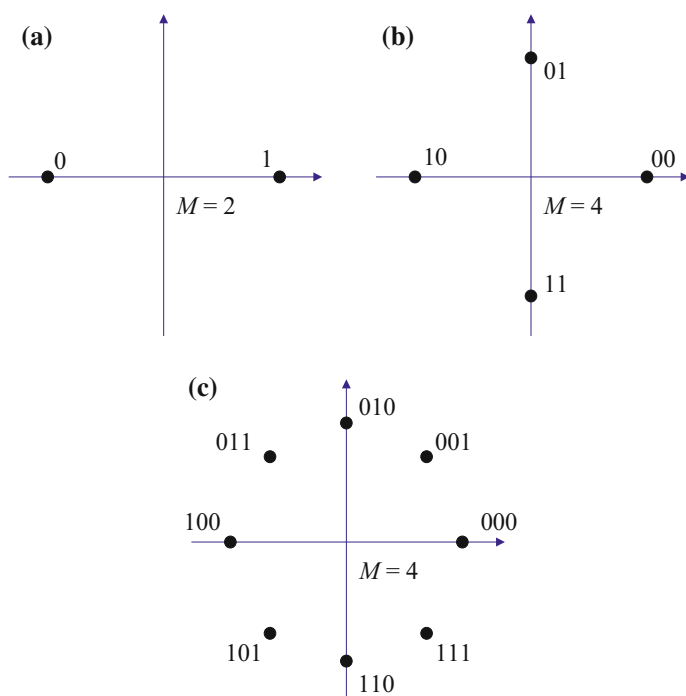


Fig. 2.7 Signal constellations of **a** BPSK, **b** QPSK and **c** octal PSK

2.4.3 Frequency Shift-Keying

If the digital information is modulated onto the carrier by changing its frequency (e.g. by deploying two LOs), then FSK is accomplished. Normally, it is not practical to use more than two frequencies to represent 0 and 1, which results in binary FSK. An FSK-modulated signal is illustrated in Fig. 2.8.

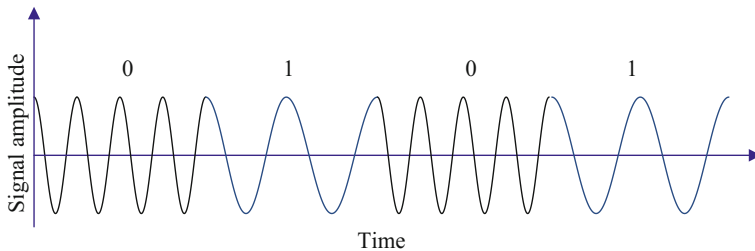


Fig. 2.8 Binary FSK-modulated signal waveform

2.4.4 Pulse-Amplitude Modulation

If the digital signal is encoded as the amplitude in series of pulses, then PAM is achieved. M signal waveforms can be represented as

$$s_m(t) = A_m g(t) \cos 2\pi f_c t \quad (2.3)$$

for $m = 1, 2, \dots, M$, $0 \leq t \leq T$, where A_m is the signal amplitude that can take discrete levels

$$A_m = (2m - 1 - M)d. \quad (2.4)$$

In the previous equation, the value of $2d$ is defined as the distance on the amplitude axis between two signal amplitudes. PAM is illustrated in Fig. 2.9.

As with PSK, one-bit, two-bit or three-bit symbol combinations are suitable for practical implementation. This corresponds to $M = 2, 4$ and 8 respectively. Constellations for 2-PAM, 4-PAM and 8 PAM are illustrated in Fig. 2.10.

2.4.5 Quadrature Amplitude Modulation

QAM can be used as a technique for bandwidth efficiency improvement. In this modulation scheme, two or more separately modulated signals are combined on the carriers that are out of phase. In essence, this results in a combination of a PSK scheme with another scheme, usually PAM. For example, if PAM containing M_1 amplitude levels is combined with PSK with M_2 phases, the resulting signal constellation will have $M = M_1 M_2$ waveforms. Mathematically, this can be represented by

$$s_m(t) = A_{mc} g(t) \cos 2\pi f_c t - A_{ms} g(t) \sin 2\pi f_c t \quad (2.5)$$

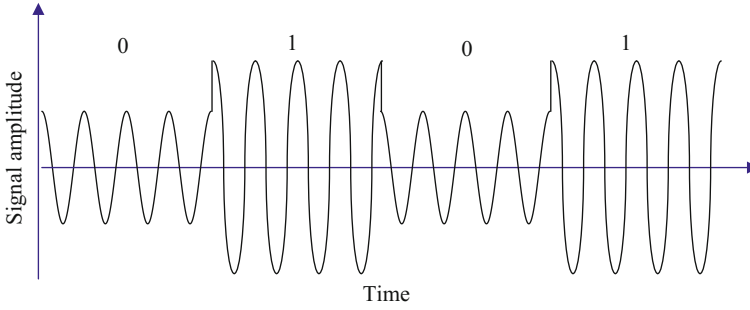
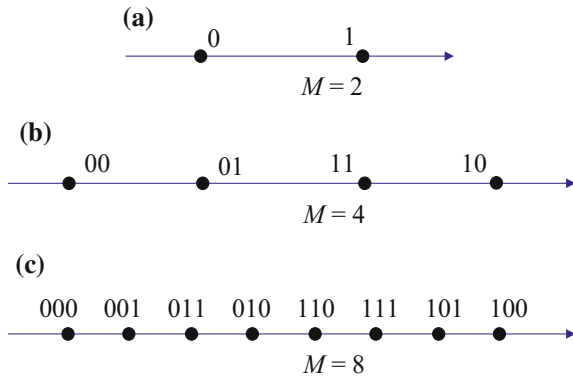


Fig. 2.9 The pulse-amplitude-modulated signal waveform

Fig. 2.10 PAM constellations for **a** $M = 2$, **b** $M = 4$ and **c** $M = 8$



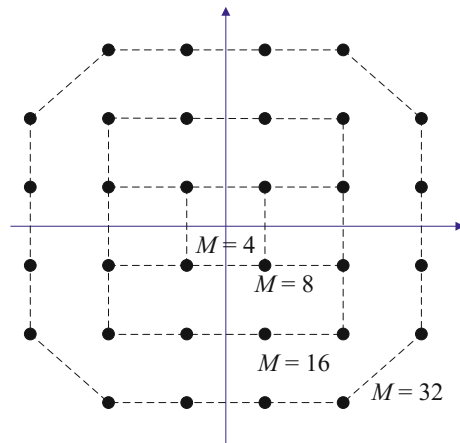
for $m = 1, 2, \dots, M$, $0 \leq t \leq T$, where A_{mc} and A_{ms} are information-bearing signal amplitudes of the quadrature carriers and $g(t)$ is the signal pulse once more. The signal constellations are illustrated in Fig. 2.11.

For example, 16-QAM can achieve 7 Gb/s data transfer rate in the 60-GHz band with 2.16 GHz of bandwidth [14]. This is double the data rate that can be achieved by QPSK, but there are several challenges in realizing 16-QAM direct-conversion transceivers, since the carrier-to-noise ratio requirement is much larger than that of QPSK.

2.4.6 Orthogonal Frequency-Division Multiplexing

OFDM is a technique that is often used to achieve high data transmission rates. This is accomplished by transmitting data over a large number of carriers simultaneously rather than using a single carrier with a high data rate. This scheme is practical for application such as HDTV or LTE networks in RF and microwave frequency ranges, and for high-data-rate transmissions in 60 GHz millimeter-wave band.

Fig. 2.11 QAM constellations for different values of M



OFDM incorporates the same modulation scheme around same-amplitude carriers separated in frequency only enough so that intermodulation products arising from one frequency are negligible at the frequencies of the other carriers. Generating these carrier signals consists of inverse fast Fourier transform (FFT) operations on blocks of M symbols at the transmitter, and they are extracted at the receiver by performing FFTs on blocks comprised of M discrete samples [15–17]. OFDM is particularly prone to nonlinear distortion that is caused in the transmitter.

2.5 Antennas for Millimeter-Waves

In a transmitter, after modulation and power amplification, the amplified signal is passed on to an antenna, which is used to radiate the electromagnetic energy into the channel effectively [18]. Such an antenna is referred to as the transmitting antenna. On the receiving side, receiving antennas are used for receiving the electromagnetic energy from the channel. Fundamentally, an antenna is a bidirectional (reciprocal) device [3]. This means that an identical antenna may be used for both transmit and receive functions. A transmitter antenna radiates spherical waves. At some distance from the antenna, the spherical waves can be approximated with plane waves, which simplifies the antenna analysis, and that region is normally referred to as the far-field of the antenna.

Although antennas are, strictly speaking, passive components, and should therefore be investigated in Chap. 4 of this book, their relation with millimeter-wave propagation requires them to be treated in this section.

2.5.1 General Antenna Theory

When analyzing electronic circuitry, the antenna is typically modeled by a resistor with a value of 50Ω . This is possible because each antenna has a characteristic impedance. The low characteristic impedance value of 50Ω has become standard because it allows a high amount of power to be transferred from the transmitter to the antenna, if the output stage (the power amplifier) is properly matched to the antenna. The same happens on the receiving side, where the aim is to transfer a lot of power to the first block in the receiver, typically an LNA.

Another important concept of antenna design is antenna efficiency. Antenna efficiency can be defined similar to the efficiency of the amplifier, and in the case of the transmitting antenna it represents the ratio of radiated power to the power fed to the antenna [19]:

$$\eta_A = \frac{P_{RAD}}{P_{FED}}. \quad (2.6)$$

The efficiency of the receiving antenna can be defined similarly as the ratio of the absorbed power to the power transferred to the first stage of the receiver.

Also similar to an amplifier, an antenna is typically designed to have a certain amount of gain. Gain of the transmitting antenna in a particular direction is typically denoted by G_T , while that of the receiving antenna is denoted by G_R . With this defined, the power density at the distance r from the antenna can be computed, and is

$$p(r) = G_T \frac{P_T}{4\pi r^2}, \quad (2.7)$$

where P_T is the transmitted power. From this equation, it is evident that the power density decreases quadratically with the distance, and consequently that high gains are needed to transmit over long distances. The amount of power received by the antenna on the receiver side with gain G_R is given by the Friis formula

$$P_R = \frac{P_T G_T G_R \lambda^n}{(4\pi r)^n}, \quad (2.8)$$

where $n = 2$. The above formula is only valid if there is a direct line of sight between the transmitter and receiver. In practice, however, there is more than one propagation path between the transmitter and receiver, which means that n in Eq. (2.8) needs to be adjusted for different path losses. Exponent n typically ranges from 1.7 measured indoor, to 5, measured in suburban areas [20].

The length of an antenna, as discussed before, is related to the signal wavelength. The length becomes very important when considering packaging. If an antenna is placed in air, the length of the antenna is simply equal to the wavelength.

However, if the antenna can be printed in a package, then the length of the antenna scales according to the following equation [21]:

$$l = \frac{\lambda}{\sqrt{\mu_r \epsilon_r}}, \quad (2.9)$$

where μ_r and ϵ_r are the relative permeability and permittivity of the substrate on which the antenna is printed, respectively. This allows antennas to be printed in SoP solutions at microwave, and even in integrated circuit packages. Different substrates allow for antennas with different lengths to be placed in different packages. High-gain antennas can be built by increasing the antenna size to several wavelengths.

An antenna also has radiation characteristics, which are mostly determined by its length and the way in which it is excited. The principle of antenna operation is based on the Ampere-Maxwell's law:

$$\Delta \times \mathbf{H} = \mathbf{J} + \frac{\partial \mathbf{D}}{\partial t}, \quad (2.10)$$

where $\frac{\partial \mathbf{D}}{\partial t}$ is the displacement current (D is the maximum dimension of the antenna), $\mathbf{J}(t)$ is the time varying current density and $\mathbf{H}(t)$ is the time varying magnetic field around the antenna. The radiation pattern is normally plotted separately for the horizontal and vertical dimensions. The antenna pattern exhibits several distinct lobes that peak in different directions. The largest of these lobes is known as the main beam, and the others are referred to as sidelobes.

If a lot of antenna power can be concentrated into one direction, then a directional antenna is built and the concept is referred to as directivity. Highly directional antennas are known as pencil beam antennas and the directivity of the pencil beam antenna can be approximated if the beamwidths in both the horizontal and vertical directions (θ_1 and θ_2 respectively, expressed in rad) are known:

$$D \approx \frac{32,400}{\theta_1 \theta_2}. \quad (2.11)$$

Directivity is thus a dimensionless quantity and is sometimes expressed in dB.

For directional antennas, the gain of the antenna can be expressed in terms of antenna efficiency and directivity:

$$G = \eta_A D. \quad (2.12)$$

2.5.2 Millimeter-Wave Antennas

Various types of antennas are commonly used [8]:

- Wire antennas (dipoles, monopoles and others) with low gains, used mostly at HF to UHF;
- Aperture antennas (open-ended waveguides, horns, reflectors) with moderate to high gains, used in microwave bands; and
- Printed antennas on various substrates (slots, dipoles or microstrip) with high gains and also used in microwave bands.

As discussed before, several antennas can be combined in antenna arrays in order to obtain more directivity and other desirable properties.

Millimeter-wave antennas are typically built by scaling the antennas built for lower than millimeter-wave frequencies [3]. Aperture and printed antennas are both suitable for scaling.

Slot Arrays

Applications that require an antenna with a steerable, directional pattern and a reasonable gain typically utilize slot array antennas. Power delivered to the antenna is thus radiated from the slot openings machined into waveguides into free space. Millimeter-wave antennas typically use substrate-integrated waveguides (SIWs).

Horn Antennas

SIWs can also be used in the design of horn antennas. Horn antennas can achieve high gain and high power by using narrow beamwidths. At millimeter-waves, horn antennas can typically be made more efficient than lower frequency antennas.

Microstrip Antennas

Printed antennas, including microstrip antennas, usually exhibit high losses when used with millimeter-waves. Specialized fabrication techniques are generally required to design well-performing antennas. Although with wavelength scaling these antennas can be used in ICs if only size is a concern, implementation of the antennas is more suitable for packaging solutions for SoP.

Leaky Wave Antennas

Leaky wave antennas use the fact that radiation originates from structures where the first of the higher order modes is known to appear at high frequencies. Leaky waves can be generated in closed waveguides by perturbing the aperture with tapered slots or any type of open aperture. At millimeter-wave frequencies, dielectric rod antennas, non-radiative dielectric guide antennas, tapered slot antennas, partially reflective surface patches and printed log-periodic dipole arrays can all be used.

Dielectric Resonator Antennas

Dielectric resonators yield greater radiation efficiencies than microstrip antennas. These antennas can thus be used in ICs, but the nature of their fabrication process introduces certain complexities.

2.6 High-Frequency Electronics: Practical Two-Port Modeling of Low-Noise Amplifiers

Two-port modeling of an amplifier and consequently LNAs is sometimes necessary as a starting point for amplifier research or design. In LF design, experience has shown that admittance parameters (Y -) proved practical, while at microwave and millimeter-wave frequencies, measuring amplifier response becomes more meaningful if the scattering (S -) parameters are deployed. Sometimes it makes more sense to use impedance (Z -) parameters when describing LNAs, specifically the input impedance of an LNA. Admittance, impedance and scattering parameters all describe the amplifier unambiguously.

2.6.1 Admittance Parameters

Admittance is defined as the reciprocal of the impedance [22]:

$$Y = \frac{1}{Z} = G \pm jB, \quad (2.13)$$

where G is the conductance and B is the susceptance, both expressed in siemens (S). An amplifier as a two-port black box with Y -parameters is shown in Fig. 2.12, where I_1 and V_1 are the input current and voltage respectively and I_2 and V_2 are the output current and voltage respectively, and $[Y]$ denotes the admittance matrix, i.e. the 2-by-2 matrix incorporating all four two-port admittance parameters. The short-circuit Y -parameters are then by definition:

$$y_i = \left. \frac{I_1}{V_1} \right|_{V_2=0}, \quad (2.14)$$

$$y_r = \left. \frac{I_1}{V_2} \right|_{V_1=0}, \quad (2.15)$$

$$y_f = \left. \frac{I_2}{V_1} \right|_{V_2=0}, \quad (2.16)$$

and

$$y_o = \left. \frac{I_2}{V_2} \right|_{V_1=0}, \quad (2.17)$$

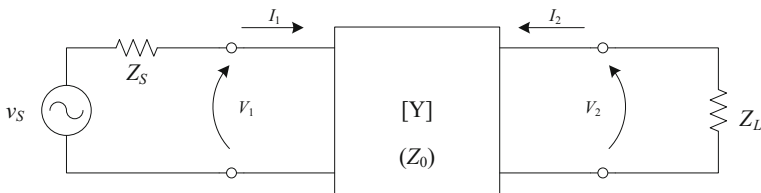


Fig. 2.12 Two port Y -parameter network

where y_i is the short-circuit input admittance, y_r is the short-circuit reverse transfer admittance, y_f is the short-circuit forward transfer admittance and y_o is the short-circuit output admittance. Thus,

$$I_1 = y_i V_1 + y_r V_2, \quad (2.18)$$

and

$$I_2 = y_f V_1 + y_o V_2. \quad (2.19)$$

Z -parameters are easily defined in the same manner, except that the analysis has to be performed in terms of impedances and not admittances.

2.6.2 S -Parameters

S -parameters are much easier to measure and work with than Y - and Z -parameters. S -parameters are also more intuitive than Y - and Z -parameters, since they are the measure of the reflection and gain, as opposed to being a measure of just an abstract quantity such as admittance. The benefits of this approach will become apparent later when gain, stability and matching are discussed. S -parameters can be defined with the aid of Fig. 2.13, where incident and reflected traveling waves are shown in relation to the scattering matrix $[S]$. A traveling wave has the following characteristics [22]:

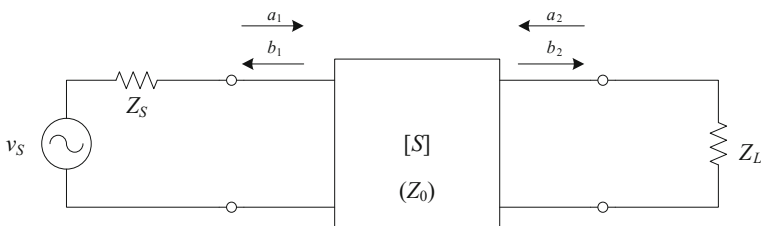


Fig. 2.13 S -parameter two-port model with incident and reflected waves shown

1. A part of the traveling wave originating from the source and incident upon the two-port device (a_1) will be reflected as b_1 and another part will be transmitted through the two-port device;
2. A part of the transmitted signal is reflected from the load and becomes incident upon the output of the two-port device (a_2); and
3. A part of the signal (a_2) is reflected from the output port back toward the load as b_2 and another part is transmitted through the two-port device back to the source.

This requires two reflection coefficients (S_{11} and S_{22}) and two gain coefficients. Thus, the input port reflection coefficient S_{11} is defined as

$$S_{11} = \left. \frac{b_1}{a_1} \right|_{a_2=0}, \quad (2.20)$$

the output port reflection coefficient S_{22} is

$$S_{22} = \left. \frac{b_2}{a_2} \right|_{a_1=0}, \quad (2.21)$$

the forward gain coefficient is

$$S_{21} = \left. \frac{b_2}{a_1} \right|_{a_2=0}, \quad (2.22)$$

and the reverse gain coefficient is

$$S_{12} = \left. \frac{b_1}{a_2} \right|_{a_1=0}. \quad (2.23)$$

Finally, the following two equations can be set

$$b_1 = S_{11}a_1 + S_{12}a_2, \quad (2.24)$$

and

$$b_2 = S_{21}a_1 + S_{22}a_2. \quad (2.25)$$

Conversion between Y -parameters and S -parameters is fairly simple but is typically not required, seeing that only one set of parameters is normally sufficient for research into a particular problem.

2.7 Practical Amplifier Gain Relationships and Stability

Amplifier power gain was already defined in the introductory chapter:

$$G = \frac{P_L}{P_{in}}. \quad (2.26)$$

The gain in Eq. (2.26) is a dimensionless quantity (ratio), but in many instances, it is useful to convert the gain into decibels, or dB, where $G(\text{dB}) = 10 \log G$. Power quantities then need to be expressed in dBm or dBW, that is, power in decibels normalized to 1 mW or 1 W respectively.

The power gain can be the result of voltage gain, current gain or both. Typically, voltage gain is associated with LNAs, and it is simply the ratio of the voltage delivered to the load from the input voltage:

$$A_v = \frac{v_L}{v_{in}}. \quad (2.27)$$

This section introduces two more gain definitions that are used in addition to the power gain relationship defined in Eq. (2.26); these are the available gain G_A and transducer gain G_T . Furthermore, amplifier stability is closely related to gain and the two concepts are usually treated together [23]. Practical gain definitions, however, require familiarity with reflection coefficients and thus the concept of the reflection coefficient is introduced first.

2.7.1 Reflection Coefficients

A two-port LNA model showing the scattering matrix and reflection coefficients is shown in Fig. 2.14. Symbols Z_S , Z_0 and Z_L denote the source impedance, two-port network characteristic impedance and load impedance, respectively, and v_S represents the source voltage.

The reflection coefficient is a parameter that describes how much of an electromagnetic power is reflected by an impedance discontinuity, and thus it can be

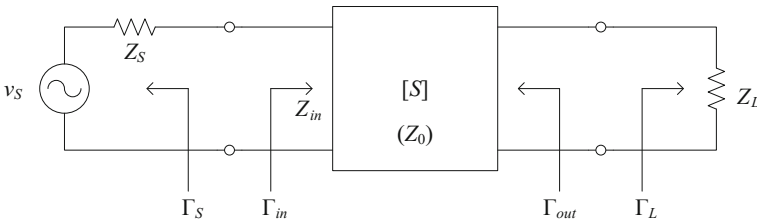


Fig. 2.14 Two-port model of an LNA showing the scattering matrix and reflection coefficients

expressed in terms of impedances. In Fig. 2.14, the reflection coefficient seen looking at the source is defined as:

$$\Gamma_S = \frac{Z_S - Z_0}{Z_S + Z_0}. \quad (2.28)$$

Similarly, the reflection coefficient seen looking at the load is defined as

$$\Gamma_L = \frac{Z_L - Z_0}{Z_L + Z_0}. \quad (2.29)$$

To compute the input and output reflection coefficients, the S -parameters of the system also need to be known. The input reflection coefficient is dependent on the reflection at the load, which gives

$$\Gamma_{in} = S_{11} + \frac{S_{12}S_{21}\Gamma_L}{1 - S_{22}\Gamma_L}. \quad (2.30)$$

Similarly, the output reflection coefficient is dependent on the reflection at the source:

$$\Gamma_{out} = S_{22} + \frac{S_{12}S_{21}\Gamma_S}{1 - S_{11}\Gamma_S}. \quad (2.31)$$

2.7.2 Gain Relationships

The available power gain describes the relationship between the power presented to the network P_N and the power available from the source P_A , and is computed as

$$G_A = \frac{P_N}{P_A}. \quad (2.32)$$

The transducer power gain G_T is defined as the ratio between the power delivered to the load and the power available from the source:

$$G_T = \frac{P_L}{P_A}. \quad (2.33)$$

The gain of the complete amplifier network is maximized when the source and load networks are both conjugately matched to the inputs and outputs of the two-port network, resulting in $G = G_A = G_T$.

Firstly, to compute the power gain, the input power P_{in} and output power delivered to the load P_L need to be known. It can be shown that

$$P_{in} = \frac{|V_S|^2}{8Z_0} \frac{|1 - \Gamma_S|^2}{|1 - \Gamma_{in}\Gamma_S|^2} (1 - |\Gamma_{in}|^2). \quad (2.34)$$

It should be noted here that the power relations are written in terms of the source voltage v_s , which is independent of the impedances at the load or the input to the amplifier. This causes power relation equations to be voltage-specific. Similarly, power to the load can be computed as

$$P_L = \frac{|V_S|^2 |S_{21}|^2 (1 - |\Gamma_L|^2) |1 - \Gamma_S|^2}{8Z_0 |1 - S_{22}\Gamma_L|^2 |1 - \Gamma_S\Gamma_{in}|^2}. \quad (2.35)$$

Finally, the power gain, defined in Eq. (2.26) can be written as

$$G = \frac{P_L}{P_{in}} = \frac{|S_{21}|^2 (1 - |\Gamma_L|^2)}{|1 - S_{22}\Gamma_L|^2 (1 - |\Gamma_{in}|^2)}. \quad (2.36)$$

To compute the available power gain, it is required to obtain the expression for P_A , as well as the expression for P_N . The power available from the source is the maximum power that can be delivered to the network, which is achieved when the input impedance to the amplifier (Z_{in}) is conjugately matched to the impedance of the source, and consequently, when $\Gamma_{in} = \Gamma_S^*$. Therefore, P_A can be calculated from Eq. (2.34) as

$$P_A = P_{in}|_{\Gamma_{in}=\Gamma_S^*} = \frac{|V_S|^2 |1 - \Gamma_S|^2}{8Z_0 |1 - |\Gamma_S|^2|}. \quad (2.37)$$

Similarly, P_N can be calculated from Eq. (2.35) when $\Gamma_L = \Gamma_{out}^*$:

$$P_N = P_L|_{\Gamma_L=\Gamma_{out}^*} = \frac{|V_S|^2 |S_{21}|^2 (1 - |\Gamma_L|^2) |1 - \Gamma_S|^2}{8Z_0 |1 - S_{22}\Gamma_L|^2 |1 - \Gamma_S\Gamma_{in}|^2} \Big|_{\Gamma_L=\Gamma_{out}^*} \quad (2.38)$$

It can be shown that [8, 23]

$$|1 - \Gamma_S\Gamma_{in}|^2 \Big|_{\Gamma_L=\Gamma_{out}^*} = \frac{|1 - \Gamma_S S_{11}|^2 |1 - \Gamma_{out}|^2}{|1 - S_{22}\Gamma_{out}^*|^2}, \quad (2.39)$$

resulting in

$$P_N = \frac{|V_S|^2}{8Z_0} \frac{|S_{21}|^2 |1 - \Gamma_S|^2}{|1 - S_{11}\Gamma_S|^2 (1 - |\Gamma_{out}|^2)}. \quad (2.40)$$

Thus:

$$G_A = \frac{P_N}{P_A} = \frac{|S_{21}|^2 (1 - |\Gamma_S|^2)}{|1 - S_{11}\Gamma_S|^2 (1 - |\Gamma_{out}|^2)}. \quad (2.41)$$

Finally, the transducer gain becomes

$$G_T = \frac{P_L}{P_A} = \frac{|S_{21}|^2 (1 - |\Gamma_S|^2) (1 - |\Gamma_L|^2)}{|1 - S_{22}\Gamma_L|^2 |1 - \Gamma_S\Gamma_{in}|^2}. \quad (2.42)$$

In some special cases, the input and output sections of the amplifier are matched for zero reflection, as opposed to conjugate matching.

This results in

$$\Gamma_L = \Gamma_S = 0. \quad (2.43)$$

As a result, the transducer gain equation simplifies to

$$G_T = |S_{21}|^2 \quad (2.44)$$

while Γ_{in} becomes

$$\Gamma_{in} = S_{11} \quad (2.45)$$

and Γ_{out} becomes

$$\Gamma_{out} = S_{22}. \quad (2.46)$$

Parameter S_{12} is, similarly, related to the concept of reverse isolation described in Chap. 1. The last few expressions therefore explain the official naming of S -parameters as discussed earlier.

2.7.3 Amplifier Stability

Depending on frequency and termination, a two-port amplifier can become unstable and begin to oscillate. Therefore, any amplifier must also meet stability conditions in the frequency range of interest.

If a new quantity Δ is defined in terms of S -parameters as

$$\Delta = S_{11}S_{22} - S_{12}S_{21} \quad (2.47)$$

then, after some manipulation, Γ_{in} and Γ_{out} can be expressed as

$$\Gamma_{in} = \frac{S_{11} - \Gamma_L \Delta}{1 - S_{22} \Gamma_L} \quad (2.48)$$

and

$$\Gamma_{out} = \frac{S_{22} - \Gamma_S \Delta}{1 - S_{11} \Gamma_S} \quad (2.49)$$

respectively.

Conditional stability implies that the magnitudes of all reflection coefficients are less than unity. In other words,

$$|\Gamma_L| < 1, |\Gamma_S| < 1, |\Gamma_{in}| < 1, |\Gamma_{out}| < 1. \quad (2.50)$$

The theory of stability circles can assist in determining unconditional stability. The system will be unconditionally stable if

$$k = \frac{1 - |S_{11}|^2 - |S_{22}|^2 + |\Delta|^2}{2|S_{12}||S_{21}|} > 1 \quad (2.51)$$

and

$$|\Delta| < 1. \quad (2.52)$$

Quantity k introduced in Eq. (2.51) is called the stability or Rollett factor.

In packaged devices, package parasitics need to be included in the stability measurement or calculation. In practice, LNAs experience linear stability phenomena, therefore the non-linear stability issues are sometimes not analyzed [24]. On the other hand, in multi-stage amplifier designs, it may not be sufficient just to confirm the global stability; one might also need to confirm the stability of various LNA stages. This can be done through simulations and measurement; in the case of the latter, probe points on the actual circuit need to be inserted.

2.8 Impedance Matching

In this section, impedance matching for maximum power transfer will be discussed, following the introductory impedance matching discussion in Chap. 1. Figure 2.15 shows a block diagram of an LNA illustrating matching on the input and output side. At millimeter-wave frequencies, where wavelengths are correspondingly small, matching can be accomplished with transmission lines [2]. Matching using lumped components can be deployed as well, once again noting the limitations discussed in Chap. 1, and later in Chap. 4.

2.8.1 Lumped Element Matching

In case of lumped component matching, two-component networks (L networks) and three-component networks (T and Π networks) are commonly used. Eight L-network configurations are possible, as shown in Fig. 2.16a and b, where X_1 and X_2 (where X is the reactance of a component) can be any combination of inductors and capacitors, Z_S is the source impedance and Z_L is the load impedance. Such an L network is a broadband (either high-pass or low-pass) network. Conversely, the T and Π networks with passives X_1 , X_2 and X_3 , shown in Fig. 2.17a and b are narrowband networks.

For maximum power transfer matching, any type of network can give a perfect match (zero reflection) at a single frequency, but multiple-element networks are fairly complex to analyze analytically, thus multiple-element tuning is usually performed graphically with the aid of Smith charts, or using computer-based

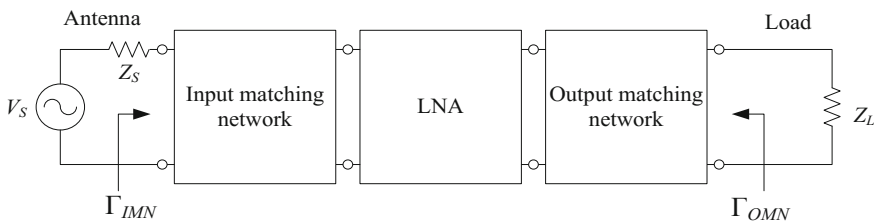
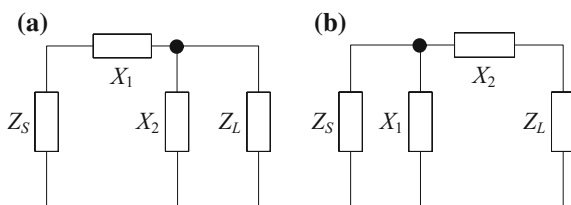


Fig. 2.15 Block diagram of an LNA showing input and output matching networks

Fig. 2.16 Two-component matching networks where passive component is parallel to **a** load and **b** source [23]



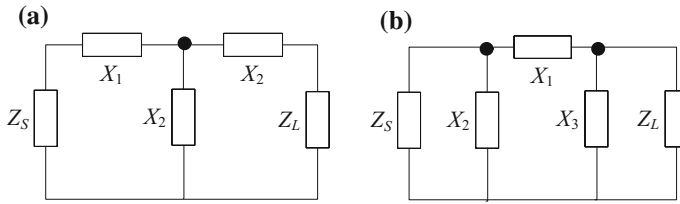


Fig. 2.17 Three-component matching networks: **a** T network and **b** Π network [23]

analysis. The main advantage of introducing multiple-element networks over two-element networks is that with these networks, which can be understood as a combination of several simpler networks, different bandwidths can be achieved.

As seen when the amplifier gain was discussed, for RF, microwave or millimeter-wave circuits, the maximum power transfer theorem states that the maximum power is transferred when the load impedance is equal to the complex conjugate of the source impedance. So, if the source impedance is $Z_S = R_S + jX_S$, the load needs to be made to look like $Z_S^* = R_S - jX_S$, resulting in the process called conjugate matching. When the source and load are both real quantities ($Z_S = R_S$, $Z_L = R_L$), an analytical solution to the matching problem is fairly simple.

For the simplest matching network (L-network) where each network consists of one capacitor and one inductor, one in shunt (parallel) and the other one in series with either the source or the load, as illustrated in Fig. 2.18, the matching procedure is as described below. If we define Q_s as the Q-factor of the series element and Q_p as the Q-factor of the shunt respectively, with

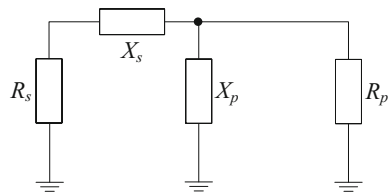
$$Q_s = \frac{X_s}{R_s} \quad (2.53)$$

and

$$Q_p = \frac{R_p}{X_p}, \quad (2.54)$$

where R_p is the parallel resistance and R_s is the series resistance, then a generic matching circuit can be used. Then, for $R_p > R_s$,

Fig. 2.18 Generic matching circuit with source and load resistances R_S and R_L replaced with series and parallel (shunt) reactances



$$Q_s = Q_p = \sqrt{\frac{R_p}{R_s}} - 1 \quad (2.55)$$

is a design equation that can be used to calculate the required Q-factor of the whole network, and from there, the reactances X_s and X_p of each of the two matching elements. The result of Eq. (2.55) implies that the designer has no control over the Q-factor of the matching network. If a precise Q-factor value is required, a multi-element network has to be used.

The condition $R_p > R_s$ means that for a solution to be possible, a parallel component needs to be placed next to the larger of the two resistance values (R_s or R_L).

Reactances are calculated with the aid of the following two equations. If the element is an inductor, then

$$L = \frac{X}{2\pi f_0} \quad (2.56)$$

and if the element is a capacitor,

$$C = \frac{1}{2\pi f_0 X}, \quad (2.57)$$

where f_0 is the operating frequency.

Maximum power transfer matching of complex lumped L-networks, as well as real or complex T or Π -networks discussed earlier, or networks involving even more elements, becomes progressively more complicated, and is done either graphically or with the help of EDA and will not be discussed further in this book.

2.8.2 Transmission-Line Matching

Another way to match impedances in an LNA system is with the aid of transmission lines. As with lumped-element matching networks, several variations of transmission-line matching networks are available. Typically, microstrip lines or waveguides are used. Microstrip transmission lines are characterized by the characteristic impedance (typically Z_0) and the transmission line length l , which will be discussed in Chap. 4. These parameters therefore allow for at least the following matching options:

- Matching networks based on parallel single-stub microstrip lines with fixed or different characteristic impedances;
- Matching networks based on double parallel stub microstrip lines;
- Matching based on a series stub transmission line (only twin conductor transmission lines); and

- Matching networks involving a quarter-wavelength transformer.

Matching networks with a combination of lumped elements and transmission lines are also sometimes deployed. Figures 2.19, 2.20 and 2.21 show some examples of matching networks.

A special type of matching network can be created if the length of the microstrip line is one quarter of the wavelength, as illustrated in Fig. 2.22. In that case, a quarter-wave transformer is created. The input impedance of a quarter wave transformer with characteristic impedance Z_1 , connected to another transmission line with characteristic impedance Z_0 , will be shown in Chap. 4 to be $Z_{in} = Z_1^2/Z_L$.

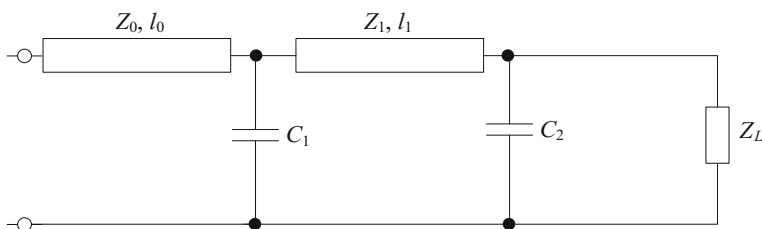


Fig. 2.19 Matching networks with a combination of transmission lines and lumped elements (in this case capacitors)

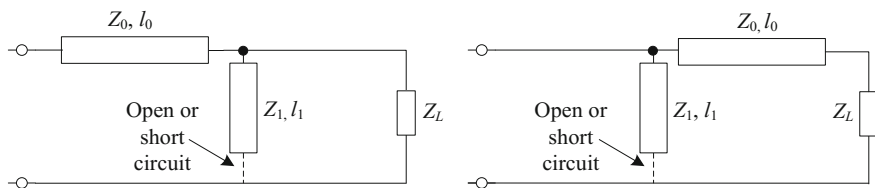


Fig. 2.20 Two variations of single-stub matching networks

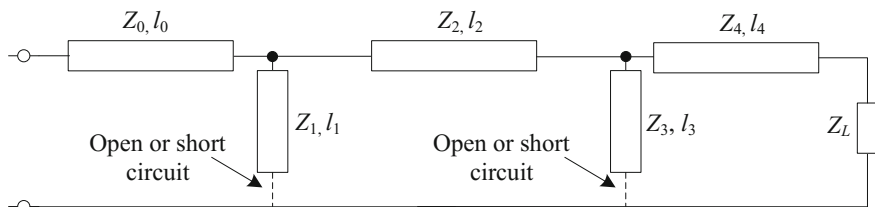
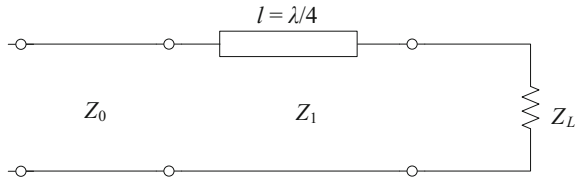


Fig. 2.21 A double-stub matching network

Fig. 2.22 A quarter-wave transformer

2.8.3 Matching and Constant Voltage Standing Wave Ratio

Another system specification that should be addressed by the impedance networks is a voltage standing wave ratio (VSWR) restriction at the input and output ports of the amplifier. The VSWR specification would typically come from the antenna or the circuit following the LNA.

If the reflection coefficient looking into the matching network from the source side is Γ_{IMN} , the input VSWR can be computed with

$$VSWR_{IMN} = \frac{1 + |\Gamma_{IMN}|}{1 - |\Gamma_{IMN}|}. \quad (2.58)$$

Similarly, if the reflection coefficient looking into the matching network from the source side is Γ_{OUT} , the output VSWR can be computed with

$$VSWR_{OMN} = \frac{1 + |\Gamma_{OMN}|}{1 - |\Gamma_{OMN}|}. \quad (2.59)$$

2.9 Biasing

Although strictly speaking, biasing is not a two-port concept, it is very closely related to impedance matching. The input impedance matching, mentioned in the previous section, ensures that the correct AC signals appear at the input of an LNA. In addition to matching, biasing at the input is important to set the correct DC operating point. Therefore, biasing provides the appropriate quiescent point for the LNA [23].

A well-designed biasing network will ensure that the quiescent point of each transistor in an LNA remains relatively constant despite parameter variations or temperature fluctuations. Active and passive biasing networks are possible. Figure 2.23 shows one-resistor and three-resistor biasing networks commonly used with MOS LNAs.

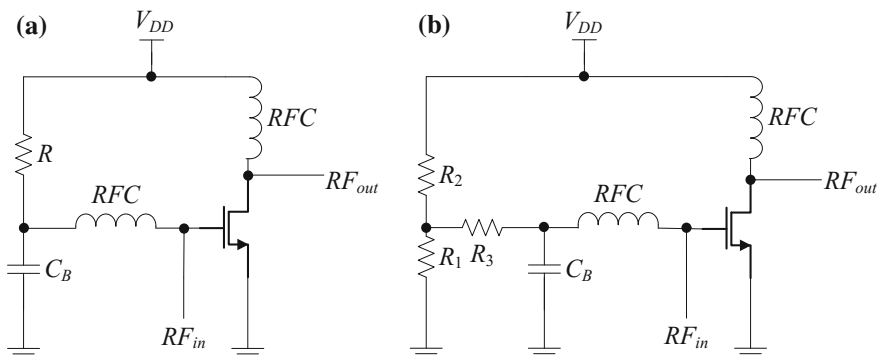


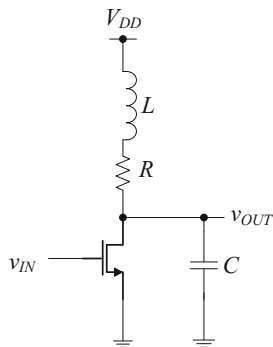
Fig. 2.23 Passive biasing networks for a MOS: **a** one-resistor configuration, **b** three-resistor configuration [23]

2.10 Broadband Amplifier Techniques

In Chap. 1 it was discussed that increasing the bandwidth of an LNA increases its reusability in numerous applications. As will be seen in later chapters of this book, LNAs for millimeter-wave applications with bandwidths greater than 20 GHz have been reported [25]. In order for the amplifier to be a truly wideband amplifier, gain needs to remain flat and matching needs to be constant over the whole band of interest. As the required bandwidth increases, this becomes increasingly difficult. The aim of this section is to present some general techniques that are used for bandwidth enhancement of amplifiers, and this topic will be explored in more detail when specific LNA configurations are investigated and when the state-of-the-art configurations are explored later in this book.

Lee [26] presents several amplifier topologies that can be used to increase amplifier bandwidth. One of the amplifier variations is the shunt peaking common-source amplifier depicted in Fig. 2.24. This topology incorporates an inductor L at the drain of the transistor. The configuration can be configured as

Fig. 2.24 A shunt-peaked amplifier



maximum bandwidth, maximally flat frequency response, best group delay or maximum power transfer. For maximum bandwidth, L is chosen as

$$L_1 = \frac{R^2 C}{\sqrt{2}} \quad (2.60)$$

and the bandwidth of the amplifier is increased by a factor of about 1.85. For maximum flatness, $\sqrt{2}$ can be replaced with ~ 3.1 . In this case, the bandwidth is increased only about 1.72 times.

Perhaps the simplest configuration is the configuration that makes use of negative feedback. In Fig. 2.25, feedback resistor R_F turns a common-source amplifier with source degeneration into a shunt-series amplifier. Note that the biasing is not shown in this figure. The name stems from the use of the combination of shunt and series feedback. The gain of this amplifier decreases to

$$A_v = -\frac{R_L}{R_1} \frac{R_F - R_S}{R_F - R_L} \quad (2.61)$$

from roughly $A_v = -R_L/R_1$ associated with a common-emitter amplifier with no feedback, but bandwidth increases to

$$BW = \left[|A_v| \left(\frac{C_{gs}}{g_m} + \frac{RC_{gd}}{2} \right) \right]^{-1} \quad (2.62)$$

where g_m , C_{gs} and C_{gd} are common emitter small signal gain and parasitic capacitances, all of which will be discussed in Chap. 3 of this book.

Another technique for increasing the bandwidth is to replace a single-ended amplifier with a differential amplifier, such as the amplifier shown in Fig. 2.26. A differential amplifier is already a good choice in many amplifier configurations, because of good common mode noise rejection. The increase in bandwidth can be understood in terms of transitional frequency f_T . In Chap. 3, the f_T of a transistor will be defined as

$$f_T = \frac{1}{2\pi} \frac{g_m}{C_{gs} + C_{gd}}. \quad (2.63)$$

Fig. 2.25 Shunt-series amplifier

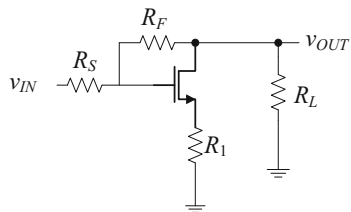
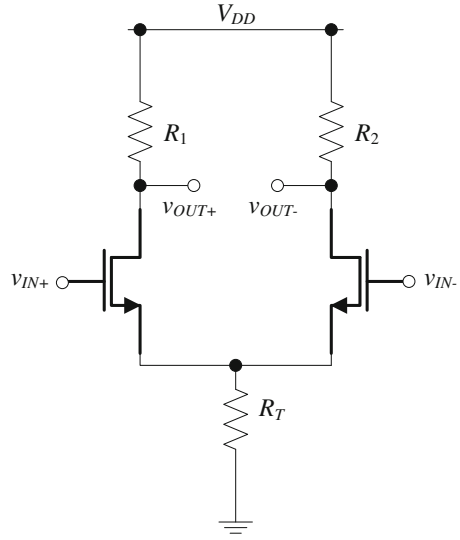


Fig. 2.26 Differential amplifier circuit



In a differential configuration, the device capacitances C_{gs} and C_{gd} of each transistor in the differential amplifier appear in series. For perfectly matched transistors,

$$(C_{gs} + C_{gd})' = \frac{1}{2}(C_{gs} + C_{gd}) \quad (2.64)$$

and thus

$$f_T' = 2f_T, \quad (2.65)$$

and the differential amplifier is therefore an f_T doubler. As a result, any amplifier built around transistors with doubled f_T will have somewhat higher bandwidth.

More complex amplifier topologies other than the ones described up to now are also used. A balanced amplifier implemented with 90° hybrid couplers is described here. This is shown in Fig. 2.27 [3]. The two 90° hybrid couplers in the balanced amplifier serve to cancel out reflections from the amplifier input ports, which leads to an improved impedance match. Furthermore, the bandwidth of the overall amplifier system is primarily determined by the coupler bandwidth, meaning that these can be optimized individually to increase the overall bandwidth.

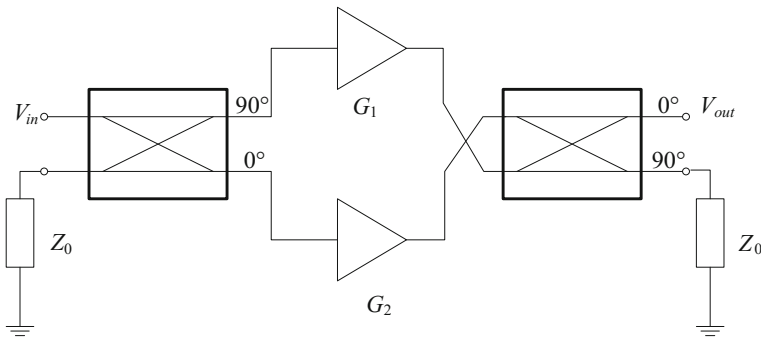


Fig. 2.27 A balanced amplifier

2.11 Narrowband Amplifier Techniques

The alternative to a broadband amplifier is the narrowband amplifiers. Narrowband amplifiers are achieved by introducing a resonant tank, which is used to pass a frequency or a group of frequencies selectively while attenuating all other frequencies [19, 22, 27]. In LNAs, the same concept can be used to resonate out the unwanted parasitic components, e.g. an inductor can be placed in parallel or in series to a parasitic capacitance.

A circuit diagram of a parallel resonant circuit is shown in Fig. 2.28.

The reactance of capacitor C at any frequency is

$$X_C = \frac{1}{\omega C}, \quad (2.66)$$

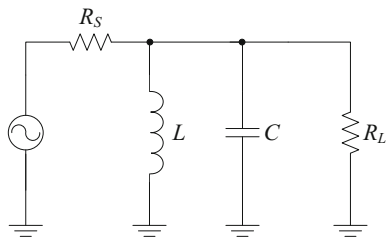
where $\omega = 2\pi f$ is the angular frequency. The reactance of the inductor is

$$X_L = \omega L. \quad (2.67)$$

Resonance is achieved at the frequency where both reactances are equal:

$$X_C = X_L = \frac{1}{\omega C} = \omega L. \quad (2.68)$$

Fig. 2.28 Schematic of a parallel resonant circuit



The resonant frequency is then

$$f_o = \frac{1}{2\pi\sqrt{LC}}. \quad (2.69)$$

An important property of a resonant tank is that it ideally does not add any attenuation to the signal at the frequency of resonance. This can be proven by assuming that the combined reactance of the capacitor and inductor close to resonance is sufficiently small to allow load resistance in parallel to be ignored. The magnitude of the output voltage of the resonant circuit, V_{out} , in terms of the magnitude input voltage V_{in} is

$$V_{out} = \frac{Z_C || Z_L}{R_S + Z_C || Z_L} V_{in}, \quad (2.70)$$

where $Z = jX$ is the impedance of a capacitor or inductor. Thus, the attenuation at any frequency ω is

$$\frac{V_{out}}{V_{in}} = \left| \frac{j\omega L}{R_S - \omega^2 R_S LC + j\omega L} \right|. \quad (2.71)$$

At ω_o ,

$$\frac{V_{out}}{V_{in}} = 1. \quad (2.72)$$

2.12 Noise in Amplifiers

In this section, the amplifier noise theory is expanded in the light of the two-port modeling described earlier in this chapter. Transistor noise modeling is deferred until Chap. 3.

2.12.1 Noise Figure

If S_i and N_i can be defined as signal and noise power respectively at the input of the amplifier, and S_o and N_o can be defined as signal and noise power respectively at the output of the amplifier, such as depicted in Fig. 2.29, then the noise factor (noise figure) can be expressed as

Fig. 2.29 Signal and noise power at the input and output of a two-port network



$$F = \frac{S_i/N_i}{S_o/N_o} = \frac{S_i}{N_i} \frac{N_o}{S_o}. \quad (2.73)$$

In noiseless circuits, input and output signal power and input and output noise power are related by the power gain G . From Eq. (2.26),

$$S_o = GS_i \quad (2.74)$$

and

$$N_o = GN_i. \quad (2.75)$$

A non-ideal amplifier will add some noise, such that

$$N_o > GN_i. \quad (2.76)$$

Equation (2.74) also allows for an alternative definition of the noise factor in terms of gain:

$$F = \frac{N_o}{GN_i}, \quad (2.77)$$

which illustrates that F is also the ratio of the total output noise to the part of output noise due to the source resistance (amplified input noise).

2.12.2 Noise Floor

The noise power that is always present at the receiver is known as the noise floor. If the signal is to be detected, the signal power always needs to be greater than the noise power by at least the value of the noise floor to result in $S_o > N_o$ and reliable signal detection. LNAs are therefore designed with a low noise figure, allowing for a lower minimum detectable signal (MDS). The lowest MDS of a receiver with bandwidth B and overall noise figure NF in dB is [20]

$$MDS = -173.83 + NF + 10 \log B. \quad (2.78)$$

The value of -173.83 dB is the noise floor at 290 K.

2.12.3 Noise Temperature

An alternative way to represent the noise is through noise temperature. The noise temperature T_n is defined as the temperature at which the source resistance must be kept to so that the noise at the output due to source resistance is equal to the noise produced by the circuit itself. At temperature T at which the noise temperature is specified, T_n can be computed by

$$T_n = T(F - 1). \quad (2.79)$$

2.12.4 Noise Bandwidth

Noise bandwidth can be defined more formally if a reference frequency is chosen. The reference frequency of the noise passband is normally chosen as the band center frequency. As a result, noise bandwidth is defined as

$$B = \int_0^{\infty} \frac{G(f)}{G_o} df \quad (2.80)$$

where $G(f)$ is the amplifier frequency response and G_o is the gain at the reference frequency.

2.12.5 Minimum Noise Figure and Practical Amplifier Design

Although every network can be designed for a minimum noise figure F_{min} by doing noise matching, some noise matching is naturally sacrificed to ensure that the required gain of the amplifier is reached and that the stability criterion is satisfied. This results in the concept of simultaneous noise and power matching, discussed before. In this case, an alternative noise figure definition may be derived that includes F_{min} and the part of the noise due to noise mismatch. In that case, if the noise is described by means of two-port conductance G_n (or alternatively, resistance $R_n = 1/G_n$) and if the source has conductance G_s , the noise figure will be minimized for some optimal source conductance G_{opt} . A two-port noise amplifier network is represented in Fig. 2.30.

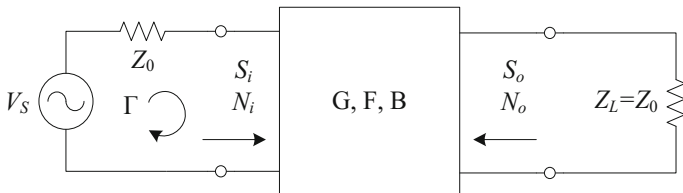


Fig. 2.30 Generic amplifier with an impedance mismatch at its input port

If $G_s = G_{opt}$, then $F = F_{min}$. At any other source admittance, the noise figure can be calculated as

$$F = F_{min} + \frac{R_n}{G_s} (G_s - G_{opt})^2. \quad (2.81)$$

This equation assumes no imaginary components in the source or noise values. A more complete formula would be in terms of admittances, i.e.

$$F = F_{min} + \frac{R_n}{G_s} |Y_s - Y_{opt}|^2 \quad (2.82)$$

where Y_s is the actual source admittance and Y_{opt} is the optimal source admittance.

For millimeter-wave designs, a representation utilizing reflection coefficients may be more applicable. Equation (2.81) can be rewritten as

$$F = F_{min} + \frac{4R_n}{Z_0} \frac{|\Gamma_s - \Gamma_{opt}|^2}{(1 - |\Gamma_s|^2)(1 + |\Gamma_{opt}|^2)} \quad (2.83)$$

where Γ_{opt} is now the optimal source reflection coefficient and Z_0 is the characteristic impedance of the amplifier [28]. Γ_{opt} and Y_{opt} are related with the expression

$$Y_{opt} = Y_0 \frac{1 - \Gamma_{opt}}{1 + \Gamma_{opt}} \quad (2.84)$$

where Y_0 is now the amplifier characteristic admittance.

2.12.6 Input Noise Power

If the input impedance is matched for the minimum noise figure, and the noise bandwidth is known, the input noise power to the amplifier can actually be computed, which simplifies noise calculations. Per definition, the input taken at the

noise power that results from a source resistor that operates at a temperature $T_0 = 290$ K is given by

$$N_i = kT_0B, \quad (2.85)$$

where k is Boltzmann's constant, which is equal to 1.380×10^{-23} J/K. If expressed in decibel and T_0 and k are substituted, the input noise power in dBm is

$$N_i(\text{dBm}) = -174 + 10 \log B. \quad (2.86)$$

2.12.7 Noise Factor in a Cascaded System

If several amplifiers are cascaded, as illustrated in Fig. 2.31, each amplifier will have its own noise figure. The noise factor of the m -th amplifier is thus from Eq. (1.49)

$$F_m = \frac{N_{om}}{G_m N_i}, \quad (2.87)$$

where G_m is the gain of the m -th amplifier and N_i is constant, per definition, for each amplifier stage. If Eq. (2.86) is rewritten for the m -th amplifier as

$$N_{om} = G_m N_i + N_m \quad (2.88)$$

where N_m is the noise added to the amplifier itself, then F_m will be

$$F_m = \frac{G_m N_i + N_m}{G_m N_i}. \quad (2.89)$$

N_m is then

$$N_m = (F_m - 1)G_m N_i. \quad (2.90)$$

For $m = 1$,

$$N_1 = (F_1 - 1)G_1 N_i. \quad (2.91)$$

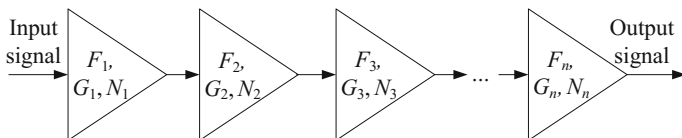


Fig. 2.31 Cascading of amplifier stages

For $m = 2$,

$$N_2 = (F_2 - 1)G_2N_i. \quad (2.92)$$

If the cascaded amplifier is denoted with index c , then the gain of the cascade will be $G_c = G_1G_2$, and the noise will be

$$N_c = (F_c - 1)G_cN_i, \quad (2.93)$$

and also

$$N_c = G_2N_1 + N_2. \quad (2.94)$$

Manipulating the previous four equations results in an equation for the noise factor for a cascaded system

$$F_c = F_1 + \frac{F_2 - 1}{G_1}. \quad (2.95)$$

If the process is repeated n number of times, using induction, the noise factor equation for the n cascaded amplifiers can be worked out to be

$$F_c = F_1 + \frac{F_2 - 1}{G_1} + \frac{F_3 - 1}{G_1G_2} + \dots + \frac{F_n - 1}{G_1G_2 \dots G_{n-1}}. \quad (2.96)$$

This leads to an important observation: the noise figure of the cascaded system will be predominantly determined by the noise figure of the first amplifier stage. In the receiver, the noise figure is thus mostly determined by the LNA, which illustrates the importance of the concept for the low-noise research. Similarly, if an LNA is designed as a multistage amplifier, the noise figure of the first stage will be the one that should be considered.

2.13 Amplifier Linearity

As discussed in Chap. 1, linearity is related to LNA performance and understanding linearity concepts aids in understanding research into LNA linearity improvements.

2.13.1 Harmonic Distortion and Intermodulation Distortion

Harmonic distortion arises owing to the unwanted multiples of the fundamental frequency of the output signal of the particular system appearing in the signal waveforms. If the fundamental frequency is denoted with f_c , then the n -th harmonic is designated as nf_c .

Harmonic distortion is typically quantified by means of total harmonic distortion (THD). THD is the ratio of the sum of the power in all harmonic components to the power contained in the fundamental frequency, expressed as [29]

$$THD = \frac{\sqrt{\sum_{n=2}^{\infty} V_{on}^2}}{V_1}, \quad (2.97)$$

where V_{on} is the root-mean square (RMS) value of the voltage of the n -th harmonic and V_1 is the RMS value of the voltage of the signal at fundamental frequency.

Similar to harmonic distortion is the IMD, where instead of multiples of one frequency appearing at the output, intermodulation products of at least two frequencies f_1 and f_2 appear at the output, i.e. $f_{IMD} = \pm nf_1 \pm mf_2$ [19] for natural values of m and n .

A technique called the two-tone test is used in practice, to measure the IMD. The test is conducted such that at least two sinusoidal waveforms are applied in series to the amplifier. The test result is expressed as the carrier-to-intermodulation ratio (C/I), which should be higher than 30 dBc (dB below carrier). If a two-tone signal

$$v_{in} = A \cos \omega_1 t + B \cos \omega_2 t \quad (2.98)$$

consisting of two signals at frequencies $\omega_1 = 2\pi f_1$ and $\omega_2 = 2\pi f_2$ spaced closely together in the frequency domain is applied to the amplifier that introduces an amount of distortion, then the output waveform can be expressed as a Taylor polynomial with an infinite number of terms:

$$v_{out}(t) = a_0 + a_1 v_{in}(t) + a_2 v_{in}^2(t) + a_3 v_{in}^3(t) + \dots + a_n v_{in}^n(t) + \dots, \quad (2.99)$$

where $a_0, a_1, a_2, \dots, a_n, \dots$ are some amplification coefficients. Typically, only the first four coefficients are used to evaluate so-called third-order intermodulation distortion, or IMD3:

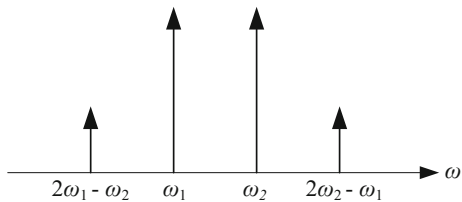
$$v_{out}(t) = a + b v_{in}(t) + c v_{in}^2(t) + d v_{in}^3(t). \quad (2.100)$$

Expanding Eq. (2.100) with Eq. (2.98) results in

$$\begin{aligned} v_{out} = & a + b(A \cos \omega_1 t + B \cos \omega_2 t) \\ & + c(A^2 \cos^2 \omega_1 t + B^2 \cos^2 \omega_2 t + 2AB \cos \omega_1 t \cos \omega_2 t) \\ & + d(A^3 \cos^3 \omega_1 t + A^2 B \cos^2 \omega_1 t \cos \omega_2 t \\ & + AB^2 \cos \omega_1 t \cos^2 \omega_2 t + B^3 \cos^3 \omega_2 t). \end{aligned} \quad (2.101)$$

In this equation, IMD3 terms are terms that give m and n such that $m + n = 3$. By definition then, the terms $dA^2 B \cos^2 \omega_1 t \cos \omega_2 t$ and $dAB^2 \cos \omega_1 t \cos^2 \omega_2 t$

Fig. 2.32 Intermodulation tones



represent the third-order intermodulation terms at frequencies $2\omega_1 - \omega_2$ and $2\omega_2 - \omega_1$, illustrated in Fig. 2.32. This computation also results in unwanted signals at other combinations of frequency sum and difference: $\omega_2 - \omega_1$, $\omega_2 + \omega_1$, $2\omega_1$, $2\omega_2$, $3\omega_1$, $3\omega_2$, $2\omega_1 + \omega_2$, $\omega_1 + 2\omega_2$. In this case, however, these appear much further from either ω_1 or ω_2 , in the frequency domain, and can be filtered out.

If power is associated to each output voltage waveform term in Eq. (2.101), then IMD3 (in dB) can be expressed as the difference in wanted power at f_2 , $P_{out}(f_2)$ and the unwanted power $P_{out}(2f_2 - f_1)$ of the intermodulation term $2f_2 - f_1$:

$$IMD_3 = P_{out}(f_2) - P_{out}(2f_2 - f_1), \quad (2.102)$$

where power quantities are expressed in dBm.

2.13.2 Gain Compression

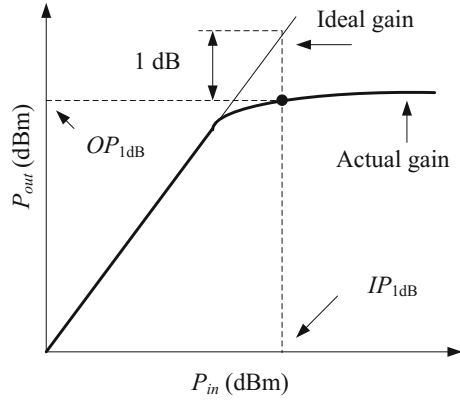
The power gain equation Eq. (2.26), introduced in Chap. 1, indicates a linear relationship between the output power and input power of an amplifier. In reality, this linear relationship only holds for a narrow range and typically, as the amount of input power increases, at a certain stage the output power will stop tracking the input power, i.e. gain value will decrease, or rather, compress. Gain compression is typically specified in terms of the 1 dB compression point, defined as the input power level at which the output power has dropped to 1 dB below the linear characteristic [8]. This effect is illustrated in Fig. 2.33, where both the x and y axes are in dBm.

The 1 dB compression point relationship can formally be expressed as

$$OP_{1dB} = IP_{1dB} + G - 1 \text{ dB}, \quad (2.103)$$

where OP_{1dB} is the output power at 1 dB compression point and IP_{1dB} is the input power at 1 dB compression point, and all parameters are expressed in either dB or dBm, as in the figure.

Fig. 2.233 Illustration of the 1 dB compression point



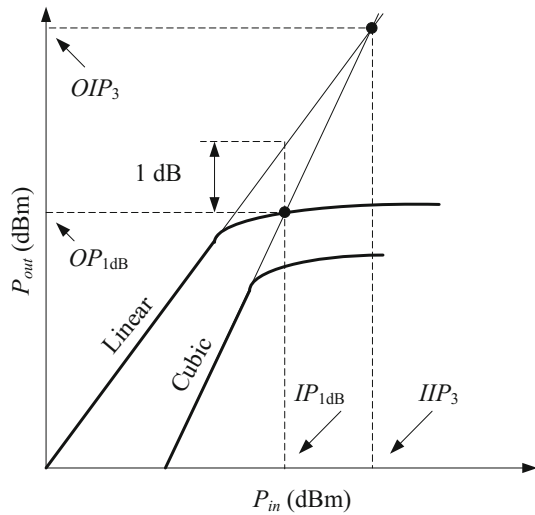
2.13.3 Third Order Intercept Point

The third order intercept point is a linearity measure that ties with both gain compression and intermodulation distortion.

With intermodulation distortion present, it is clear that in a non-ideal amplifier, the power will also be generating terms of the third order (cubed terms). Thus, assuming that all the other terms are attenuated by a filter, the wanted power $P_{out}(f_2)$ and the unwanted power $P_{out}(2f_2 - f_1)$ defined previously can be plotted on the same set of axes against the input power $P_{in}(f_2)$, as illustrated in Fig. 2.34, with both axes once more in dBm.

From Eq. (2.26) the slope of the $P_{out}(f_2)$ versus $P_{in}(f_2)$ curve on the dBm scale will have a value of 1. Since the curve for $P_{out}(2f_2 - f_1)$ versus $P_{in}(f_2)$ contains

Fig. 2.34 Illustration of the third-order intercept in a non-linear device



third-order terms, and $\log(x^3) = 3\log x$, the slope of the curve on the dB scale will have a value of 3, as shown in Fig. 2.34. At low signal levels, the third-order products are negligibly small, meaning that for low P_{in} , $P_{in}(f_2) > P_{out}(2f_2 - f_1)$. As the input power increases, the slope of the $P_{out}(2f_2 - f_1)$ versus $P_{in}(f_2)$ curve will cause it to come closer and closer to the $P_{out}(2f_2)$ versus $P_{in}(f_2)$ curve until the two curves eventually intercept. The intersection of these two curves will mark the third-order intercept point (IP_3). Another intercept point can be derived for f_1 and $2f_1 - f_2$ curves.

In practical amplifiers, because of gain compression (which is typically lower by about 10 to 15 dB than IP_3), this point will remain hypothetical; however, it still makes for a useful linearity metric. Depending on whether the value of the IP_3 is read on the x or the y axis of the graph illustrating the concept, both output (OIP_3) and input (IIP_3) third-order intercept points can be defined.

Unlike the noise factor, the IIP_3 gets better as the frequency increases. Also, the IIP_3 of the last stage of a multi-stage system dominates the overall IIP_3 of the system.

2.13.4 Amplifier Dynamic Range

The amplifier dynamic range is another measure of linearity and is generally defined as the operating range over which a particular component exhibits desirable performance. It typically refers to the linear portion of the power amplification curve, that is, the top limit is taken as the point at which the intermodulation distortion becomes intolerable. In an LNA, often operating close to the noise floor, its lower limit can be established by the noise. This leads to the definition of the spurious-free dynamic range (SFDR), which is the range for which the spurious responses are minimal. This is illustrated in Fig. 2.35.

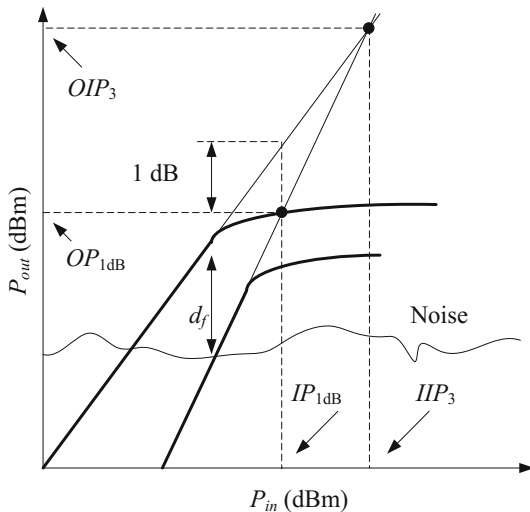
The SFDR, d_f , in terms of output third-order intercept point OIP_3 and noise N_o (dBm) is

$$d_f = \frac{P_{f_2}}{P_{2f_2-f_1}} \bigg|_{P_{2f_2-f_1}=N_o} = \left(\frac{IP_3}{N_o} \right)^{\frac{2}{3}} \quad (2.104)$$

and if both quantities in the equation are in dB or dBm,

$$d_f = \frac{2}{3} (OIP_3 - N_o). \quad (2.105)$$

Fig. 2.35 Spurious-free dynamic range illustration



2.14 Performance Measure of a Low-Noise Amplifier

Instead of looking at the parameters such as noise figure, gain and bandwidth of an LNA separately, the overall performance of an LNA can be assessed by a single figure of merit (*FOM*) [30]. The LNA figure of merit can be expressed in terms of the GBP, the magnitude of the noise factor F_{mag} and the DC power consumption P_{dc} . Thus,

$$FOM = \frac{GBP}{(F_{mag} - 1) \cdot P_{dc}}, \quad (2.106)$$

where GBP can be expressed in terms of the magnitude of the forward gain S_{21} and bandwidth B :

$$GBP = S_{21,mag} \cdot B \quad (2.107)$$

and all the quantities are expressed as ratios and not as dB values.

2.15 Concluding Remarks

This chapter served as a first step in contextualizing the research into LNAs for millimeter-wave applications covered in this book. It covered several important aspects that need to be considered in the design of an LNA, such as the band in which the LNA is to operate and therefore the design frequency of operation and

signal wavelength. The wavelength in turn influences the antenna design and also the propagation characteristics of the signal. All of these influence the quality of the signal reaching the LNA. The modulation scheme deployed at the transmitter also determines the type of signals that will reach the LNA and thus the types of modulation schemes were also discussed to place the research in the telecommunication context.

To simplify the connection of the subsystems, such as the antenna to the LNA or LNA to the demodulator, two-port modeling is often used. Thus, in this chapter, two-port modeling was discussed and LNA design parameters, such as gain and the noise figure, were discussed once again from the two-port perspective. Design specifications also govern parameters describing the linearity of LNA, and the final portion of this chapter was dedicated to linearity and distortion.

In addition to the specifications posed by the application, some of the LNA research is initiated by the active device technology that is available to the researcher. Understanding fabrication technologies is therefore of crucial importance for research into LNAs. This is discussed in Chap. 3.

References

1. Rappaport TS, Murdock JN, Gutierrez F (2011) State of the art in 60-GHz integrated circuits and systems for wireless communications. *Proc IEEE* 99(8):1390–1436
2. International Telecommunication Union (2000) Nomenclature of the frequency and wavelength bands used in telecommunications. ITU-R Recommendation V.431 (internet). Available from: <http://www.itu.int/rec/R-REC-V.431/en>. Cited 19 May 2015
3. du Preez J, Sinha S (2016) Millimeter-wave antennas: configurations and applications. Springer, Berlin
4. Baykas T, Sum CS, Lan Z, Wang J, Rahman MA, Harada H, Kato S (2011) IEEE, 802.15. 3c: the first IEEE wireless standard for data rates over 1 Gb/s. *IEEE Commun Mag* 49(7):114–121
5. Perahia E, Cordeiro C, Park M, Yang LL. IEEE 802.11 ad: defining the next generation multi-Gbps Wi-Fi. In: 2010 7th IEEE consumer communications and networking conference; 2010; Las Vegas, pp 1–5
6. Hsiao YH, Chang YC, Tsai CH, Huang TY, Aloui S, Huang DJ, Chen YH, Tsai PH, Kao JC et al (2016) A 77-GHz 2T6R transceiver with injection-lock frequency sextupler using 65-nm CMOS for automotive radar system application. *IEEE Trans Microw Theory Tech* 64(10):3031–3048
7. Hasch J, Topak E, Schnabel R, Zwick T, Weigel R, Waldschmidt C (2012) Millimeter-wave technology for automotive radar sensors in the 77 GHz frequency band. *IEEE Trans Microw Theory Tech* 60(3):845–860
8. Pozar M (2012) *Microwave engineering*, 4th edn. Wiley, Hoboken
9. Adhikari P (2008) Understanding millimeter wave wireless communication. White paper. Loea Corporation
10. du Preez J, Sinha S (2017) *Millimeter-Wave Power Amplifiers*. Springer, Cham
11. Proakis JG, Salehi M (2008) *Digital communications*, 5th edn. McGraw-Hill, New York
12. Raab FH, Asbeck P, Kenington PB, Cripps S, Popovic ZB, Potheary N, Sevic JF, Sokal NO (2003) RF and microwave power amplifier and transmitter technologies—Part I. *High Freq. Electron.* 2:22–36

13. Sinha S, Božanić M, Schoeman J, du Plessis M, Linde LP. A CMOS based multiple-access DSSS transceiver. In: 2009 South African conference on semi and superconductor technology; 2009; Stellenbosch, pp 19–24
14. Okada K, Li N, Matsushita K, Bunsen K, Murakami R, Musa A, Sato T, Asada H, Takayama N, Ito S, Chaivipas W (2011) A 60-GHz 16QAM/8PSK/QPSK/BPSK Direct-Conversion Transceiver for IEEE802.15.3c. *IEEE J Solid State Circ* 46(12):2988–3004
15. Cimini L (1985) Analysis and simulation of a digital mobile channel using orthogonal frequency division multiplexing. *IEEE Trans Commun* 33(7):665–675
16. Falconer D, Ariyavisitakul SL, Benyamin-Seeyar A, Eidson B (2002) Frequency domain equalization for single-carrier broadband wireless systems. *IEEE Commun Mag* 40(4):58–66
17. Thompson SC, Ahmed AU, Proakis JG, Zeidler JR, Geile MJ (2008) Constant envelope OFDM. *IEEE Trans Commun* 56(8):1300–1312
18. Cheng DK (1993) *Fundamentals of engineering electromagnetics*, 1st edn. Reading: Addison-Wesley Publishing Company
19. Kazimierczuk MK (2015) *RF Power amplifiers*, 2nd edn. Wiley, Chichester
20. Robertson I, Somjit N, Chongcheawchamnan M (2016) *Microwave and millimetre-wave design for wireless communications*, 1st edn. Wiley, Chichester
21. Tummala RR, Swaminathan M (2008) *System-on-package: miniaturization of the entire system*, 1st edn. McGraw-Hill Professional, New York
22. Bowick C, Blyler J, Ajluni C (2008) *RF circuit design*, 2nd edn. Elsevier, Burlington
23. Ludwig R, Bretchko P (2000) *RF circuit design: theory and applications*, 1st edn. Prentice Hall, Upper Saddle River
24. Pantoli L, Barigelli A, Leuzzi G, Vitulli F (2014) Analysis and design of a Q/V-band low-noise amplifier in GaAs-based 0.1 μm pHEMT technology. *IET Microw Antennas Propag* 10(14):1500–1506
25. Feng G, Boon CC, Meng F, Yi X, Li C (2016) An 88.5–110 GHz CMOS low-noise amplifier for millimeter-wave imaging applications. *IEEE Microw Wirel Compon Lett* 26(2):134–136
26. Lee TH (2004) *The design of CMOS radio-frequency integrated circuits*, 2nd edn. Cambridge University Press, Cambridge
27. Grebennikov A, Sokal NO, Franco MJ (2012) *Switchmode RF and microwave power amplifiers*, 2nd edn. Elsevier, Burlington
28. Agilent (2010) *Fundamentals of RF and microwave noise figure measurement*. Application note. Agilent, Santa Clara
29. Chen FY, Chen JF, Lin RL (2007) Low-harmonic push-pull class-E power amplifier with a pair of LC resonant networks. *IEEE Trans Circ Syst I Regul Pap* 54(3):579–589
30. Lee YT, Chiong CC, Niu DC, Wang H (2014) A high gain E-band MMIC LNA in GaAs 0.1- μm pHEMT process for radio astronomy applications. In: 9th European microwave integrated circuit conference (EuMIC), Rome, pp 456–459

<http://www.springer.com/978-3-319-69019-3>

Millimeter-Wave Low Noise Amplifiers

Božanić, M.; Sinha, S.

2018, XVIII, 334 p. 264 illus., 46 illus. in color.,

Hardcover

ISBN: 978-3-319-69019-3