

Chapter 2

Class D Audio Amplifiers and Data Conversion Fundamentals

Abstract This chapter provides a theoretical overview of the relevant aspects addressed in this book. Therefore, its purpose is to cover all aspects of the developed work. First, a brief presentation of Class D audio amplifiers is made. The reasons behind its high theoretical efficiency, in regards to traditional audio amplifiers, are presented. Also, the most common issues with Class D amplifiers are also discussed. Afterwards, the main concepts behind signal conversion in general and $\Sigma\Delta$ in particular are addressed. These include the matters of Sampling, Oversampling and Quantization. The advantages of Oversampling over Nyquist Sampling are given. A comparison between CT and DT designs is performed, the noise-shaping concept is defined and the known stability issues of $\Sigma\Delta$ s are addressed. Finally, an analysis of several $\Sigma\Delta$ architectures is performed as well as the use of 1.5-bit quantization.

2.1 Class D Audio Amplifiers

Audio amplifiers are used to amplify input audio signals in order to drive output elements (like speakers) with suitable volume and power levels, with low distortion. These amplifiers must have a good frequency response over the range of frequencies of the human ear (20 Hz to 20 kHz).

Power is dissipated in all linear output stages, because the process of generating the output signal unavoidably causes non-zero voltages and currents in at least one output transistor. The amount of power dissipation strongly depends on the method used to bias the output transistors.

Traditional Class A audio amplifiers have a maximum efficiency of about 25 % (50 % if inductive coupling is used), which is considerably low. Class B audio amplifiers can reach an efficiency of 78.5 % (theoretically), but have known disadvantages (cross-over distortion being the main one). The combination of both, Class AB audio amplifiers, can reach a similar efficiency, while practically eliminating the crossover [13].

In a Class D amplifier, the output transistors are operated as switches. They are either fully on (the voltage across it is small, ideally zero) or fully off (the current through it is zero). This leads to very low power dissipation, which results in high efficiency (ranging from 90 % to 100 % [7, 8]).

The basic block diagram of a Class D amplifier is shown in Fig. 2.1.

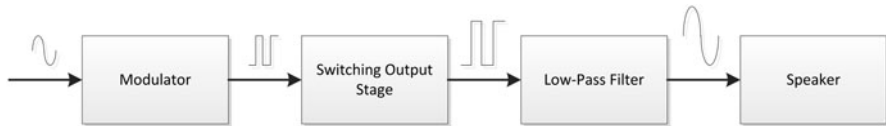


Fig. 2.1 Class D amplifier block diagram

Fig. 2.2 Half-bridge output stage

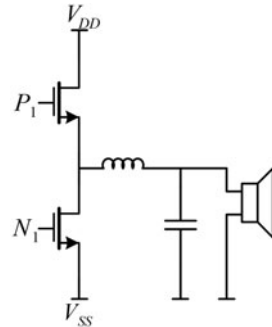
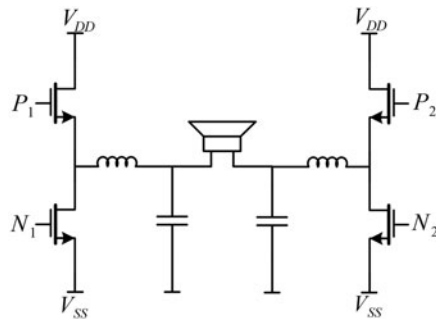


Fig. 2.3 Bridge-tied-load output stage



The input audio signal is modulated into a digital control signal which drives the power devices in the output stage. This signal can be modulated, normally, using pulse-width modulation (PWM) or pulse-density modulation (PDM). The output stage can be implemented using a Half-Bridge or a Bridge-Tied-Load (BTL) topology, illustrated in Figs. 2.2 and 2.3. Class D amplifiers are often operated in a bridged configuration to increase the output power without increasing the power supply voltages. The last stage, the low pass filter, is used to remove the high frequency PWM/PDM carrier frequency, thus retrieving the audio signal.

In a BTL amplifier, both sides of the speaker load are driven in opposite phase. Thus, a single supply can be used, while doubling the voltage swing across the load, yielding four times more output power than a Half-Bridge amplifier. Since this is a balanced operation, even order distortion is cancelled. However, a BTL amplifier needs twice the number of power switches and inductors [8].

Concerning the Half-Bridge amplifier, since it is a single-ended circuit the output signal contains a Direct-Current (DC) component with a $\frac{V_{cc}}{2}$ amplitude that might damage the speaker due to the high output power. Moreover, in the Half-Bridge Class D amplifier the energy flow can be bi-directional, which leads to the *Bus pumping* phenomena that causes the bus capacitors to be charged up by the energy flow from the load back to the supply. This occurs mainly at low audio frequencies and can be limited by adding large decoupling capacitors between both supply voltages (V_{cc} and V_{ss} , the latter being typically ground) [15].

Considering both topologies, the BTL output stage is often used as the primary solution for high quality audio applications since it provides superior audio performance and output power. Nevertheless, neither topology can reach a power efficiency of 100 % in reality since there is always switching and conduction losses that need to be considered and that limit the output stage's power efficiency.

A problem called *shoot-through* can reduce the efficiency of class-D amplifiers and lead to potential failure of the output devices [8, 15]. This results from the simultaneous conduction of both output stage's complementary transistors (when one is being "turned off" and the other is "turned on"), during which a low impedance path between V_{cc} and V_{ss} is created, leading to a large current pulse that flows between the two. This is caused by each transistor's response time, which is never immediate. The power loss that comes from shoot-through is given by

$$P_{ST} = I_{ST} \cdot (V_{cc} - V_{ss}), \quad (2.1)$$

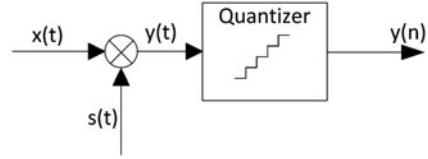
where I_{ST} is the average current that flows through both transistors. This can be eliminated by driving the output stage transistors with non-overlapping signals, avoiding simultaneous conduction.

The high switching frequency used in class-D amplifiers is a potential source of RF interference with other electronic equipment. The amplifiers must be properly shielded and grounded to prevent radiation of the switching harmonics. In addition, low-pass filters must be used on all input and output leads, including the power supply leads [8].

Another concern related to the use of Class D audio amplifiers is their Power Supply Rejection Ratio (PSRR). Due to the very low resistance that the output stage transistors have when connecting the power supplies to the low-pass filter, there is little to no isolation to any noise or voltage variation from these sources. If this problem is not properly addressed, the output signal will present a considerable level of distortion. A way of taming this problem is through the use of feedback directly from the output stage [16].

The pulses from the modulator and output stage contain not only the desired audio signal but also significant high-frequency energy (originated in the modulation process). As stated before, the low-pass filter removes this high frequency, allowing the speaker to be driven without such energy, thus minimizing the electromagnetic interference (EMI). If the modulation technique used is PWM, EMI is produced within the AM radio band. However, if PDM is employed (by a $\Sigma\Delta M$) much of the high-frequency energy will be distributed over a wide range of frequencies. Therefore, $\Sigma\Delta M$ present a potential EMI advantage over PWM [15].

Fig. 2.4 Sampling and quantizing the $x(t)$ input signal



2.2 Signal Conversion Fundamentals

In this section, the ADCs theoretical behaviour will be presented. Most of the information presented is also valid for Digital-to-Analog Converters (DACs). Both sample an input signal at a certain sampling frequency and convert it to a bitstream. For ADCs, the input signal is analog and the resulting bitstream is the digital representation of the analog signal. In the DACs case, the input signal is digitally represented by an N bit word, which is converted into a bitstream (that is a digital representation of the input signal as well) that is then applied to a filter that recovers the analog version of the input signal.

Signal conversion in ADCs is performed when an analog input signal is transformed into a digital output signal. Since an analog signal can assume infinite values in a finite time interval and it is impossible to process infinite samples, it is necessary to acquire a finite number of values. This is done by *sampling* the analog signal (usually with a constant sampling period T_s) and *quantizing* its amplitude (so that it assumes one of a finite number of values) [10]. A representation is shown in Fig. 2.4, where $x(t)$ is the input signal, $s(t)$ the sampling function and $y(t)$ the output signal.

The output signal is a result of the product of the input signal by the sampling function (Eq. 2.2).

$$y(t) = x(t) \cdot s(t) \quad (2.2)$$

This sampling function is a periodic pulse train (Eq. 2.3), where $\delta(t)$ is the Dirac delta function and T_s is the sampling interval.

$$s(t) = \sum_{n=-\infty}^{+\infty} \delta(t - nT_s) \quad (2.3)$$

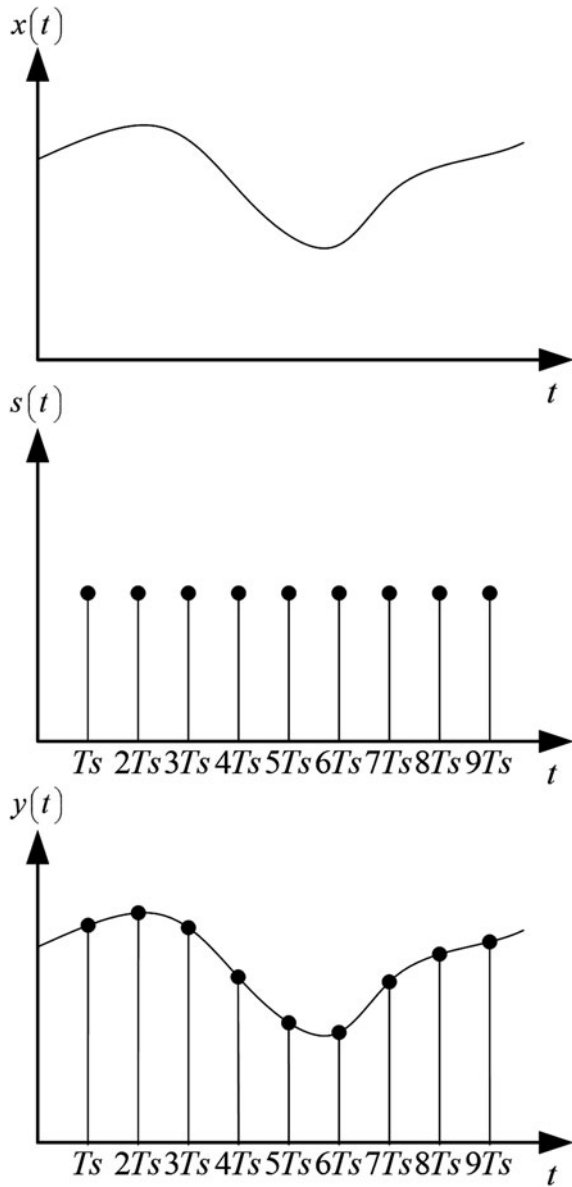
The Fourier transform of a periodic impulse train is another periodic impulse train. Thus,

$$S(j\omega) = \frac{2\pi}{T_s} \sum_{k=-\infty}^{+\infty} \delta\left(\omega - k\frac{2\pi}{T_s}\right) \quad (2.4)$$

From Eq. 2.2, since the multiplication procedure in the time domain is the equivalent of performing a convolution in the frequency domain, it is possible to write Eq. 2.5.

$$Y(j\omega) = \frac{1}{2\pi} X(j\omega) \otimes S(j\omega) = \frac{1}{T} \sum_{k=-\infty}^{+\infty} X\left(j\omega - \frac{jk2\pi}{T_s}\right) \quad (2.5)$$

Fig. 2.5 Sampling process in the time domain



Figures 2.5 and 2.6 illustrate the sampling process in the time and frequency domain respectively. Concerning Fig. 2.5, $x(t)$ represents the input signal, $s(t)$ the sampling function (Dirac pulses) and $y(t)$ the sampled signal. In regards to Fig. 2.6, the spectrum of the input signal $X(f)$, the sampling signal $S(f)$ and the sampled output $Y(f)$ are shown.

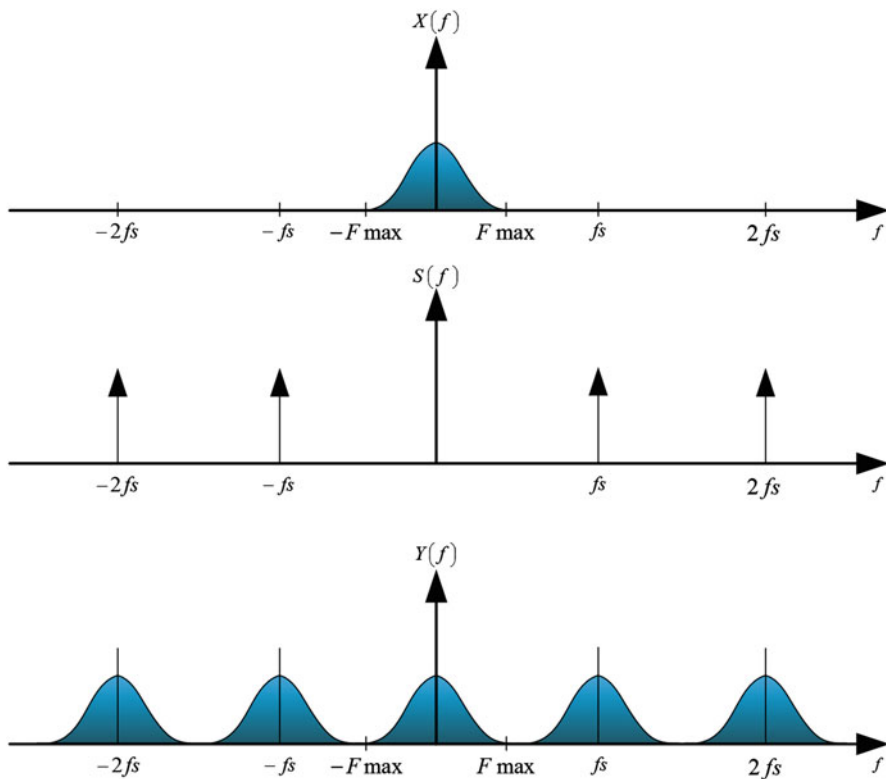


Fig. 2.6 Sampling process in the frequency domain

The spectrum separation between each spectral repetition of the input signal depends directly on the sampling frequency ($f_s = \frac{1}{T_s}$), as Eq. 2.5 shows. Thus, smaller sampling frequencies narrow the gap between each spectral repetition. If f_s is too low, an effect called *aliasing* occurs (shown in Fig. 2.7). This refers to a high-frequency component in the spectrum of the input signal apparently taking on the identity of a lower frequency in the spectrum of a sampled version of the signal (spectral overlap). Figure 2.7 shows that it is impossible to recover the original input spectrum without distortion, due to *aliasing* [9, 11, 14].

In order to avoid this effect, as stated by the *Nyquist-Shannon Theorem* [9], a signal can be sampled with no loss of information if the sampling frequency (f_s) is at least two times higher than the maximum signal frequency (B)¹. This sampling

¹ This can be enforced if a low pass pre-alias filter is used, prior to sampling, to attenuate the high-frequency components of the signal that lie outside the band of interest.

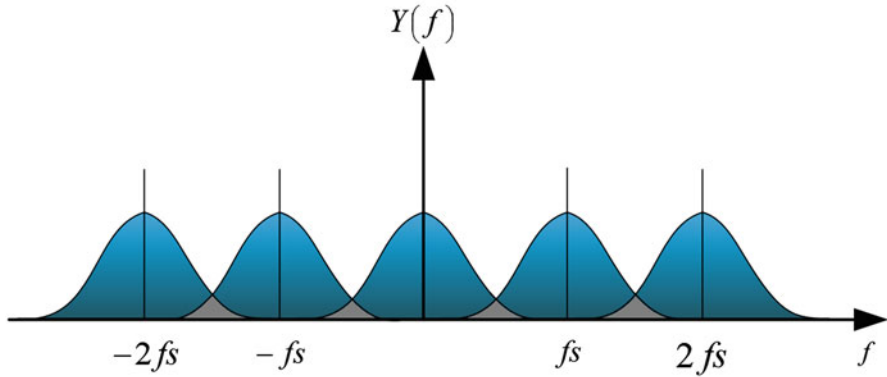


Fig. 2.7 Aliasing effect

frequency is called the Nyquist Frequency (f_N):

$$f_N = 2 \cdot B \quad (2.6)$$

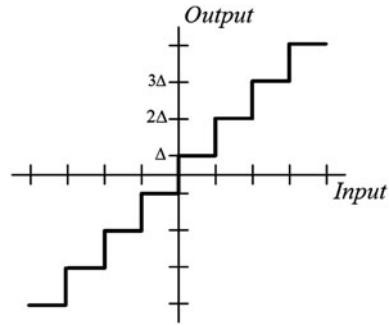
It is possible to recover the original signal through the use of an ideal low pass filter with a cut-off frequency of B .

Data converters that use $f_s = 2 \cdot B$ are called Nyquist converters, while converters that use sampling frequencies that are greater than two times the maximum signal frequency ($f_s \gg 2 \cdot B$) are called Oversampling converters. The latter will be further explored in Sect. 2.4.

Data converters in general have common performance metrics, which are usually classified in two categories: static and dynamic. The latter can be further subdivided into spectral, frequency domain and power metrics [18]. Those most used to evaluate the overall performance of a $\Sigma\Delta$ M and used in this work are briefly discussed below.

- **Signal-to-Noise Ratio (SNR):** The SNR of a converter is given by the ratio of the signal power to the noise power at the output of said converter, for a certain input amplitude. It doesn't take into account the harmonically related signal components.
- **Signal-to-Noise-and-Distortion Ratio (SNDR):** The SNDR of a converter is the ratio of the signal power to the noise and all distortion power components. Thus, it takes into account several of the harmonics (at least the 2nd and the 3rd harmonic) that lie inside the band of interest.
- **Dynamic Range (DR):** The range of signal amplitudes over which the structure operates correctly, i.e. within acceptable limits of distortion. It is determined by the maximum amplitude input signal and by the smallest detectable input signal.
- **Total Harmonic Distortion (THD):** Ratio of the sum of the signal power of all harmonic frequencies above the fundamental frequency to the power of the fundamental frequency. The harmonic distortion generated by a specific n^{th} harmonic can also be determined and is given by the ratio between the signal power and the power of the distortion component at that n^{th} harmonic of the signal frequency.

Fig. 2.8 Quantizing characteristic



2.3 Quantization Noise

The Nyquist-Shannon Theorem ensures that if the sampling frequency is at least two times the signal bandwidth, no distortion will be introduced. The same cannot be said about quantization. While sampling concerns time, quantization deals with amplitude [11].

After the samples of the analog input signal are acquired (through sampling), they must then be converted into a digital signal. This is done through the quantizing process. It has a staircase characteristic (shown in Fig. 2.8) and the difference between two adjacent quantized values is called *step size* (Δ). For large input values the quantizer output may saturate. The conversion range for which the quantizer doesn't overload is referred as the full scale (*FS*) range of the quantizer. A quantizer with a N number of bits, can represent up to 2^N amplitude levels, resulting in a Δ given by

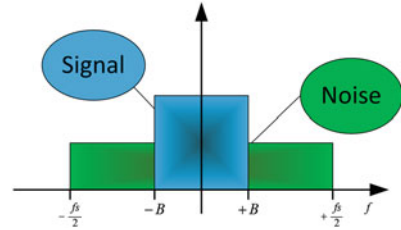
$$\Delta = \frac{FS}{2^N} \quad (2.7)$$

All input values will be rounded off to the nearest corresponding amplitude level, when applied to a quantizer with such characteristic (Fig. 2.8). Since these amplitude levels cannot possibly be exactly equal to each input value, there is always some error associated with the quantizing process [10, 11, 14]. This error is called *Quantization Error* (q_e)². Also, since the input and output range of the quantizer are not necessarily equal, the quantizer can show a non-unity gain k .

From Fig. 2.8, it's apparent that this quantization error has a maximum value of $\frac{\Delta}{2}$ and the total range of variation of this error is distributed over a range of values that go from $-\frac{\Delta}{2}$ to $+\frac{\Delta}{2}$ [10, 11]. Since the quantization error is considered a random process, uncorrelated with the input signal, it can be regarded as white noise, spread across the considered range with equal likelihood.

² If the input signal exceeds the valid input range of the quantizer, the result is a monotonously increasing quantization error.

Fig. 2.9 Spectral density of quantization noise when using oversampling



The average power of q_e (P_{q_e}) can thus be determined by averaging q_e^2 over all possible values of q_e , as follows

$$P_{q_e} = \frac{1}{\Delta} \int_{-\Delta/2}^{\Delta/2} q_e^2 dq_e = \frac{\Delta^2}{12} \quad (2.8)$$

From Fig. 2.8, the higher the resolution of the quantizer, the smaller Δ is. Thus, by increasing the number of bits (N) of the quantizer, the noise power decreases. Specifically, for each additional bit in the quantizer, the noise power decreases by 6.02 dB [9–11]. Hence the SNR is increased by 6 dB.

When a signal is sampled at a frequency f_s , the quantization noise distribution is uniform along the frequency range $\left[-\frac{f_s}{2}; +\frac{f_s}{2}\right]$, since it is considered white noise. Thus, the spectral noise power distribution is given by Eq. 2.9.

$$E(f) = \frac{\frac{\Delta^2}{12}}{f_s} = \frac{\Delta^2}{12 \cdot f_s} \quad (2.9)$$

It is clear from Eq. 2.9 that an increase of the sampling frequency will result in the quantization noise being spread over a wider band, thus reducing the noise in the band of interest, as shown in Fig. 2.9. This is one of the major advantages of using a sampling frequency higher than the Nyquist Frequency (in other words, using oversampling).

2.4 Oversampling

In many applications, such as digital audio, high resolution and linearity are required. Most standard Nyquist converters cannot provide the accuracy needed, and those that do, are too slow for signal-processing applications. Oversampling converters are capable of trading conversion time for resolution, by using simple high-tolerance analog components. Nevertheless, they require fast and complex digital signal processing stages [10, 11].

This tradeoff becomes acceptable since high-speed digital circuitry can be easily manufactured in less area, while high-resolution analog circuitry is harder to realize due to low power-supply voltages and poor transistor output impedance caused by short-channel effects [9].

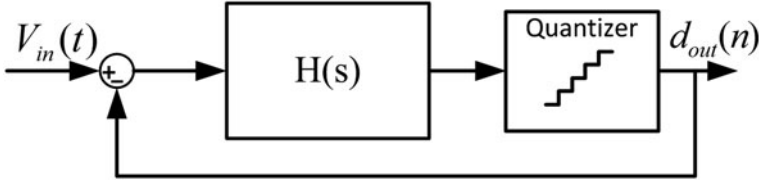


Fig. 2.10 Block diagram of a generic CT- $\Sigma\Delta$ M

For an input signal bandlimited to B , oversampling occurs when f_s is higher than two times B ($2B$ being the Nyquist frequency). The ratio between f_s and the Nyquist frequency is called Oversampling Ratio (OSR) [11] and is given by

$$OSR = \frac{f_s}{2B} \quad (2.10)$$

From Eq. 2.10, $f_s = 2 \cdot B \cdot OSR$. Applying it to Eq. 2.9 leads to Eq. 2.11,

$$E(f) = \frac{\frac{\Delta^2}{12}}{2 \cdot B \cdot OSR} = \frac{\Delta^2}{24 \cdot B \cdot OSR} \quad (2.11)$$

Equation 2.11 shows that if the OSR is doubled (i.e., sampling at twice the rate), the spectral noise power decreases by 3 dB.

As seen in Sect. 2.1, to recover the information from the sampled signal, a low pass filter is used to remove the spectral repetitions. This low-pass filter requires a sharp cut-off frequency response when Nyquist converters are used, since they cause the spectral repetitions to be very close to each other. This is another advantage of Oversampling converters: by using greater sampling frequencies, the spectral repetitions are more distant from one another. Therefore, simpler filters can be used.

Since the audio band is composed by low frequencies (ranging approximately from 20 Hz to 20 kHz), the use of Oversampling converters in audio applications should not pose a problem since today's electronics maximum frequency can be very high (up to GHz). However, due to the parasitics of the output stage, f_s should not be very high (above 2 MHz) in order to achieve high efficiency.

2.5 Continuous-Time (CT) Sigma-Delta Modulators ($\Sigma\Delta$ M)

The block diagram of a generic $\Sigma\Delta$ M is shown in Fig. 2.10. It's composed of a certain filter (typically designated as the *loop filter*) with a $H(s)$ transfer function, followed by the quantizer and a feedback loop.

In the last decades most of the work regarding $\Sigma\Delta$ M has been realized in discrete-time (DT) circuits, through a switched-capacitor (SC) implementation. DT- $\Sigma\Delta$ M are used due to the design ease of SC filters and the high degree of linearity obtained in the circuits realized. However, the design of CT- $\Sigma\Delta$ M is also feasible. In fact, $\Sigma\Delta$ M

was first proposed in the 1960s [19] through a CT implementation. CT- $\Sigma\Delta$ M can be distinguished from their discrete-time counterparts by the following:

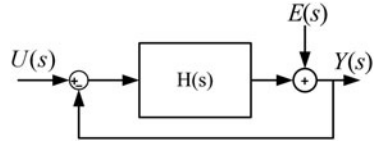
- **Sampling Operation:** In CT- $\Sigma\Delta$ M, the sampling operation takes place inside the $\Sigma\Delta$ loop, whereas in the DT- $\Sigma\Delta$ M a sample-and-hold circuit is placed before the input of the converter. Therefore, all the non-idealities of the sampling process are included, when using a CT- $\Sigma\Delta$ M. However, as it will be shown next, the errors inside the loop of a CT- $\Sigma\Delta$ M are filtered out. Moreover, the use of a CT- $\Sigma\Delta$ M can result in some kind of implicit antialiasing filtering, making unnecessary the use of a pre-alias filter. In regards to the DT case, every error that the sample-and-hold circuit generates is added to the input signal. Concerning clock jitter, CT- $\Sigma\Delta$ Ms are more susceptible to timing errors and the resulting SNR can be degraded, unlike DT- $\Sigma\Delta$ Ms who use SC circuits to realize the loop filter. However, the sampled noise may be aliased if the switch and capacitance time constant are much smaller than the sampling period. Therefore, the gain-bandwidth product (GBW) of the loop filter amplifiers in a DT- $\Sigma\Delta$ M realization limits the sampling frequency used.
- **Filter Realization:** As it will be explained ahead in this section, the loop filter is typically implemented with integrators (which can be DT or CT). While in DT implementations, based on SC circuits, the gain is determined by a capacitor ratio and is very precise, CT integrator gains typically depend on a RC or gmC product. These are subject to large process dependent variations, that may lead to mismatches and, in worst case, an unstable system. Also, the non-linearity of the passive/active components used contributes to the harmonic distortion at the modulator output. In CT- $\Sigma\Delta$ Ms, the first stage generally defines the overall accuracy of the system.
- **Quantizer Realization:** In both DT and CT implementations, the non-idealities of the quantization process are averaged out, since the quantizer resides within the feedback loop. However, the decision time of the quantizer has a different influence for each implementation: CT- $\Sigma\Delta$ Ms ideally require a very fast quantization since the result is needed right away to generate the correct CT feedback signal. DT systems, on the flip side, allow this decision time to last until half of a sampling period.
- **Feedback Realization:** In the DT system the feedback signal is applied by charging a capacitor to a reference voltage and discharging it onto the integrating capacitance, while in the CT realization the feedback waveform is integrated over time. Therefore, CT- $\Sigma\Delta$ Ms are sensitive to every mismatch on the feedback waveform that may occur, due to clock jitter or loop delay, for example. Nevertheless, there are a set of solutions to overcome this problem [2, 1].

These differences and many more that distinguish DT and CT- $\Sigma\Delta$ Ms are presented and exploited in literature [10, 11, 18]. Table 2.1 highlights the main ones.

In order to analyse the theoretical behaviour of a CT- $\Sigma\Delta$ M, a linear model of the block diagram presented in Fig. 2.10 is shown in Fig. 2.11. Since quantization is a non-linear operation (thus impossible to include directly in the model), the model is

Table 2.1 Main differences between CT and DT $\Sigma\Delta$ s

	CT- $\Sigma\Delta$	DT- $\Sigma\Delta$
Sampling frequency	Not very sensitive to the amplifiers GBW	Limited by the GBW of loop filter amplifiers
Power consumption	Lower	Higher
Anti-aliasing	Inherent	External filter needed
Sampling errors	Shaped by the loop filter	Appear directly at the ADC output
Clock jitter	Sensitive	Robust
Process variations	Sensitive	Accurate

Fig. 2.11 Linear model representation for the $\Sigma\Delta$ 

shown having two independent inputs, the quantized input signal and the quantization error (the latter is perceived as noise).

Two separate transfer functions can be derived from Fig. 2.11. *STF* concerns the signal transfer function, while *NTF* represents the noise transfer function.

$$S_{TF}(s) = \frac{Y(s)}{U(s)} = \frac{H(s)}{1 + H(s)} \quad (2.12)$$

$$N_{TF}(s) = \frac{Y(s)}{E(s)} = \frac{1}{1 + H(s)} \quad (2.13)$$

Equation 2.13 shows that when $H(s)$ goes to infinity, *NTF* goes to zero. This is due to the feedback loop, capable of averaging out the quantization error, since it is fed back negatively. Therefore, it's possible to greatly attenuate the noise over the frequency band of interest, while leaving the signal itself unaffected, by choosing $H(s)$ such that it has a large gain over the signal band. This process is called *noise-shaping* [9]. The use of noise-shaping applied to oversampling signals is commonly referred to as $\Sigma\Delta$ modulation. It refers to the process of calculating the difference (Δ) between both the output and the input signal and then integrating (Σ) it [9, 11].

To perform noise-shaping, *NTF*(s) should have a zero at dc (i.e., $s = 0$) so that it has a high-pass frequency response. Since the zeros of *NTF* correspond to the poles of $H(s)$, letting $H(s)$ be a continuous-time integrator (having a pole at $s = 0$) allows the *NTF* to present such response, as shown in Eqs. 2.14 and 2.15. For a single stage $\Sigma\Delta$, this results in a noise-shaping with a +20 dB/dec slope.

$$S_{TF}(s) = \frac{1}{s + 1} \quad (2.14)$$

$$N_{TF}(s) = \frac{s}{s + 1} \quad (2.15)$$

Equations. 2.14 and 2.15 also show that the use of integrators allows the *STF* to have unitary gain. Thus, the quantization noise is reduced over the frequency band of interest and the signal is left largely unaffected. Therefore, a greater SNR/SNDR can be achieved, as desired. A $\Sigma\Delta M$ can be implemented through a cascade of integrator stages. For each stage, the noise-shaping slope increases by +20 dB/dec [10].

To design the *NTF* of the loop filter, several noise-shaping functions can be used. The most common are the Butterworth High-Pass response and the Inverse Chebyshev High-Pass Response (although others could be used, like pure differentiators [11]).

Intuition would lead to think that in order to nearly eliminate the noise from the frequency band of interest, all that is needed to do is to add further integrator stages. However, this cannot be done in a direct fashion.

Through all the advantages that $\Sigma\Delta M$ bring, as all systems employing feedback they present some *stability* issues that need to be addressed. These issues are justified, since each integrator stage adds a -90 phase shift. Thus, $\Sigma\Delta M$ s can become unstable whenever the loop filter order is larger than 2.

Unfortunately, there is no certain solution for the stability issues of the $\Sigma\Delta M$ s, since they include a quantizer, which is a non-linear element. Therefore, circuit designers usually follow a *rule of thumb* (Lee's Criterion), which states that the peak frequency response gain of the *NTF* should be less than 1.5 [10, 11]. In mathematical terms,

$$|N_{TF}(e^{j\omega})| \leq 1.5 \quad (2.16)$$

for $0 \leq \omega \leq \pi$. This *rule* is not mandatory, as its value may range from 1.3 to 1.8 depending, but not limited, to the order of the modulator.

Another way of performing a stability analysis of the system is through the use of the root-locus graph. As with any CT system, the poles must be positioned on the left-half of the complex plane and should not cross over to the right-half of the imaginary axis. Each integrator stage of the $\Sigma\Delta M$ adds a pole at the origin. It's shown in [18, 17] that each of these poles moves directly into the right-half plane, if a linear model is used and the quantizer is modelled as $ke^{s\theta}$ (where k is the quantizer gain and θ the phase shift), for any value of k and θ .

Hence, to stabilize the $\Sigma\Delta M$, a zero (or more than one) can be introduced in the loop filter transfer function to counter the -90 phase shift that each pole of each integrator stage causes. This allows the $\Sigma\Delta M$ to be stable for a certain quantizer gain k where the poles are placed in the left-half plane, and can be achieved through feedforward or feedback compensation which can be implemented through several architectures (that in turn implement the noise-shaping functions). These are explored in Sect. 2.6. Also, the amplitude of the input signal should not be too large since it may push the poles to non-stable regions [10, 11, 17].

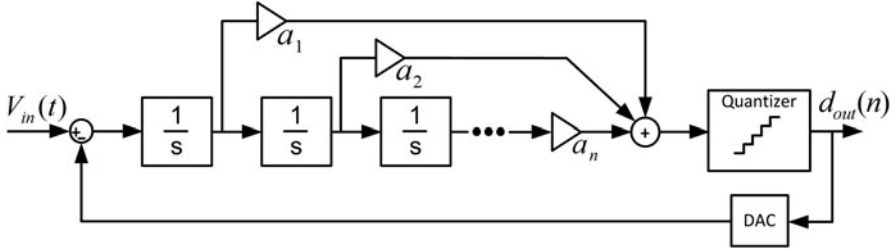


Fig. 2.12 Block diagram of a n^{th} order $\Sigma\Delta\text{M}$ with distributed feedforward

2.6 Analysis of the $\Sigma\Delta\text{M}$ Architecture

In this section, the design of the loop filter of the $\Sigma\Delta\text{M}$ through several architecture options is presented. Both feedforward and feedback techniques are described, and a comparison between the two is made. Although both techniques provide the same *NTF*, they have different signal transfer characteristics. As seen in Sect. 2.5, the *NTF* determines how much of the quantization noise is attenuated, thus determining the overall SNDR of the converter. Furthermore, the increase of the quantizer's resolution by 0.5 bit (from traditional 1-bit to 1.5-bit quantization) is discussed and its major advantages presented.

2.6.1 Feedforward Summation

The first topology that is subject of analysis is one in which a cascade of several integrators is put together, where a fraction of the output of each integrator stage is added to the output of the last stage, by means of a weighted feedforward path. It is commonly known as a cascade of integrators with feedforward (CIFF) structure.

The block diagram of a n^{th} order $\Sigma\Delta\text{M}$ with distributed feedforward is shown in Fig. 2.12. The *STF* of this structure is given by Eq. 2.17, while the *NTF* is given by Eq. 2.18.

$$STF = \frac{\sum_{i=1}^n a_{n-i+1} \cdot s^{i-1}}{s^n + \sum_{i=1}^n a_{n-i+1} \cdot s^{i-1}} \quad (2.17)$$

$$NTF = \frac{s^n}{s^n + \sum_{i=1}^n a_{n-i+1} \cdot s^{i-1}} \quad (2.18)$$

In this topology, only the error signal is fed into the loop filter, which consists mainly on quantization noise. Therefore, the first integrator can have large gain to suppress the subsequent stage's noise and distortion [10].

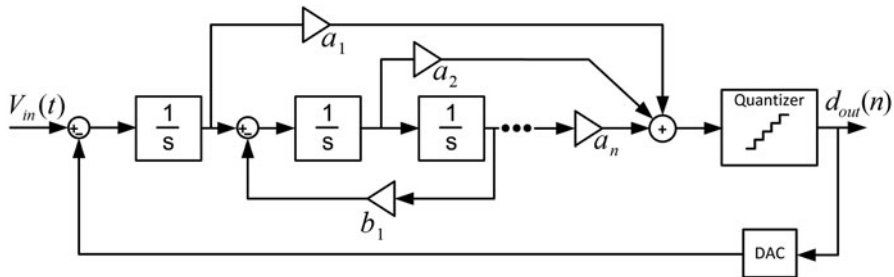


Fig. 2.13 Block diagram of a n^{th} order $\Sigma\Delta$ M with distributed feedforward and local resonator feedback

From Eq. 2.18, the zeros of the *NTF* are all placed at dc. To implement such architecture Butterworth high-pass filters are used, since Inverse Chebyshev *NTFs* have stopband zeros at non-zero frequencies and thus cannot be used. The cut-off frequency of this filter function is selected in order to limit the maximum gain of the *NTF* and eliminate the instability of the $\Sigma\Delta$ M.

A drawback of adding zeros to the *STF* is that it will create peaking at a certain frequency due to the resulting filter characteristic [17]. If input signals with these frequencies are applied to this structure, the modulator could overload due to the gain of this peaking. Possible solutions to this issue are the use of a pre-filter or the modification of the *NTF* such that flat *STFs* are obtained [11, 18].

2.6.2 Feedforward Summation and Local Resonator Feedback

In the previous structure the *NTF* zeros are all placed at dc, which limits the effectiveness of noise-shaping only to low frequencies. However, a much better performance can be achieved if these zeros are optimally distributed inside the signal bandwidth, as shown in [11].

This can be achieved by adding a negative-feedback term around pairs of integrators in the loop filter, creating local resonator stages, which allows to move the open-loop poles (that become the *NTF* zeros when the loop is closed). This structure is commonly known as a cascade of resonators with feedforward summation (CRFF).

The block diagram of a n^{th} order $\Sigma\Delta$ M with distributed feedforward is shown in Fig. 2.13. The general transfer function of a resonator is given by Eq. 2.19, where k_u is the unity-gain frequency of the integrators.

$$H(s) = \frac{k_u \cdot s}{s^2 + k_u^2} \quad (2.19)$$

In this case, it is possible to implement the inverse Chebyshev *NTF*, by picking the stopband edge frequency of the filter. With the zeros spread across the signal bandwidth a better SNR/SNDR can be obtained, when compared to the Butterworth

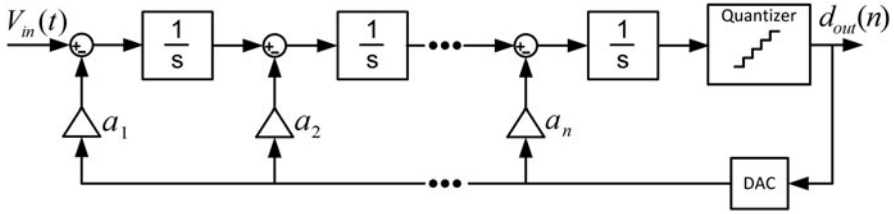


Fig. 2.14 Block diagram of a n^{th} order $\Sigma\Delta\text{M}$ with distributed feedback

alignment, as the filter presents greater attenuation over the frequency band of interest. From inspection, the *NTF* of odd order $\Sigma\Delta\text{Ms}$ have always one zero placed at dc, while the *NTF* of even order $\Sigma\Delta\text{Ms}$ have none, when using this architecture since odd order $\Sigma\Delta\text{Ms}$ have always a plain integrator beyond the cascade of resonators. Typically, this integrator is the input stage of the loop filter in order to minimize the input-referred contributions of noise sources from following stages, as stated in the previous subsection.

In this architecture the drawback of high frequency peaking of the feedforward architecture is alleviated. However, the local resonator coefficients scale with OSR^{-2} . Therefore, they rapidly decrease with the *OSR* and this technique is more viable for low *OSR*. The resonators themselves are unstable, due to their pole locations. But since they are inside a stable feedback system local oscillations are prevented. The *NTF* magnitude response will exhibit one or more notches in its frequency response [10, 11, 18].

2.6.3 Distributed Feedback

In the previous Sects. 2.6.1 and 2.6.2, feedforward was used to improve stability and provide a higher performance for the $\Sigma\Delta\text{M}$. Onwards, alternative methods recurring to feedback paths are presented. These paths also create the zeros of the *NTF*, as in the previous subsections. The structure presented in Fig. 2.14 is a group of cascaded integrators with distributed feedback (CIFB), with each integrator stage receiving a fraction of the output from the DAC, by means of a weighted feedback path.

The block diagram of a n^{th} order $\Sigma\Delta\text{M}$ with distributed feedback is shown in Fig. 2.14. The *STF* of this structure is given by Eq. 2.20, while the *NTF* is given by Eq. 2.21.

$$STF = \frac{1}{s^n + \sum_{i=1}^n a_i \cdot s^{i-1}} \quad (2.20)$$

$$NTF = \frac{s^n}{s^n + \sum_{i=1}^n a_i \cdot s^{i-1}} \quad (2.21)$$

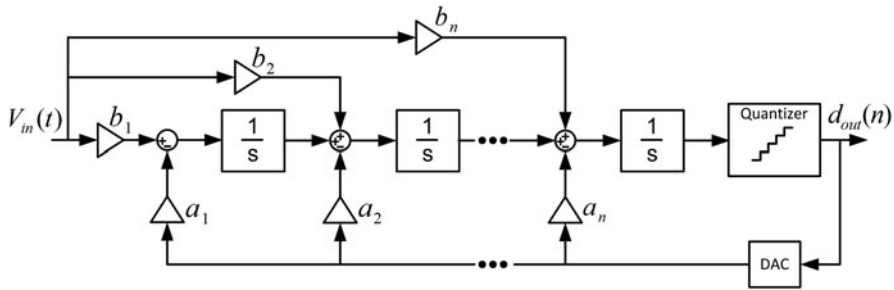


Fig. 2.15 Block diagram of a n^{th} order $\Sigma\Delta$ with distributed feedback and distributed feedforward inputs

While in the CIFF and CRFF structures only the error signal was fed into the loop filter, in the distributed feedback topology the entire output signal (including the input signal and the quantization noise) is fed back to every internal node of the filter.

As in the CIFF structure (Sect. 2.6.1), all zeros of the *NTF* lie at $s = 0$ (dc). Again, the *NTF* can be seen as a Butterworth high-pass filter and the *STF* as a Butterworth low-pass filter (note that the *STF* has no zeros in this structure). As in the CIFF structure, the cut-off frequency of this filter function is selected in order to limit the maximum gain of the *NTF* and eliminate the instability of the $\Sigma\Delta$.

One of the downsides of this architecture is that the integrator outputs contain significant amounts of the input signal as well as filtered quantization noise [11].

In this architecture, the *STF* is somewhat dependent of the *NTF*: by determining the latter, the former is automatically fixed. To overcome this, feedforward paths can be added between the input node and each integrator's summing junction. This is depicted in Fig. 2.15. The *STF* of this structure is given by Eq. 2.22, while the *NTF* is given by Eq. 2.23.

$$STF = \frac{\sum_{i=1}^n b_{n-i+1} \cdot s^{n-i}}{s^n + \sum_{i=1}^n a_i \cdot s^{i-1}} \quad (2.22)$$

$$NTF = \frac{s^n}{s^n + \sum_{i=1}^n a_i \cdot s^{i-1}} \quad (2.23)$$

Eq. 2.22 shows that the addition of these paths allows the *STF* to be independent from the *NTF*, by properly choosing the values of the b coefficients. Notice that the order of the numerator of Eq. 2.22 is one less than that of the denominator. The zeros of Eq. 2.22 can be placed in a way that it cancels some of the poles, allowing the *STF* to have a lower roll-off rate [11]. The *NTF* is exactly the same for both structures (Figs. 2.14 and 2.15).

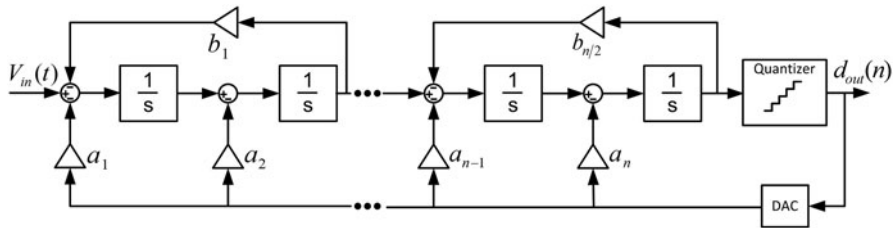


Fig. 2.16 Block diagram of a n^{th} order $\Sigma\Delta\text{M}$ with distributed feedback and local resonator feedback

2.6.4 Distributed Feedback and Local Resonator Feedback

As in the CRFF structure, local resonator stages can be added to the CIFB structure in order to shift the *NTF* zeros away from dc and spread them over the signal band. Again, this will improve the SNR/SNDR of the modulator. This structure is commonly known as a cascade of resonators with distributed feedback (CRFB).

The block diagram of a n^{th} order $\Sigma\Delta\text{M}$ with distributed feedback and local resonator feedback is shown in Fig. 2.16.

The filter coefficients are determined using the same method as in the CRFF structure, i.e. considering the *NTF* as a representation of a Chebyshev type II filter and choosing its stopband edge frequency.

2.6.5 Comparison Between Feedforward and Feedback Compensation

As seen in the previous Sects. (2.6.1–2.6.4), both feedforward and feedback compensation yield similar results, since both are capable of improving the stability of the loop while the SNR/SNDR of the modulator can be improved by adding local resonator stages. Yet, these topologies are not entirely equal since some differences exist between the two.

The feedforward summation topology only feeds the quantization noise into the loop filter while the distributed feedback topology feeds both the input signal and the quantization noise. Also, the dissipated power of the feedforward compensation is less than in the feedback compensation due to lower internal signal swings.

However, while the *STF* of an n^{th} order CIFB structure has n poles and $n-1$ zeros, the *STF* of an n^{th} order CIFB structure has only n poles. Thus, the former can be interpreted as a first order lowpass filter and the latter as a n^{th} order lowpass filter, meaning that the CIFB structure can provide much stronger filtering for high frequency signals.

Also, due to the non-ideal compensation of *STF* poles and zeros, the CIFB *STF* shows peaking at a certain frequency. In contrast, the CIFB *STF* has no zeros and this peaking is practically non-existent.

Therefore, for the reasons stated above, the architecture chosen for the loop filter is the feedback compensation (both the CIFB and the CRFB structures) since it presents no peaking and provides much stronger filtering. Also, the addition of local resonator stages allows the distribution of the *NTF* zeros along the signal bandwidth in such a way that the filter presents greater attenuation over the frequency band of interest.

2.6.6 1.5-bit Quantization

As stated in Sect. 2.1, Class D amplifiers typically achieve a power efficiency of at least 90 %. However, these efficiency values are obtained through tests that assume ideal situations, where a maximum amplitude signal is applied to the load. Thus, there is a maximum power transfer to the load and the ratio between transferred power and total power consumption leads to a very high power efficiency. When subjected to non-ideal situations, i.e. the input signal's amplitude is not maximum, the power efficiency decreases due to the reducing of power delivered to the load and the impact of switching losses that remain mostly unaltered.

For input signals of zero or near zero amplitude, the power transferred to the load is roughly non-existent but the switching activity remains high. In this case, a very low power efficiency is obtained. Since under normal working conditions the input signal's amplitude may vary from low to high levels, the resulting average Class D power efficiency will be lower than the maximum efficiency obtained through ideal conditions.

Several solutions can be employed to tackle this issue: reducing the power device's parasitic capacitances and conduction resistances, the use of a smaller switching frequency, among others [15]. However, most of them consist in some sort of trade-off where a slight increase in efficiency leads to a decrease in another important factor or the proposed solution although feasible is not viable.

Still, there is an option that has not yet been discussed: the use of multi-bit quantization instead of traditional 1-bit quantizers. This can reduce unnecessary switching of the output stage power devices, due to the existence of a number of other quantization levels besides the two levels that 1-bit quantization provides.

However, the number of quantization levels that a Class D amplifier can provide is limited by the number of amplitude levels that the output stage can represent. The Half-Bridge circuit can only represent two amplitude levels, those being when the load is connected to either supply voltage (V_{cc} and V_{ss}), making multi-bit quantization impossible. Fortunately, the BTL circuit can provide three levels (1.5-bit), the third being when the load is connected to the same potential on both terminals. This state is generally denominated the *zero-state*, since zero power is transferred to the load.

1.5-bit quantization in general and the zero-state in particular ease the representation of zero/near-zero amplitude input signals, significantly reducing switching activity for both high and low amplitude signals. Therefore, greater power efficiency can be achieved.

Multi-bit quantization exhibits far more advantages than only a power efficiency improvement. For multi-bit quantizers, the loop filter is inherently more stable since the quantizer gain (k) is well-defined (i.e. it can be approximated to be unity) and the no-overload range of the quantizer is increased, improving the linearity of the feedback in the modulator.

Also, as stated in Sect. 2.3, the higher the resolution of the quantizer, the lower the quantization noise is [9, 10, 11, 18]. This will improve the SNR/SNDR of the circuit.

In the particular case of the 1.5-bit quantizer, going from two to three levels leads to the decrease of the quantization error by a factor of two (6 dB). Also, the input range of the modulator is increased by 1.6 dB. A total improvement of around 7.6 dB of the SNDR can be achieved [20].

Design and Implementation of Sigma Delta Modulators
($\Sigma\Delta$ M) for Class D Audio Amplifiers using Differential
Pairs

Pereira, N.; Paulino, N.

2015, XIII, 76 p. 65 illus., 18 illus. in color., Softcover

ISBN: 978-3-319-11637-2