

Architectures and Synthesizers for Ultra-low Power Fast Frequency-Hopping WSN Radios

ANALOG CIRCUITS AND SIGNAL PROCESSING SERIES

Consulting Editor: Mohammed Ismail. Ohio State University

For other titles published in this series, go to
www.springer.com/series/7381

Emanuele Lopelli • Johan van der Tang •
Arthur van Roermund

Architectures and Synthesizers for Ultra-low Power Fast Frequency-Hopping WSN Radios

Dr. Emanuele Lopelli
Broadcom Corporation
Kosterijland 14
3981 AJ Bunnik
The Netherlands
elopelli@gmail.com

Dr. Johan van der Tang
Broadcom Corporation
Kosterijland 14
3981 AJ Bunnik
The Netherlands
johan@icrf.nl

Series Editors:

Mohammed Ismail
205 Dreese Laboratory
Department of Electrical Engineering
The Ohio State University
2015 Neil Avenue
Columbus, OH 43210, USA

Prof. Arthur van Roermund
Electrical Engineering
Eindhoven University of Technology
Den Dolech 2
5600 MB Eindhoven
The Netherlands
A.H.M.v.Roermund@tue.nl

Mohamad Sawan
Electrical Engineering Department
École Polytechnique de Montréal
Montréal, QC, Canada

ISBN 978-94-007-0182-3
DOI 10.1007/978-94-007-0183-0
Springer Dordrecht Heidelberg London New York

e-ISBN 978-94-007-0183-0

© Springer Science+Business Media B.V. 2011

No part of this work may be reproduced, stored in a retrieval system, or transmitted in any form or by any means, electronic, mechanical, photocopying, microfilming, recording or otherwise, without written permission from the Publisher, with the exception of any material supplied specifically for the purpose of being entered and executed on a computer system, for exclusive use by the purchaser of the work.

Cover design: VTEX, Vilnius

Printed on acid-free paper

Springer is part of Springer Science+Business Media (www.springer.com)

Contents

1	Introduction	1
1.1	Application Field	2
1.1.1	One-Way Link	3
1.1.2	Two-Way Link	3
1.2	System Requirements	4
1.2.1	One-Way Link	5
1.2.2	Two-Way Link	6
1.3	Energy Scavenging Techniques	7
1.3.1	Super-capacitor Size Estimation	9
1.3.2	Battery Size Estimation	10
1.4	General Wireless Node Requirements	10
1.4.1	Link Robustness	10
1.4.2	Data Rate	11
1.4.3	Range and Sensitivity	11
1.4.4	Turn-on and Synchronization Time	13
1.4.5	Technology Comparison and Trade-offs	13
1.5	State of the Art	14
1.5.1	Research in Industries	14
1.5.2	Research in Universities	15
1.6	The Objectives of This Book	16
1.7	Outline of the Book	17
2	System-Level and Architectural Trade-offs	19
2.1	Modulation Schemes for Ultra-low Power Wireless Nodes	19
2.1.1	Impulse Radio Transceivers	20
2.1.2	Back-scattering for RFID Applications	21
2.1.3	Sub-sampling	21
2.1.4	Super-regenerative	22
2.1.5	Spread-Spectrum Systems	22
2.2	Optimal Data-Rate	31
2.2.1	Constant Duty-Cycle	33
2.2.2	Constant Time Between Two Consecutive Transmissions	34

2.3	Transmitter Architectures	36
2.3.1	Direct Conversion	36
2.3.2	Two-Step Conversion	37
2.3.3	Offset PLL	38
2.4	Receiver Architectures	39
2.4.1	Zero-IF	39
2.4.2	Super-heterodyne	40
2.4.3	Low-IF	41
2.5	Conclusions	42
3	FHSS Systems: State-of-the-Art and Power Trade-offs	45
3.1	Synchronization	45
3.1.1	Stepped Serial Search	47
3.1.2	Matched Filter Acquisition	48
3.1.3	Two-Level Acquisition	49
3.1.4	Acquisition Methods Comparison	49
3.2	State-of-the-Art FHSS Systems	51
3.3	FH Synthesizer Architectures	54
3.4	Specifications for Ultra-low-power Frequency-Hopping Synthesizers	55
3.4.1	PLL Based	56
3.4.2	DDFS Based	58
3.5	PLL Power Estimation Model	62
3.5.1	VCO	62
3.5.2	Loop Filter	64
3.5.3	Charge Pump	64
3.5.4	PFD and Frequency Divider	65
3.5.5	Complete PLL Power Model	65
3.6	DDFS Power Estimation Model	70
3.6.1	DDFS Specifications for Frequency-Hopping Synthesizers	70
3.6.2	AA-filter Power Consumption	73
3.6.3	Phase Accumulator and ROM Power Consumption Estimation	75
3.6.4	DAC Power Consumption Estimation	78
3.6.5	Power Dissipation of the Whole DDFS	85
3.7	Summarizing Discussion	88
3.8	Conclusions	90
4	A One-Way Link Transceiver Design	93
4.1	General Guidelines for Transmitter Design	93
4.2	Transmitter Architecture	94
4.2.1	Concepts and Block Diagrams	96
4.2.2	Frequency Planning and Pre-distortion	99
4.2.3	Transmitter Specifications	106
4.2.4	Oscillator-Divider Based Architecture	111
4.2.5	Power-VCO Based Architecture	130

4.3	Receiver Architecture	133
4.3.1	RX-TX Center Frequency Alignment Algorithm	134
4.3.2	Residual Frequency Error after Pre-distortion	143
4.4	Implementation and Experimental Results	151
4.4.1	TX Node Implementation	151
4.4.2	RG Implementation	159
4.4.3	Measurement Results	159
4.4.4	Benchmarking	164
4.5	Conclusions	164
5	A Two-Way Link Transceiver Design	167
5.1	Transmitter Design General Guidelines	168
5.2	Transmitter Architecture	169
5.3	Synthesizer Design	170
5.3.1	Baseband Frequency Hopping Synthesizer Specifications	171
5.3.2	Baseband Frequency-Hopping Synthesizer Architecture	172
5.3.3	Baseband Frequency Hopping Synthesizer Implementation	179
5.4	Generation of a 288-MHz Reference Clock	189
5.5	Receiver Design at System Level	191
5.5.1	Receiver Link Budget Analysis	192
5.5.2	Receiver Building Blocks State-of-the-Art	202
5.6	Simulation and Experimental Results	207
5.6.1	Baseband Synthesizer Without the LP-notch Filter	208
5.6.2	Stand Alone Tunable LP-notch Filter	211
5.6.3	Complete Frequency Hopping Baseband Synthesizer	212
5.6.4	Benchmarking	214
5.7	Conclusions	217
6	Summary and Conclusions	219
	Appendix Walsh Based Harmonic Rejection Sensitivity Analysis	223
	References	227
	Index	233

Acronyms

AA	Anti-Aliasing
ADC	Analog to Digital Converter
ADM	Arctan-Differentiated deModulator
AFC	Automatic Frequency Control
AGC	Automatic Gain Control
AWGN	Additive White Gaussian Noise
BAW	Bulk Acoustic Wave
BER	Bit Error Rate
BFSK	Binary Frequency Shift Keying
BJT	Bipolar Junction Transistor
BPF	Band-Pass Filter
BPSK	Binary Phase Shift Keying
CDM	Cross-Differentiate Multiply
CMOS	Complementary Metal Oxide Semiconductor
CP-FSK	Coherent-Phase Frequency Shift Keying
CW	Continuous Wave
DAC	Digital to Analog Converter
DCDM	Digital Cross-Differentiate Multiply
DDFS	Direct Digital Frequency Synthesizer
DDS	Direct Digital Synthesizer
D-FF	D Flip-Flop
DPSK	Differential Phase Shift Keying
DQPSK	Differential Quadrature Phase Shift Keying
DR	Dynamic Range
DSP	Digital Signal Processor
DSSS	Direct Sequence Spread Spectrum
EEPROM	Electrically-Erasable Programmable Read-Only Memory
ENOB	Effective Number Of Bits
EOC	End Of Conversion
FBAR	thin-Film Bulk Acoustic Resonator
FCC	Federal Communication Commission

FCW	Frequency Control Word
FD	Frequency Detector
FE	Front-End
FEC	Forward Error Correction
FET	Field Effect Transistor
FF	Flip-Flop
FFT	Fast Fourier Transform
FH	Frequency Hopping
FHSS	Frequency Hopping Spread Spectrum
FLL	Frequency-Locked Loop
FM	Frequency Modulation
FOM	Figure Of Merit
FSK	Frequency-Shift Keying
FSR	Full-Scale Range
HPF	High Pass Filter
laD	Integration and Dump
IC	Integrated Circuit
INL	Integral Non-Linearity
IP	Intellectual Property
ISM	Industrial Scientific and Medical
LBR	Low Bit Rate
LEO	Low-Earth Orbit
LNA	Low Noise Amplifier
LO	Local Oscillator
LOS	Line of Sight
LP	Low Power
LPF	Low-Pass Filter
LSB	Least Significant Bit
MAC	Media Access Control
MEMS	Micro Electro Mechanical System
MOS	Metal Oxide Semiconductor
MSB	Most Significant Bit
NF	Noise Figure
NLOS	Non Line of Sight
OOK	On-Off Keying
OPAMP	OPerational AMPlifier
PA	Power Amplifier
PCB	Printed Circuit Board
PCI	Peripheral Component Interconnect
PER	Packet Error Rate
PFD	Phase Frequency Detector
PG	Processing Gain
PLL	Phase-Locked Loop
PNC	Pseudo-random Noise Code
PSD	Power Spectral Density

PVT	Process and Temperature Variation
Q	Quality factor
QoS	Quality of Service
RAM	Random Access Memory
RFID	Radio Frequency IDentification
RG	Residential Gateway
ROM	Read Only Memory
SAR	Successive Approximation Register
SAW	Surface Acoustic Wave
SFDR	Spurious Free Dynamic Range
SH	Sample and Hold
SNDR	Signal to Noise and Distortion Ratio
SNR	Signal to Noise Ratio
SOA	Silicon on Anything
SOI	Silicon on Insulator
SPI	Serial Programmable Interface
SS	Signal Spreading
SSB	Single Side-Band
ST-DFT	Short-Time Discrete Fourier Transform
TCXO	Temperature Compensated Crystal Oscillator
THSS	Time Hopping Spread Spectrum
TWD	Traveling-Wave Divider
UWB	Ultra Wide-Band
VCO	Voltage Controlled Oscillator
VGA	Voltage Gain Amplifier
WLAN	Wireless Local Area Networks
WPAN	Wireless Personal Area Networks
WSN	Wireless Sensor Networks

Chapter 1

Introduction

Wireless sensor area networks are networks that are typically limited to a small cell radius and with low aggregate data-rate. In the recent years, thanks to several years of technological advances, the field of Wireless Personal Area Networks (WPAN) and Wireless Sensor Networks (WSN) is gaining much attention.

Different standards have been developed to satisfy the requirements for medium data-rate applications (ZigBee, Bluetooth) or for identification purposes (ISO 15693). While these active and passive solutions are well matched to the requirements of a wide range of applications, there exists a gap between them. This gap is in the area of low bit-rate communications.

Figure 1.1 shows typical examples of passive and active radio technology. The passive tag is ultra low power in the sense that it doesn't require a battery, because energy is provided by an EM-field that is also used to exchange information. However, its functionality is very limited and normally only an identification tag is transmitted over a very short range. On the right side of Fig. 1.1, a typical active radio is shown with its Printed Circuit Board (PCB) application. In total it usually requires several tens of milli-Watts power in active mode, provided by a battery, but its range can be more than 100 m and its bandwidth may support transmission of multi-media content (all dependent on the standard). What is required for low bit-rate applications is simplified functionality and technology compared to Bluetooth or Zigbee, but at a cost level more closely related to passive tagging, and at much reduced power levels. The area on which this book will focus in terms of data-rate versus power dissipation and complexity is shown in Fig. 1.2.

Low bit-rate communications can be found useful in applications such as ambient intelligence, sensor networking, and control functions in the home of the consumer (domotica) as well as in the industrial and military scenario or in the environmental control. Clearly a single hardware solution will not be sufficient to support the wide range of applications while minimizing the overall power consumption.

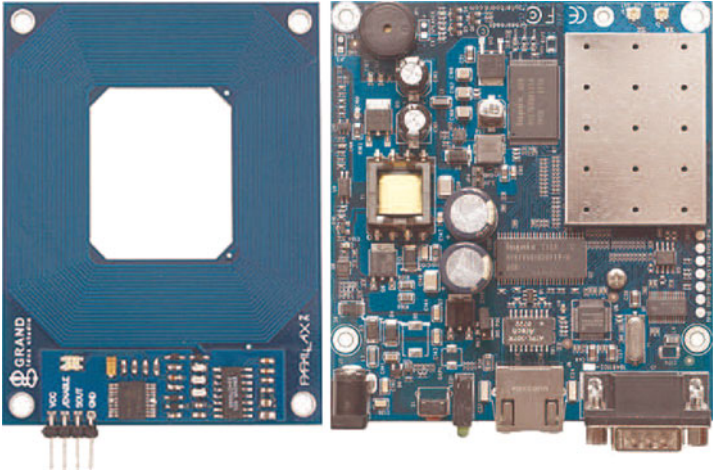


Fig. 1.1 Typical passive and active radio technologies

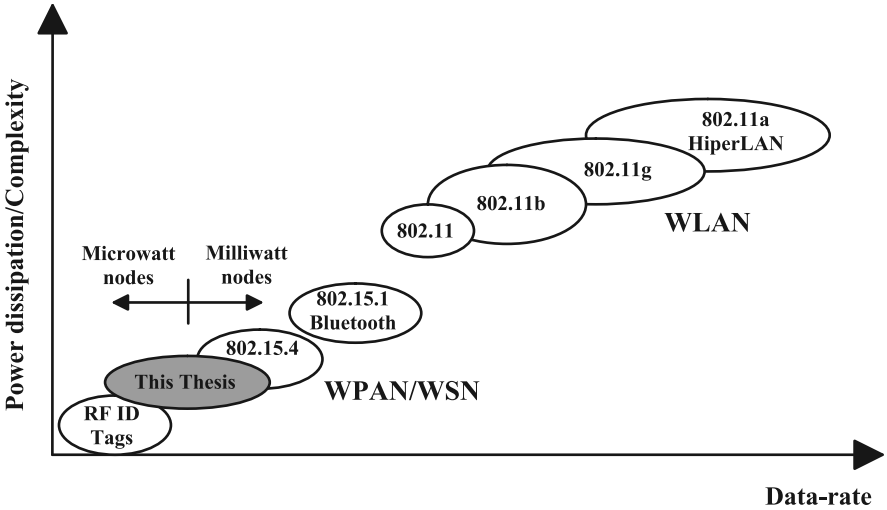


Fig. 1.2 The relative power dissipation/complexity versus data rate of various Wireless Local Area Networks (WLAN) and WSN standards and the targeted operation space of the transceiver architectures that will be investigated in this book

1.1 Application Field

It is possible to divide the application range for ultra-low power radios into two sub-parts. One part will cover all the applications, which require only to send data. These applications can use a fixed infrastructure-based network and wireless nodes can communicate only with the base station. Other applications require an ad-hoc

network, in which wireless nodes can communicate with each other without any pre-arranged infrastructure.

1.1.1 One-Way Link

The one-way link requires the presence of an infrastructure to allow the nodes to send their data. The infrastructure consists of a RG, which is mains supplied and therefore, it has a virtually unlimited power budget. In this section a number of application are highlighted, which can use an infrastructure-based network to fulfill their tasks. Link robustness can be achieved either by using an acknowledge signal or by retransmitting the data a few times in order to increase the probability of correct reception of the data packets.

In [1] a sensor network is used to monitor the sub-glacier environment in Norway. Drills are made at different depths inside the glacier ice and nodes are placed in these holes. All the wireless nodes are equipped with temperature, pressure and tilt sensors and they periodically send data to an RG placed at the top of the glacier.

In [2] sensors are used to monitor the status of the cold chain and warn the final customer in case the cold chain has been broken. It uses a four level wireless network in which the first two levels are the sensor nodes and the relay units. The sensor nodes are responsible to collect the temperature data from each product. Relay nodes collect the temperature data from several sensor nodes. They are in general more power hungry devices and can be battery powered or mains supplied. The other two levels collect all the data from different production sites and they make use of internet connections.

In [3] a WSN can be used to help rescuers in locating victims of an avalanche. For this purpose, persons at risk carry a sensor node, which allows to measure and transmit to an RG carried by the rescuers, the heart rate and the orientation of the victim. In this way the rescue team can also prioritize the rescue activity based on the status of the victim.

Another potential field of application is the automotive domain in which several kilos of cables can be saved (reducing the car weight and therefore, the average fuel consumption) by just having a sensor network which transmits data to an RG placed inside the car. In this way several parameters can be checked like tire pressure as shown in Fig. 1.3.

1.1.2 Two-Way Link

Some applications cannot use any kind of infrastructure as a support for communications. This is due to a difficult access at the site or a too large cost for the infrastructure deployment. Therefore, they require an ad-hoc network in which every wireless node can communicate with other nodes within a certain range. Using

Fig. 1.3 Siemens' tire pressure monitoring system



a multi-hop approach a large network is able to organize itself without using any base-station or RG.

For example in [4] a WSN is used to implement virtual fences. An acoustic signal is given to any animal, which is passing the virtual fences. The fence lines can be dynamically shifted by using the data collected from the WSN. This saves the cost of installation and maintenance of real fences and improve the usage of feed lots. Because the fence needs to be dynamically changed a transceiver is required.

In [5] a WSN is used to monitor a large vineyard in Oregon, USA. Several parameters like temperature, light and humidity are collected by a multi-hop network in which some of the nodes are required to act as data routers as well as simple transmitters.

Besides these applications more and more applications can be foreseen which do not require any infrastructure. For example for fire control in the forests thousands of nodes can be deployed from an airplane. These nodes will assemble themselves in an ad-hoc network detecting any fire before it spreads through the whole forest. Ad hoc networks can be used in highly contaminated areas where human presence is impossible.

This scenario presents different tasks and trade-offs with respect to the one-way link scenario. In order to optimize the overall wireless node power consumption it is necessary to conceive a dedicated architecture for the one-way link which exploits the particular features of that scenario. The same approach will be used for the two-way link in which the architecture will exploit the particular features of this different scenario. The first step consists to address the system requirements in these two different cases.

1.2 System Requirements

Some basic requirements are foreseen in WSNs. Depending on the application these requirements can vary and therefore, the system needs to be conceived in a way to optimize the power consumption for a given application range.

1.2.1 One-Way Link

Some basic requirements need to be satisfied for asymmetrical networks.

- Large number of sensors
- Very-low energy use for each node
- Very-low cost for each node

The required large number of deployed devices often goes with a small-size requirement. This means that the required device form factor has to be very small. In Fig. 1.4 is shown a 433 MHz transmitter, which uses a SAW based reference.

The SAW resonator accounts for about 40% of the whole area limiting de facto the form factor. Therefore, the one-way link device should be capable to work without using a resonator or a crystal based reference signal (i.e. it has to be a crystal-less node). The synchronization of the network will be handled mainly at the RG side where power consumption can be virtually unlimited. Besides the small form factor, the presence of a large number of devices within the communication range of each other increases the probability that two or more nodes can collide during the transmission. Therefore, the network must be able to cope with a high level of possible congestion.

Battery replacement on such a large network cannot be considered feasible. Therefore, the wireless device must be self-contained and harvest the energy from the ambient. Energy scavenging techniques will be discussed in the next section, but obviously the amount of energy that can be harvested is limited by ambient conditions and by technology limitations in the energy conversion process. The wireless node must be therefore duty-cycled; it wakes up, it transmits the required data and it falls back into an ultra low-power mode called idle mode. If the peak power consumption is too high, the duty-cycle becomes too low, shrinking the possible range of applications. Therefore, peak power consumption has to be minimized in such a way that at around 1% duty cycle the average power consumption does not exceed 100 μ W.

The last requirement regards the cost of the node. An infrastructure based network has a drawback in terms of cost given the fact that the RG has a non negligible cost. This means that to keep the concept convenient from an economical point of view the wireless device must have an extremely low cost. This requirement points again to the necessity to remove the crystal. While in high-end costly application areas, such as cellular telephones and televisions, the cost of a crystal accounts for

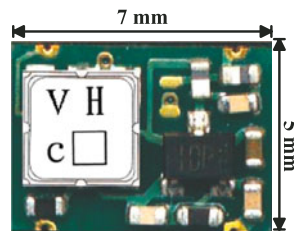


Fig. 1.4 On Shine 433 MHz transmitter

no more than 1% of the product cost, in low-end applications the crystal cost can account for as much as 10% of the unit cost.

1.2.2 Two-Way Link

A two-way link network has more stringent hardware requirements compared to the one-way link. Without an infrastructure the network needs to be an ad-hoc network. Therefore, all the wireless devices must be able to transmit as well as to receive data. The requirements in this case are the following:

- Large number of sensors
- Low energy use
- Low cost

Also ad-hoc networks must allow for deployment of a large number of devices. Given the fact that the network must be self-organized, the energy consumption can be increased at the cost of using a non-replaceable battery and of decreasing the duty cycle. A battery is used to store the energy during inactive time and can be replaced by a less bulky and more reliable capacitor in a one-way link network. Certainly the battery cannot be very big to not degrade excessively the form factor. Front-edge technology [6] has developed an ultra-thin battery (thickness between 0.1 mm and 0.3 mm depending on the capacity) which can store enough energy to allow the wireless device to operate during transmission or reception. The size of the battery is 20×25 mm and can store 0.1 mAh in the 0.1 mm version and 1 mAh in the 0.3 mm version. This battery is shown in Fig. 1.5.

A two way link wireless node will probably use a combination of harvesting technologies for a faster battery charging. Its average power consumption can be higher than in the one-way link and it will use a crystal. This requirement becomes mandatory because no infrastructure is present in the network. Therefore, also the receiver part is power constrained. This means that also the processing power of

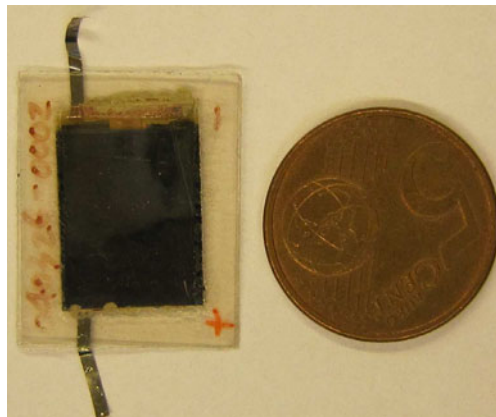


Fig. 1.5 NanoEnergy battery from FrontEdge Technology, Inc.

the receiver cannot be too large, which limits the amount of digital computation the node can do. This, of course, translates in a higher costs and in a larger form factor with respect to the one-way link. The form factor can be optimized by an optimal design of the wireless node. The cost of a single node can be higher without necessary affecting the overall network cost because the two-way link scenario does not require any infrastructure.

1.3 Energy Scavenging Techniques

Several scavenging techniques have been studied in the recent years. However, it is unlikely that a single solution will satisfy the total application space. For example a solar cell requires minimum lighting conditions, a piezoelectric generator sufficient vibration, and a Carnot-based generator sufficient temperature gradient.

In Fig. 1.6 various scavengeable energy sources that can be used in autonomous wireless nodes are shown. When considering one of these sources as a possible harvesting field, one main characteristic that should be considered is the power density of the harvesting technology. One of the most common scavenging techniques is to harvest energy from an RF signal. An electric field of 1 V/m yields only $0.26 \mu\text{W}/\text{cm}^2$, but such field strengths are quite rare [8]. This technique is generally used for RFID tags, which have a power consumption between 1 and $100 \mu\text{W}$. Energy can be harvested by using solar cells. While 1 cm^2 of standard solar cells produces around 100 mW under bright sun, it only generates no more than $100 \mu\text{W}$

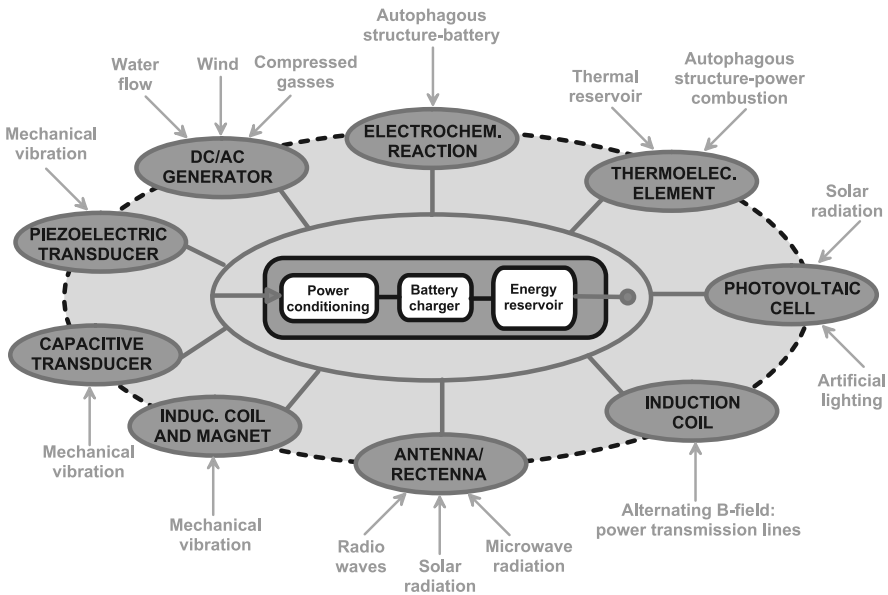


Fig. 1.6 Various scavengeable energies that can be used for portable systems (from [7])

Table 1.1 Power densities harvesting technologies

Harvesting technology	Power density
Solar cells (outdoor, at noon)	15 mW/cm ²
Solar cells (indoor)	<100 μ W/cm ²
Piezoelectric (shoe inserts)	40 μ W/cm ³
Vibration (small microwave oven)	116 μ W/cm ³
Thermoelectric (10°C gradient)	330 μ W/cm ³
Acoustic noise (100 dB)	0.96 μ W/cm ³
RF signal (1 V/m)	0.26 μ W/cm ²

in a typically illuminated office [8]. Also thermoelectric conversion can be used as an energy scavenging technique. Unfortunately, the Carnot cycle limits the use of this technique for small temperature gradients by squeezing the efficiency below 5% for about 15 degrees temperature difference [8]. Another possible solution to the scavenging problem can be found in the vibrational energy. If 1 cm³ volume is considered, then up to 4 μ W power can be generated from a typical human motion, whereas 800 μ W can be harvested from machine-induced stimuli [8].

Power density examples for the most common harvesting technologies, which can be used in autonomous wireless nodes are listed in Table 1.1 [9]. Depending on the technology it is possible to estimate the maximum duty-cycle of the wireless node as a function of the peak power consumption of the node.

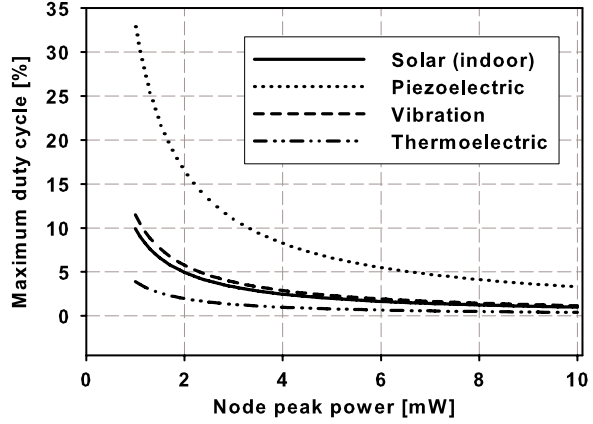
$$d_{\max} = \frac{\overline{P_{\text{harv}}} - P_{\text{sleep}}}{P_{\text{act}}} \quad (1.1)$$

where $d_{\max}$ is the maximum duty-cycle, $\overline{P_{\text{harv}}}$ is the average harvested energy, P_{sleep} the power consumption in the idle mode and P_{act} the peak value of the active power consumption.

In Fig. 1.7 the maximum duty-cycle is plotted versus the node peak power consumption for the harvesting technologies in Fig. 1.6. Given the required small form factor, and depending on the harvesting technology used, either a one cm² area or a one cm³ volume is considered in the following example. The idle power consumption is fixed to 1 μ W, and therefore, acoustic harvesting and RF signal harvesting are not possible within a volume of 1 cm³ and an area of 1 cm² respectively. Therefore, these two harvesting technologies will not be considered further in the discussion. As can be seen from Fig. 1.7 the node peak power consumption should not exceed a few milli-Watts. In this way almost every harvesting technology can guarantee an ever-standing source of energy. If the peak power consumption is too high, then the maximum duty-cycle has to be quite small and this will narrow the allowed application range.

While some harvesting technologies can generate power the all day long, some others are limited to a part of the day. For example a vibrational based harvesting element can harvest energy continuously while a solar cell needs the light. Therefore, it can harvest energy only when office light is on or during the day in outdoor environments.

Fig. 1.7 Maximum duty-cycle versus node peak power consumption for different harvesting technologies



While a combination of such technologies can guarantee a full day harvesting capability, it can be quite costly. Therefore, another possible solution is to have an energy reservoir (see also Fig. 1.6) which can take the form of either a battery or a super-capacitor. In this situation the wireless node cannot work at its maximum duty-cycle otherwise during the day the battery will not be recharged for the coming night. This means that a reduced duty-cycle has to be used in order to save some energy during two consecutive transmissions. This energy, stored in the battery or in a super-capacitor will be used during night time.

1.3.1 Super-capacitor Size Estimation

Remembering that the energy stored on a capacitor is equal to $\frac{1}{2}C(V_H^2 - V_L^2)$ where C is the capacitor value, V_H is the maximum voltage and V_L is the minimum voltage¹ the following relation holds:

$$C = 2T(P_{\text{act}}d + (1 - d)P_{\text{sleep}}) \frac{1}{V_H^2 - V_L^2} \quad (1.2)$$

where T is the time interval between two consecutive transmissions. The value C is the capacitance value required to allow the node to transmit a data packet. Supposing that a solar cell is used, there is the possibility that no harvesting is possible during part of the day. Supposing that the largest amount of time the node cannot harvest any energy from the ambient is $T_{\text{no-harv}}$, then the capacitance value C has to be multiplied by $\frac{T_{\text{no-harv}}}{T}$. For a 2 mW peak power consumption, 2 μ W idle power consumption, 10 hours without any possible harvesting (a full night), a duty-cycle of 1%, a V_H of 3 V and a V_L of 1 V a 0.1 F super-capacitor is required.

¹The maximum voltage is the voltage at full charge and the minimum voltage is the voltage when the capacitor is considered fully discharged.

1.3.2 Battery Size Estimation

The calculation of the required mAh for a battery is a bit more cumbersome. The required energy during the time T is the same as for the super-capacitor example:

$$E = T(P_{\text{act}}d + (1 - d)P_{\text{sleep}}) \quad (1.3)$$

The energy stored in a battery of capacitance X expressed in mAh is equal to

$$E = 3.6 \times X \times RAV \quad (1.4)$$

where RAV is the remaining average voltage of the battery and in the case of a perfectly linear battery discharge curve is $(V_1 + V_2)/2$ where V_1 and V_2 are the fully charged and fully discharged battery voltages.

The energy required for the same situation as in Sect. 1.3.1 equals 0.79 J. The Front-Edge Technology battery discussed in Sect. 1.2.2 has a RAV equal to roughly 3.95 V. The required capacity of the battery is about 0.055 mAh. In a two way link it is reasonable to suppose that the receiver consumes the same peak power as the transmitter. Therefore, for a two-way link a battery of roughly 0.11 mAh can allow the node to work during the night.

1.4 General Wireless Node Requirements

This section summarizes, starting from the unique challenges an ultra-low power wireless node needs to face in order to be able to maximize its life time, the node requirements at system level. These requirements are set in order to allow the wireless node to last “forever” by harvesting energy from the ambient using any of the form mentioned in Fig. 1.6. Given the specific range of applications targeted, the requirements are very different from other wireless low power standards like Bluetooth or Zigbee. Those standards, though targeting the low power application area, are still too power hungry to allow to build a wireless network able to harvest the required energy from the ambient at a reasonably duty-cycle (i.e. 0.1% to 1%).

1.4.1 Link Robustness

Though the overall power consumption has to be very limited it is important to have a wireless network which is reliable. The robustness of the wireless link depends on the type of modulation used as well as on diversity schemes. It is well known, for example, that all amplitude modulations are very weak in highly fading scenarios like in an indoor environment. Phase modulation exhibits a higher level of robustness and therefore, is preferable for wireless nodes in indoor environments.

Diversity can assume various forms:

- Spatial diversity

- Time diversity
- Frequency diversity

Spatial diversity makes use of different receivers and antennas in such a way to constructively add the arriving signals enhancing the receiver output. Though very effective, the use of multiple receiver paths rules out this possibility given the power budget available.

Time diversity implies that the same data is transmitted multiple times, or a redundant error code is added. Transmitting data multiple times especially during strong fading condition can result in a non-optimal use of the scarce energy resources available.

Signal Spreading (SS) works quite well in situations with strong narrow band interference signals since the SS signal has a unique form of frequency diversity. The actual signal spreading may be achieved with one of three basic techniques. These include direct sequence, frequency hopped and time hopped forms. Spread Spectrum techniques will be analyzed in Sect. 2.1.5 more into details.

1.4.2 Data Rate

WSNs are unique when the data-rate is considered. Generally the amount of data collected is very low, packets are very small and most of the time a packet needs to be sent at unpredictable times. Therefore, while increasing the data-rate can seem a good solution to reduce the time the transmitter must be on, a few drawbacks need to be taken also into account.

After a certain data-rate the system will be limited by its turn-on time. This behavior is more accentuated for systems which require synchronization in code and frequency like SS systems. Moreover if the data-rate is too high, to not be limited by the turn-on time, the amount of data the node needs to gather between two consecutive transmissions has to be bigger. This in practice means that a larger data packet needs to be sent at each transmission time. Therefore, for a given average amount of data generated by the network (which depends on the application), the delay time of the entire network will increase. Concluding, the ending effect will be a narrower application range that can be covered by the wireless node (for example in a burglar alarm the police needs to know immediately about any intrusion).

These considerations set a big difference between common low power networks and WSNs.

1.4.3 Range and Sensitivity

A network meant for WPANs or WSNs is generally considered a short-range wireless network. With short-range, in this book, a communication range smaller than

10 meters is considered. The environment is an indoor environment where reflections can pose a severe problem.

Starting from the communication range it is possible to estimate the required sensitivity on the receiver side as a function of the transmitted power. Several parameters affect the signal on its path between the transmitter and the receiver. The signal is subjected to a gain L_{path} ² due to the propagation of the signal in the air. The transmitter and receiver antennas have characteristic gains that can be denoted as G_{TX} and G_{RX} . Furthermore, especially in an indoor environment, the signal is subjected to multipath reflections, which can corrupt the received signal. Likewise, obstacles will greatly reduce the strength of the received signal.

Unity gain antennas are supposed at the transmitter as well as at the receiver. Therefore, expressing all the parameters in decibel units, the received signal power is:

$$P_{\text{RX}} = P_{\text{TX}} + L_{\text{path}} \quad (1.5)$$

where P_{RX} and P_{TX} are respectively the received and the transmitted power. The attenuation losses due to propagation and fading need to be calculated. When no objects are present between the transmitter and the receiver, a Line of Sight (LOS) condition is present, while when there are objects in between the path we generally refer to a Non Line of Sight (NLOS) condition. Generally, all wireless links have both a LOS and NLOS propagation paths. To derive the attenuation when a NLOS condition occurs we first calculate the propagation loss in a LOS situation. The free space attenuation can be expressed as (in [10] a full derivation of the following formula can be found):

$$L_{\text{path,LOS}} = 27.56 \text{ dB} - 20 \log_{10}(f_c) - 20 \log_{10}(r_0) \quad (1.6)$$

where f_c is the carrier frequency expressed in MHz and r_0 is the un obstructed communication distance between the transmitter and the receiver expressed in meters. The NLOS path loss can be approximated as [10]

$$L_{\text{path}} = L_{\text{path,LOS}} - 10 \cdot n \cdot \log_{10}\left(\frac{r}{r_0}\right) \quad (1.7)$$

where n is the path loss exponent, which indicates how fast the path loss increases with distance, $L_{\text{path,LOS}}$ is the corresponding propagation loss of the LOS path, and r is the distance between transmitter and receiver.

The required sensitivity depends on the carrier frequency. At 915 MHz and 2.4 GHz, given $n = 4$, a transmitted power equal to -6 dBm translates in -68.1 dBm and -76.5 dBm respectively for an unobstructed communication distance of three meters and a communication distance of 10 meters.

²This term, will have a negative sign because the signal is attenuated while traveling from the transmitter to the receiver.

1.4.4 Turn-on and Synchronization Time

As already mentioned in Sect. 1.2, any wireless device to be used in WPAN or WSN networks needs to be duty-cycled. Therefore, between the time in which the node can send the data and the time in which the node turns on some time is required to settle the correct operating point of the system as well as for synchronization.

This time, including the synchronization time, must be minimized in order to reduce the power consumption wasted during this process. This becomes very important especially for low data-rate systems given the fact that packets are generated very rarely and their size is very small (maximum a few hundreds of bits). Therefore, while link robustness is an important requirement, it is not allowed to come at the expense of a too long synchronization time to not degrade too much the device average power consumption.

1.4.5 Technology Comparison and Trade-offs

Two transistor operation principles are generally used in analog and digital circuit design. The distinction is based on the type of carrier transportation mechanism involved:

- Field Effect Transistor (FET)
- Bipolar Junction Transistor (BJT)

FET devices are unipolar devices in which only the majority carriers are responsible for the transportation mechanism. The majority carriers are drifted from the source to the drain via an electric field, while the current is modulated via the gate voltage through channel width modulation. For the designer the control parameter is the trans-conductance (g_m). Therefore, the FET device acts as a trans-conductance amplifier.

On the other hand, in BJTs, both electrons and holes are involved in the transportation mechanism. The collector current is modulated through the base current and therefore, the BJT acts as a current amplifier with an amplification constant often referred as β .

This book takes into account only silicon based technologies and therefore, two transistor technologies will be considered:

- Complementary Metal Oxide Semiconductor (CMOS)
- Silicon BJT

Regarding low voltage operations for low power, BJT devices have a disadvantage because they are limited by the base-emitter voltage V_{be} . CMOS is limited by the so called threshold voltage V_{th} which is nowadays quite smaller than V_{be} . On the other hand, V_{be} is much more stable with process variation than V_{th} allowing a better control over fundamental building blocks like differential pairs.

CMOS technology allows also to make smaller devices. These devices are self-isolating but they do not have the beneficial current scaling property of the BJT devices. This means that CMOS devices need to be scaled correctly as the operating current changes.

Digital blocks in a bipolar process require a DC current. This is due to the low input impedance of BJT devices. On the other hand the input impedance of a CMOS transistor is very high. Because complementary technology is used there is generally never a direct path between the supply voltage and ground in steady state conditions. This means that CMOS digital blocks consume power only during switching transients, while bipolar ones continue to waste power also in steady state conditions.

For the one-way link the Silicon on Anything (SOA), which is a bipolar technology, has been chosen. The SOA technology has been optimized for low power applications. The active layer is glued onto any kind of substrate after processing. For the design in this book, glass has been used, because it is cheap and has low losses over a wide range of frequencies.

Some inherent advantages of this particular bipolar process are [11]:

- lateral NPN transistor with $0.1 \mu\text{m}^2$ emitter area using a $0.5 \mu\text{m}$ lithography
- 5 to 20 times smaller interconnection parasitic to ground
- Integrated inductors with Quality factor (Q) values up to 40

Therefore, it offers a very low cost solution ($0.5 \mu\text{m}$ lithography) and the possibility to reduce the power thanks to the high Q values of passive components. Therefore, this technology has been used in the one-way link scenario in which cost and power are both heavily constrained.

On the other hand, given the higher complexity, the standard CMOS 90 nm technology has been used for the two-way link wireless node. The costs are higher and the Q values of passive components are lower, but a large digital back-end at very low power (scalable with advances in technology) and in a relatively small area can be easily implemented in CMOS.

1.5 State of the Art

Starting with university research, the interest in ultra-low power wireless devices has increasingly spread among companies as well. In the wide scenario of ultra low-power devices, various pioneering investigations have been conducted to prove the feasibility of an ultra-low power wireless node in terms of power consumption and robustness of the communication link.

1.5.1 Research in Industries

Several wireless products, which claim to be ultra-low power, are present on the market. Rarely these products can be used as core block for an autonomous node.

Pioneering researches toward the development of this kind of wireless nodes can be found in [12, 13]. The Eco node [13] has been designed to monitor the spontaneous motion of preterm infants using the 2.4 GHz Industrial Scientific and Medical (ISM) band at 1 Mbps data-rate. While showing a good form factor (648 mm³ by 1.6 grams), its power consumption is still far away from the minimum target required by a truly ultra-low power node. Indeed, even at 10 kbps and -5 dBm output power, it consumes 20.4 mW in Tx mode and 57 mW in Rx mode (at 1 Mbps) considering only the radio device. Robustness of the link by frequency diversity is achieved by using an FHSS technique.

The Telos node [12] complies with the ZigBee standard (which makes use of an SS technique). As a result, while having a reduced data-rate (250 kbps), it has an overall power consumption of around 73 mW from 1.8 V power supply at 0 dBm transmitted power.

1.5.2 Research in Universities

Different universities are involved in pioneering research on ultra-low power devices and networks. At Berkeley university an ultra-low power Micro Electro Mechanical System (MEMS)-based transceiver has been developed [14]. Whereas using a 1.9 GHz carrier frequency and only two channels, the receiver power consumption is 3 mA from a 1.2 V power supply. The data rate is 40 kbps at 1.6 dBm output power. The low receiver power consumption is mainly obtained by using a high- Q MEMS resonator implemented as a thin-Film Bulk Acoustic Resonator (FBAR). If more channels are needed, like in the case of an FHSS transceiver, the hardware requirement will increase linearly with the number of channels, making this choice impractical from a low-power point of view. The transmitter part adopts direct modulation of the oscillator and MEMS technology, eliminating power hungry blocks like PLL and mixers, therefore, reducing the overall power consumption. Two major drawbacks can be foreseen in the proposed architecture. While reducing the circuit and technological gap toward an autonomous node, it relies on non-standard components (MEMS), which will increase the cost and will require higher driving voltage. Furthermore, it lacks on robustness due to the use of only two channels, while requiring a linear increase of the power consumption with the channels' number, if a more robust frequency diversity scheme has to be implemented.

At the CSEM institute the WiseNet [15] project aims to optimize both the Media Access Control (MAC) and the physical layer to obtain a robust, low-power solution for sensor networks. Whereas not using the worldwide available 2.4 GHz ISM band but the lower 433 MHz ISM band, it achieves a power consumption of only 1.8 mW from a minimum supply voltage of 0.9 V in RX mode. This result was achieved by a combination of circuit and system innovative techniques and the use of the low-frequency 433 MHz band, which reduces the power consumption of the most power hungry blocks like the frequency synthesizer. In TX mode a high power consumption of 31.5 mW was reported mainly due to the choice of a high output power of 10 dBm. Data-rate is 25 kbps with Frequency-Shift Keying (FSK) modulation.

The proposed solution, while relying partially on the lower frequency band to reduce the power consumption, still requires external components like high- Q inductors for the LC-tank circuit and external RF filters, which will deteriorate form-factor and power consumption at higher frequencies.

1.6 The Objectives of This Book

The primary goal of this book is to propose new guidelines, concepts and design techniques that can be used for future ultra-low power wireless radio links. To achieve this scope, innovative integrated transceiver front-ends, that are suitable for low bit rate data transfer and are consuming a very small amount of power, will be realized in the two different technologies proposed in Sect. 1.4.5. Many applications and WSN nodes require only low-bit rates (only a few bit/s up to 1 kb/s). Low Power (LP)/Low Bit Rate (LBR) transceivers require a new architectural approach, compared to moderate and high-speed multi-media wireless links, in order to make battery lifetime and thus battery replacement practical for the consumer or to completely avoid any battery replacement. This book investigates front-end architectures including frequency and modulation schemes, baseband complexity versus front-end complexity (and thus power), frequency synchronization and frequency recovery algorithm, synthesizers architectures and concepts that are suitable for LP/LBR transceivers.

A unique technology comparison is part of this book. The comparison between a mainstream CMOS technology and a bipolar Silicon on Insulator (SOI) technology transferred to glass points out the advantages and disadvantages of low-power front-end building blocks in each technology. An important investigation on transceiver architectures without an absolute frequency reference (especially on the transmit side) is one of the main goals of this book. This would eliminate the need for integration of an absolute frequency reference and is expected to yield significant savings in power, and results in a significant form factor reduction of a microwatt node.

The transmitter can be implemented as a direct conversion or two-step conversion or by using a PLL system. The receiver can employ a super-heterodyne architecture, a zero-IF or a low-IF architecture. Moreover different new concepts can be applied to those basic architectures to improve the reliability of the link, to increase the data-rate or to decrease the overall power consumption. Therefore, this book aims to analyze the design space in the transceiver field and to propose a suitable combination of system architecture, modulation scheme, synchronization algorithm and frequency synthesizer for an ultra-low power wireless network in two cases, which will cover most of the application space of WSNs:

- One-way link
- Two-way link

The proposed aforementioned combinations exploit specific characteristics of the two links in order to minimize the overall power consumption while allowing a reliable communication link.

1.7 Outline of the Book

In the following chapters the one-way link and the two-way link scenarios will be explored. In Chap. 2 a survey of a number of possible architectures suitable for an ultra-low power implementation will be discussed. Different modulation schemes will be analyzed from a low power point of view and an optimal data-rate will be proposed, which minimizes the overall power consumption including start-up power consumption. Transmitter and receiver architectures will be also analyzed trying to find an optimal solution from a realization and a power point of view.

In Chap. 3 an FHSS system will be discussed in detail as an optimal solution, which combines link robustness and potential for a low power implementation. Several topics will be analyzed like, for example, synchronization mechanisms and synthesizer architectures. Starting from a state-of-the-art analysis on FHSS synthesizers, a power estimation model for a Direct Digital Synthesizer (DDS) based synthesizer will be developed. Together with a PLL model already available [16] the need for a new approach in frequency synthesis is demonstrated.

This brings to the next two chapters of this book, which will focus on new synthesizer concepts for FHSS systems. In Chap. 4, a transmitter for a one-way link will be disclosed. This transmitter exploits the specific characteristics of the asymmetric wireless scenario in order to minimize the power consumption of the power constrained node (the transmitter). A novel frequency pre-distortion concept is introduced together with a new fast and precise frequency recovery algorithm. These two new concepts allow a simplified and low-power FHSS synthesizer and will make a crystal-less wireless node a reality. The receiver will also be disclosed not from a hardware point of view but from a higher system level point of view. Indeed, given the fact that the receiver is mains supplied, it is not power constrained and its analysis down to hardware implementation is out of the scope of this book. Nevertheless, it will be proven that a receiver can be implemented so that a robust wireless link can be implemented at a very low transmitter power consumption. This very low power consumption makes realistic the implementation of an autonomous wireless node.

Chapter 5 deals with a low power transceiver implementation for the two-way link scenario. Given the fact that the most power hungry block is the hopping synthesizer while the RF part is power constrained by the high frequency operation, a novel suitable architecture for the FHSS synthesizer is disclosed. This architecture takes the maximum benefit from the use of a crystal and from the use of CMOS technology. Therefore, at implementation level it is mostly digital allowing a very low power consumption, which scales down with shrinking in technology. A system level analysis for the receiver is carried out showing the possibility to implement it in a low power fashion allowing in this way the wireless node to be autonomous.

Chapter 6 concludes this book summarizing the innovation steps involved and discussing future developments in the direction of an ultra-low power transceiver design for WSNs.

Chapter 2

System-Level and Architectural Trade-offs

This chapter focuses on high level design of ultra-low power wireless nodes. First, different system architectures are compared in order to assess advantages and drawbacks of each architecture from a power consumption point of view. Second, different modulation formats are compared and an optimal data-rate is chosen in order to minimize the average power consumption of the node. Finally, the most common transmitter and receiver architectures are reviewed.

2.1 Modulation Schemes for Ultra-low Power Wireless Nodes

Different radio architectures have been recently studied in order to reduce the power consumption. Some of these architectures comprise Ultra Wide-Band (UWB) transceivers, Back-scattering transceivers, Sub-sampling and Super-Regenerative transceivers, as well as spread-spectrum based transceivers (both frequency hopping and direct sequence).

Though spread spectrum techniques are also ultra-wideband modulation schemes, in this book, “ultra-wideband modulation” is used to refer to a spectrum that is larger than 500 MHz (e.g. impulse radio based schemes). Although a spread-spectrum modulated signal can have a bandwidth larger than 500 MHz, this is not a necessary condition. Therefore, “with spread-spectrum modulated signal”, in this book, we refer to any signal in which the transmitted bandwidth is much larger than the signal bandwidth (e.g. the transmitted bandwidth is larger than 10 times the signal bandwidth).

Looking finally to regulations, Federal Communication Commission (FCC) rules specify UWB technology as any wireless transmission scheme that occupies more than 500 MHz of absolute bandwidth or more than 20% of the carrier frequency.

2.1.1 Impulse Radio Transceivers

Among different architectures suitable for an ultra-low power implementation, UWB based systems are gaining more and more attention.

The most important characteristic of UWB systems is the capability to operate in the power-limited regime. In this regime, the channel capacity increases almost linearly with power, whereas at high Signal to Noise Ratio (SNR) it increases only as the logarithm of the signal power as shown by the Shannon theorem

$$C = BW \times \log_2 \left(1 + \frac{P_S}{P_N} \right) \quad (2.1)$$

where P_S is the average signal power at the receiver, P_N is the average noise power at the receiver and BW is the channel bandwidth. For low data-rate applications (small C), it can be seen from (2.1) that the required SNR can be very small given an available bandwidth in excess of several hundreds MHz. A small SNR translates in a small transmitted power and as a result in a reduction of the overall transmitter power consumption.

Although UWB transceivers can have reduced hardware complexity, they pose several challenges in terms of power consumption. In Fig. 2.1 a schematic block diagram of an UWB transceiver is shown. The biggest challenge in terms of power consumption is the Analog to Digital Converter (ADC). If all the available bandwidth is used, the sampling rate has to be in the order of several Gsamples per second. Furthermore, the ADC should have a very wide dynamic range to resolve the wanted signal from the strong interferers. This implies the use of low-resolution full-flash converters. It can be proven [17] that a 4-bit, 15 GHz flash ADC can easily consume hundreds of milliwatts of power. Even if a 1-bit ADC at 2 Gsample/s is used, the predicted power consumption of the ADC remains around 5 mW [18]. Furthermore, the requirement on the clock generation circuitry can be very demanding in terms of jitter.

Besides these drawbacks, wideband Low Noise Amplifier (LNA) and antenna design are challenging when the used bandwidth is in excess of some Gigahertz. The antenna gain, for example, should be proportional to the frequency [19], but most conventional antennas do not satisfy this requirement. LNA design appears quite challenging when looking at power consumption of state-of-the-art wideband LNAs [20]. A wideband LNA consumes between 9 and 30 mW making it very difficult to fulfill a constraint of maximum 10 mW peak power consumption for the overall transceiver. Although several successful designs are recently published

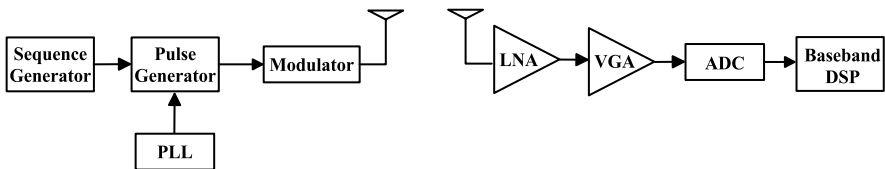


Fig. 2.1 The building blocks of an impulse based UWB transceiver

showing the potential of UWB systems, their power consumption remains too high to be implemented in a “micro-Watt node”. In [20] the total power consumption is around 136 mW at 100% duty cycle. In [21] a power consumption of 2 mW has been reported for the pulse generator only.

2.1.2 Back-scattering for RFID Applications

In the wide arena of low-power architectures, RFIDs represent a good solution when the applications scenario requires an asymmetric network. In this case the “micro-Watt node” needs to transmit data and to receive only a wake-up signal. The required energy is harvested from the RF signal coming from the interrogator. In [22] the interrogator operates at the maximum output power of 4 W, while generating by inductive coupling 2.7 μ W. This power allows a backscattering-based transponder to send On-Off Keying (OOK)-modulated data back to the interrogator in a 12 meters range using the 2.4 GHz ISM band. Unfortunately the limited amount of intelligence at the transmitter side makes this architecture not flexible and only suitable in a highly asymmetric wireless scenario.

There are however other drawbacks in this kind of modulation, mainly shadowed regions. There are two types of such regions. One occurs when the phase of the reflected signal is in opposition with the phase of the RF oscillator. The second occurs due to multiple reflections in an indoor environment. In this situation multiple paths can add destructively at the receiver. Therefore, an increase in the complexity of the receiver is required which can be very severe if the link should be robust enough. Finally the backscattering technique has an increased sensitivity to the fading. The small scale fading observed on the backscattered signal has deeper fades than a conventional modulated signal.

2.1.3 Sub-sampling

The Nyquist theorem has been explored in sub-sampling based receivers in order to reduce the overall power consumption. The power consumption of analog blocks mainly depends on the operating frequency. Applying the theory of bandpass sampling [23], it can be proven that the analog front-end can be considerably simplified reducing the operating frequency. This has the potential to lead to a very low power receiver implementation. Unfortunately due to the noise aliasing, it can be proven that the noise degradation in decibels is:

$$D = 10\text{Log}_{10}\left(1 + \frac{2MN_p}{N_0}\right) \quad (2.2)$$

where M is the ratio between the carrier frequency and the sampling frequency, N_0 is the white noise spectral density and N_p is the Band-Pass Filter (BPF) filtered

version of N_0 . In this sense, the choice of the BPF filter as well as the choice of the sampling frequency become quite critical. Beside this, the phase noise specification of the sampling oscillator becomes quite demanding. Indeed, the phase noise is amplified by M^2 requiring a careful design of the VCO. Accordingly, when interferers are present, a poor phase noise characteristic can degrade the Bit Error Rate (BER) through reciprocal mixing considerably. Consequently, up to now, this architecture has been used mainly in interferer-free scenarios (space applications) [24].

2.1.4 Super-regenerative

Super-regenerative architectures date back to Armstrong, who invented the principle. Despite many years of development, they still suffer from poor selectivity and lack of stability, while having the potential to be low power. Furthermore, they are restricted to OOK modulation techniques only.

In [25] Bulk Acoustic Wave (BAW) resonators are used to reduce the power consumption and to provide selectivity. In spite of achieving an overall power consumption of 450 μ W, it relies on non-standard technologies (BAW resonators), which will increase cost and form factor of the “micro-Watt node”.

In [26] a 1.2 mW super-regenerative receiver has been designed and fabricated in 0.35- μ m CMOS technology. Even though the power consumption is very close to the requirements of a “micro-Watt node”, selectivity is quite poor. Indeed, to demodulate the wanted signal in the presence of a jamming tone placed 4 MHz far from the wanted channel with a BER of 0.1%, the jamming tone has to be no more than 12 dB higher than the desired signal. Generally, to achieve a reliable communication, the receiver should be able to handle interferers which have a power level 40 dB higher than the wanted signal with a BER smaller than 0.1%. This specification is very demanding for a super-regenerative architecture and it requires the use of non-standard components like BAW resonators to achieve a better selectivity.

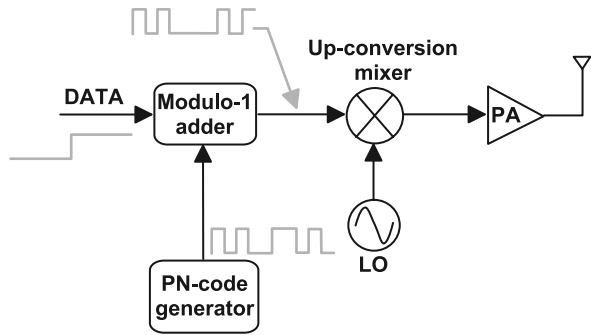
2.1.5 Spread-Spectrum Systems

Any transmission technique in which a Pseudo-random Noise Code (PNC) is used to spread the signal energy over a bandwidth much larger than the information bandwidth is defined as an SS type of transmission. SS techniques are mainly of three types:

- Direct Sequence Spread Spectrum (DSSS)
- Frequency Hopping Spread Spectrum (FHSS)
- Time Hopping Spread Spectrum (THSS)

Sometimes these techniques are combined to form hybrid systems. These hybrid systems are outside the scope of this system given their high complexity. The most widely used systems are of the first two kinds and therefore, this section is restricted to the analysis of both DSSS and FHSS systems.

Fig. 2.2 Schematic block diagram of a DSSS transmitter



Direct Sequence Spread-Spectrum

In DSSS systems, the spreading code is applied to the incoming data. In this way the data symbol is chopped in several parts following a pseudo-random code. Each of these slices within the same symbol period is called a chip. Two quantities are defined in DSSS systems, which are the chip rate $R_c = 1/T_c$ where T_c is the chip duration and the symbol rate $R_s = 1/T_s$, where T_s is the symbol period. The chip rate is an integer multiple of the symbol rate. At every moment, the “instantaneous bandwidth” is equal to the average bandwidth and it is proportional to the chip rate.

A simplified block diagram of a DSSS transmitter is depicted in Fig. 2.2. The same principle is applied on the receiver side where after despreading the data is recovered. To despread the data the receiver must know the PNC sequence and must synchronize in time with it. Therefore, if the receiver does not know the PN sequence, then the received signal will continue to be spread spectrum and the transmitted data cannot be recovered. In this sense a DSSS system is a secure system, which broadens the range of applications in which an ultra-low power node can be used. Another very important parameter in DSSS systems is the so called Processing Gain (PG). The PG is defined as follows:

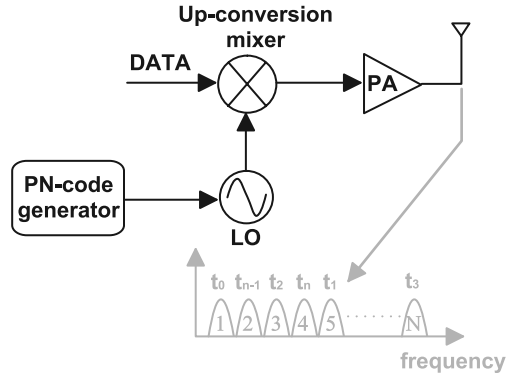
$$G_p = \frac{BW_{ss}}{BW_{info}} = N_c \quad (2.3)$$

where BW_{ss} is the occupied bandwidth after spreading, BW_{info} is the information bandwidth and N_c is the number of chips per symbol period. This parameter has great importance when a DSSS system needs to cope with an in-band interferer usually called a jammer. The effect of interferers will be analyzed later in this section when a comparison between DSSS and FHSS systems will be performed.

Frequency Hopping Spread Spectrum

Differently from DSSS systems, in FHSS systems the spreading code is applied to the frequency domain rather than to the time domain. Therefore, the system hops after a certain amount of time, called dwell time, to another frequency. An FHSS

Fig. 2.3 Schematic block diagram of an FHSS transmitter



system is instantaneously a narrowband system but on the average it is a wideband system.

Important parameters of an FHSS system are the number of channels, the dwell time (T_h) and if the system is a slow hopping or a fast hopping system. The definition of slow hopping or fast hopping is not given in the absolute sense but only in conjunction with the data-rate. A system is considered to be slow hopping if the hopping rate is smaller than the data-rate. When the hopping rate is faster than the data rate the system is called fast hopping.

A simplified block diagram of an FHSS system is given in Fig. 2.3. As in a DSSS system, an FHSS needs a PNC synchronization. To successfully recover the transmitted data the receiver needs to hop coherently with the transmitter. Any FHSS is, therefore, a secure system against intentional jammers trying to steal any information.

DSSS Versus FHSS

In this section the two most common SS techniques are analyzed more in detail and a comparison between them is performed. The reason for such a comparison is driven by the idea to find the optimal SS solution for a power constrained environment. Of course this solution must take into account an interferer scenario composed not only by radios of the same network but also from radios of other standards using the same allocated bandwidth.

Different radio characteristics are further analyzed for both a DSSS radio and an FHSS radio. Those characteristics are the followings:

- Power spectral density and probability of collision
- Susceptibility to the near-far problem
- Radio selectivity
- Robustness to fading conditions
- Robustness to narrowband jammers
- Modulation format and power efficiency
- Acquisition time

A DSSS system is an instantaneously wideband system. Therefore, its transmitted power is spread over a very large bandwidth. Though an FHSS system is on the average a wideband system, instantaneously it operates as a narrowband system. This means that the power spectral density of a DSSS system is lower than that of an FHSS system for the same transmitted power.

In a DSSS system, if two nodes communicate at the same time, they will always interfere with each other. On the other hand, the probability of collision in an FHSS system is the following:

$$P_{\text{coll.,FHSS}} = P_{\text{concurrent-communication}} \times P_{\text{same-frequency}} \quad (2.4)$$

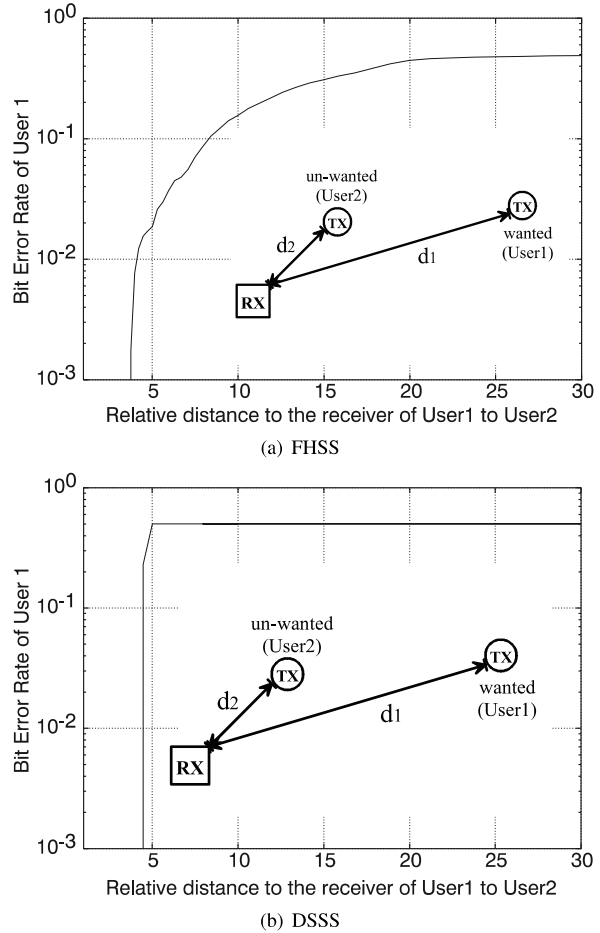
where $P_{\text{concurrent-communication}}$ is the probability that two nodes communicate at the same time (and it is the same for a given network in both DSSS systems and FHSS systems) and $P_{\text{same-frequency}}$ is the probability that two nodes occupy the same frequency bin. This probability is smaller than one and in general it is inversely proportional to the number of available frequency bins in the FHSS system. Therefore, it can be much smaller than one. This translates in a much smaller probability of collision of an FHSS system with respect to a DSSS system.

When a wanted and an unwanted node are communicating at the same time (and on the same frequency bin for an FHSS system) a collision occurs. Depending on the relative received power of the wanted and unwanted node, the data can be retrieved or can be lost because the unwanted node is overwhelming, in terms of received power, the wanted node. This is the so called near-far problem, that mainly appears when the wanted node is much farther than an unwanted node and a collision occurs. Indeed, in this situation the received power of the wanted node can be substantially lower than that of the unwanted node.

Of course the processing gain helps the receiver to distinguish between the wanted node and the unwanted node. For the moment no power control mechanism is supposed. If the wanted node is 3 times further away with respect to the unwanted node then the difference in power equals (see (1.7)) roughly 19 dB. In general DSSS systems use a Binary Phase Shift Keying (BPSK) modulation technique which, for a BER of 1%, requires ideally an $\frac{E_b}{N_0} \simeq 5$ dB where E_b is the energy per bit and N_0 is the noise spectral density. This means, for the above example, that the unwanted signal must be attenuated at least by 24 dB to correctly recover the transmitted data. Supposing now that all the received signals have the same power and considering a processing gain of roughly 12 dB, a maximum of 5 wireless nodes can transmit at the same time without affecting considerably the quality of the link.

FHSS systems, on the other hand, are on the average wideband systems, but instantaneously they are narrowband. Therefore, the susceptibility to the near-far problem is avoided as soon as the data can be transmitted on the next frequency bin and this is not jammed by another interferer. Therefore, while the only way to cope with interferers, for a DSSS radio, is to improve its processing gain, an FHSS system can hop around the problem avoiding it. An increase in the processing gain translates in a much more power hungry digital back-end because of the higher operating speed required. An FHSS system can increase its interferer suppression capability by increasing the number of frequency slots available. This is bounded only by the

Fig. 2.4 Near-far sensitivity comparison between FHSS and DSSS

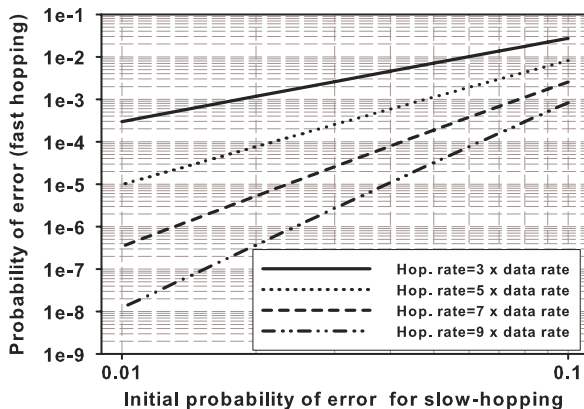


tuning range of the synthesizer, which does not affect power consumption in a first approximation. In Fig. 2.4 two DSSS systems and two FHSS systems are supposed to interfere continuously. The simulation results obtained by using a Simulink model show the clear advantage of an FHSS system over a DSSS system.

Selectivity in an FHSS system is assured by the baseband filter. This means that more nodes can be allocated in the same net. In [27] it is shown that the effective throughput of an FHSS network peaks at a certain number of nodes which is generally larger than in a DSSS system. For example, for the given number of wireless nodes and frequency bins available in the network described in [27], the FHSS network throughput peaks at around 13 nodes. On the other hand, these networks must be placed further away than in the case of a DSSS system. Still, on a single network an FHSS system is more robust than a DSSS system.

SS systems are also known for their capability to cope with fading conditions. A fading condition translates in a very poor SNR, which can be 20 dB below the

Fig. 2.5 Probability of error at variable hopping rate using a majority decision criteria



theoretical SNR calculated in a Additive White Gaussian Noise (AWGN) condition. DSSS suppresses the multipath using again the decorrelation properties of the system. Two PN sequences are uncorrelated only if they are delayed by more than one chip. The delays in a common office environment are in the order of few nanoseconds. This will imply that at low data-rate a very-high processing gain must be applied. Another way is to increase the data-rate but in both cases the digital back-end will suffer from an increase in the power consumption.

An FHSS system copes with the same problem by simply changing the frequency. Because the fading is frequency dependent it is possible by changing frequency to have less adverse fading conditions. This is related to the tuning range and the available bandwidth, more than to power consumption and therefore, it constitutes an advantage of FHSS over DSSS.

The FHSS system has an intrinsic capability to cope against strong narrowband interferers and fading by increasing its hopping rate. Supposing that N hopping channels are available and J out of N are jammed because of strong fading condition or from a large interferer, then if a single bit is transmitted on different channels and a majority decision criteria is used, the probability of error for a given bit is:

$$P_{\text{th}} = \sum_{x=r}^c \frac{c!}{r!(c-r)!} p^x (1-p)^{(c-x)} \quad (2.5)$$

where $p = J/N$, c is the number of bins in which the same bit is sent and r is the number of errors necessary to cause a bit error.¹ In Fig. 2.5 the probability of error versus the initial error probability is shown for different hopping rates. By varying the hopping rate it is possible to decrease by an order of magnitude the BER without increasing the transmitted signal power. In an indoor and interferer crowded scenario this is one of the greatest advantages an FHSS system has.

Another important point of discussion is the modulation format. As it will be shown more into detail in Sect. 2.1.5, an FHSS system generally employs an

¹For example if the hopping rate is three times the data-rate, then 2 errors will cause a bit error. For a hopping rate 5 times the data-rate, 3 errors are necessary to cause a bit error.

FSK modulation while a DSSS system uses a BPSK modulation. The power efficiency of the FSK modulation is larger than that of a BPSK. Therefore, though the BPSK has an advantage over the FSK regarding the robustness against interferers, this advantage is zeroed by the relatively lower power efficiency (see for details Sect. 2.1.5).

As already mentioned, the wake-up time is an important parameter in ultra-low power radios. This wake-up time includes also the synchronization time commonly required in every SS system. When two wireless nodes try to communicate with each other, synchronization of the PNCs must be achieved. At the beginning an offset between the transmitter PN sequence and the receiver PN sequence exists. Therefore, a space of uncertainty exists between the phases of the two PNCs. This space is sliced into “small” pieces called cells. Each cell has a time width so that, when synchronization is achieved, the residual time difference between the transmitter PNC and the receiver PNC code is within half of the chip period or half of the symbol period for a DSSS and an FHSS system respectively. The synchronization algorithm needs to explore those cells in order to find the one for which the phase difference between transmitter PNC and receiver PNC is within the aforementioned range.

It can be proven that the average acquisition time for a SS system, in the case of a serial search synchronization technique with non-coherent detection² is [28]

$$\overline{T_s} = (C - 1)T_{da} \left(\frac{2 - P_d}{2P_d} \right) + \frac{T_i}{P_d} \quad (2.6)$$

where T_i is the integration time for the evaluation of each cell in the time-frequency plane, P_d is the probability of detection when the correct cell is being evaluated, T_{da} is the average dwell time at an incorrect phase cell, C is the total number of cells. This formula can be intuitively explained supposing that the probability of false alarm is zero. In this case the average dwell time at an incorrect phase cell equals the integration time on the cell. The total number of incorrect cells is $C - 1$ and on each one the integration time is spent in the evaluation. Of course, this is just the worst case condition. On the average, there is a certain probability to find the correct cell before all the C cells are evaluated. If the probability of a cell to be the correct one is uniform, and the probability of detecting the correct cell is one (given the fact that the probability of false alarm is zero), then the expression between brackets in (2.6) approaches 0.5. This value makes sense given the fact that an uniform distribution is supposed. This means that, for each cell, there is always a 50% probability that it is the correct one. When the probability of false alarm is non zero, then the average dwell time at an incorrect cell increases. The probability of detection decreases, increasing the term between brackets in (2.6). Therefore, globally the acquisition time increases. The last term in (2.6) reflects the time spent in the evaluation of the correct cell. It is equal to the integration time if

²As it will be proven in Chap. 3 the serial acquisition technique is the most power efficient algorithm for PNC acquisition.

Table 2.1 Summary of the comparison between DSSS and FHSS systems

Radio characteristic	DSSS	FHSS
Power spectral density	Low	High
Probability of collision	High	Very low
Near-far robustness	Low	High
Selectivity	Medium	High
Fading robustness	Medium	High
Narrowband jammer robustness	Medium-low	Very high
Modulation power efficiency	Medium	High
Acquisition time	High	Low

the probability of detection is 1 and larger if there is a chance to skip the correct cell. For this reason, it has to be inverse proportional to the probability of detection.

Now, assuming that no frequency uncertainty is present, there will be a time misalignment between the two PN sequences at the transmitter and receiver side equal to ΔT_i . Therefore, while for a DSSS the system has to be synchronized within $\pm T_c/2$, an FHSS system needs to be synchronized within $\pm T_s/2$.³ Due to the fact that in a DSSS system the processing gain is related to the ratio between the chip rate and the symbol rate, the chip period is at least an order of magnitude smaller than the symbol period. As a result, the number of cells that must be evaluated in a DSSS system is considerably larger than in an FHSS system. From (2.6) the mean DSSS synchronization time is larger than in the case of an FHSS system. This will increase the wake-up time and therefore, the overall system power consumption.

A summary of the comparison between FHSS and DSSS systems is given in Table 2.1. It is clear that the FHSS technique presents a clear advantage over the DSSS technique for most of the characteristics listed in Table 2.1. For this reason, it is more suitable for power constrained, indoor, interferer crowded scenario like the one foreseen in wireless sensor networks.

Modulation Formats

Several modulation formats can be used in digital communications. The relative implementation complexity of various modulation schemes is depicted in Fig. 2.6. Given the power constrained environment it is important to select a low complexity modulation technique. Therefore, three modulation formats are analyzed more in detail:

³The remaining part of the synchronization consists in what is generally called tracking. During tracking, the phase difference between the two PNCs is reduced to virtually zero from a closed-loop system (like a PLL). This system also tracks any instantaneous variation of the phase of the transmitter PNC in order to assure a constantly aligned PN sequences between the transmitter and the receiver when the nodes are communicating with each other.

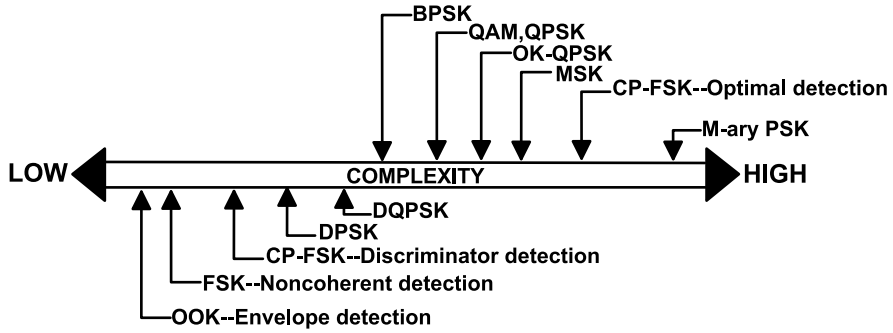


Fig. 2.6 Relative complexity of various modulation scheme (adapted from [29])

- OOK with envelope detection
- FSK with non-coherent detection
- BPSK

Coherent-Phase Frequency Shift Keying (CP-FSK), Differential Phase Shift Keying (DPSK) and Differential Quadrature Phase Shift Keying (DQPSK) are derivatives of those formats and therefore, though their complexity is not high, they will not be further analyzed in this book. OOK, FSK and BPSK modulations cover also all of the most common types of modulation formats. In fact OOK is a type of amplitude modulation, FSK is a frequency modulation and BPSK is a phase modulation. A comparison between these modulation schemes can be done based on an ideal required SNR for a given BER,⁴ signaling speed, robustness against Continuous Wave (CW) interferers, and robustness in a Rayleigh fading channel. The summary is shown in Table 2.2. The performances of OOK and FSK in a AWGN environment are very similar while the BPSK modulation has roughly 4 dB better performance. In a Rayleigh fading environment this gain reaches roughly 6 dB. When interferers are present, as it will happen in the 915 MHz and 2.4 GHz ISM bands, then the OOK modulation happens to be a very weak scheme requiring 5 dB more SNR than FSK and almost 10 dB more than BPSK. Therefore, FSK is more suitable than OOK in an interferer crowded scenario. BPSK has a 4 dB advantage in an interferer dominated scenario over the FSK modulation.

Every frequency modulated signal is a truly constant envelope signal, while BPSK modulated signals contain some amplitude modulation in their modulated envelope. Therefore, while FSK signals can be amplified by a Power Amplifier (PA) operating near the saturation level, BPSK signals require a 3 to 6 dB back-off from this level. This is necessary in order to eliminate the spectral regrowth, which will cause adjacent channel interference. This translates in a smaller power efficiency and therefore, the gain of BPSK modulated signal over FSK signal is practically more than compensated by this drawback.

⁴With the word “ideal” here it is supposed that the channel is AWGN.

Table 2.2 Performance comparison of some low complexity modulation formats for a $BER = 10^{-4}$ (adapted from [29])

Modulation	Speed [bitHz/s]	$\frac{E_b}{N_0}$ (AWGN) [dB]	$\frac{E_b}{N_0}$ ($S/I = 10$ dB) [dB]	$\frac{E_b}{N_0}$ (Rayleigh) [dB]
OOK-Envelope det.	0.8	11.9	20	19 ^a
FSK-Non-coh. ($m = 1$ ^b)	0.8	12.5	14.7	20
BPSK	0.8	8.4	10.5	14

^aWith optimum variable threshold^b m = modulation index

As stated in Sect. 2.1.5 this is a big advantage for an FHSS system over a DSSS system. Implementation of an FSK modulation format on an FHSS system is trivial given the fact that frequency hopping is very close to a frequency modulation scheme.

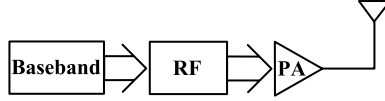
2.2 Optimal Data-Rate

The aim of this section is to find the data-rate, which minimizes the average node power consumption. The transceiver is composed by a receiving section and a transmitting section.

The receiver sensitivity depends on the Noise Figure (NF), noise bandwidth and the required SNR of the demodulator. The NF depends on the receiver architecture and technology used and in an asymmetric scenario can be considered to be smaller than 10 dB. For a chosen modulation scheme and a certain desired BER, the required SNR is fixed, apart from implementation losses in the demodulator. The only parameter left is the noise bandwidth which ultimately affects the system data rate.

A transmitter can be partitioned into three domains in terms of power consumption. This partitioning is shown in Fig. 2.7. The baseband part has a power consumption, which is proportional to the data-rate. The RF portion, is needed to upconvert the baseband information to the wanted high frequency band. Therefore, its power consumption depends on the operation frequency and it is independent of the data rate. Finally, the PA section is needed to transmit the information over the medium. Its power consumption depends on its efficiency and on the required transmission range. The efficiency, is strictly related to the modulation format used and it is optimal for constant envelope formats like FSK. Generally the most power hungry blocks are the PA and the RF section which consists of the synthesizer and a mixer for up-conversion. The first two domains are what is defined later in this section as pre-PA domain.

Fig. 2.7 Transmitter power domains



In general a transmitter first needs a time T_{wu} to wake up, then T_{tx} to transmit, and after transmission it remains T_{idle} in the idle mode till the next transmission cycle is started. So, the time between two consecutive transmissions, T , equals $T_{wu} + T_{tx} + T_{idle}$. The duty cycle of the system, “ d ” can be defined as the ratio between the time required to transmit the data and the time between two consecutive transmissions.

Several parameters are involved in the derivation of the optimal data-rate. Some parameters are fixed. Other parameters, although can be varied and optimized for low power, are also considered fixed. Finally the variable we need to optimize is the data-rate. Parameters, which are fixed are the following:

- $\frac{B_{noise}}{B_{data}}$
- N_0
- $\frac{E_b}{N_0}$
- $L_{path,nat}$

where B_{noise} ⁵ is the noise bandwidth B_{data} ⁶ is the data bandwidth, $L_{path,nat}$ represents the path losses due to propagation expressed in natural units.

Parameters which can be optimized for low power but are considered fixed in this discussion are the following:

- L_{pack}
- P_{diss}
- P_{idle}
- T_{wu}

where L_{pack} is the data packet length, P_{diss} is the power consumption of the remaining transmitter circuitry during wake up and transmission excluding the PA and P_{idle} is the power consumption in the idle mode.

It is possible to suppose that the duty-cycle of the network is constant [30], or that the time between two consecutive transmissions is constant [31]. These two cases will be further analyzed.

⁵The 2-FSK noise bandwidth can be approximated by the Carson rule as $B_{noise} = 2(\Delta f + f_m)$ where $f_m = \frac{2}{T_s}$ with $\frac{1}{T_s}$ the data rate and Δf the frequency deviation.

⁶For a 2-FSK modulated signal it equals four times the data rate.

2.2.1 Constant Duty-Cycle

The average power consumption of the transmitter node can be approximated by the following equation:

$$P_d = P_{tx} \frac{T_{tx}}{T} + P_{diss} \frac{(T_{tx} + T_{wu})}{T} + P_{idle} \frac{(T - T_{tx} - T_{wu})}{T} \quad (2.7)$$

where P_{tx} is the PA power required for the transmission of data. The transmission time depends on the packet length L_{pack} and on the data rate “ D ”:

$$T_{tx} = \frac{L_{pack}}{D} \quad (2.8)$$

The required transmitted power can be written as

$$P_{tx} = N_0 \times \frac{B_{noise}}{B_{data}} \times \frac{E_b}{N_0} \times NF \cdot D \cdot L_{path,nat} = K \times D \quad (2.9)$$

Given the fact that both B_{noise} and B_{data} are proportional to the data rate, their ratio is constant and the transmitted power P_{tx} is direct proportional to the data rate D via the constant K in (2.9). From these considerations, (2.7) can be re-written in the following way:

$$P_d = (K D + P_{diss}) \times d + P_{idle}(1 - d) + (P_{diss} - P_{idle}) \times d \frac{T_{wu}}{L_{pack}} D \quad (2.10)$$

where the duty cycle “ d ” is a constant in this discussion.

From (2.10), it can be seen that for a fixed transmission distance ($\propto P_{tx}$), the power consumption can be reduced by reducing the duty cycle, the data-rate, and by making the wake-up time small compared to the transmission time. This last requirement becomes difficult to achieve in SS systems at high data rates due to PNC synchronization. Therefore, reducing the data rate will help to relax the wake-up time for a given $\frac{T_{wu}}{T_{tx}}$. The transmitter average power consumption as a function of the data-rate for different duty-cycles is plotted in Fig. 2.8. At high data rates, the average power consumption is dominated by the transmitted power. At data rates below a threshold value the average power consumption is dominated by the pre-PA power. This threshold value depends on the pre-PA power dissipation and it is lower for lower values of the pre-PA power dissipation.

From Fig. 2.8, when the pre-PA power dissipation (P_{diss}) is 2 mW, this threshold value is around 100 kbps, while at pre-PA power of 10 mW it is located around 1 Mbps. At higher data rates, the wake-up time has to decrease considerably to keep the contribution to the average power consumption negligible. Therefore, from the previous analysis it is possible to conclude that a good strategy toward the reduction in the average transmitter power consumption consists of reducing the data-rate and decreasing the synchronization time for a given node duty-cycle.

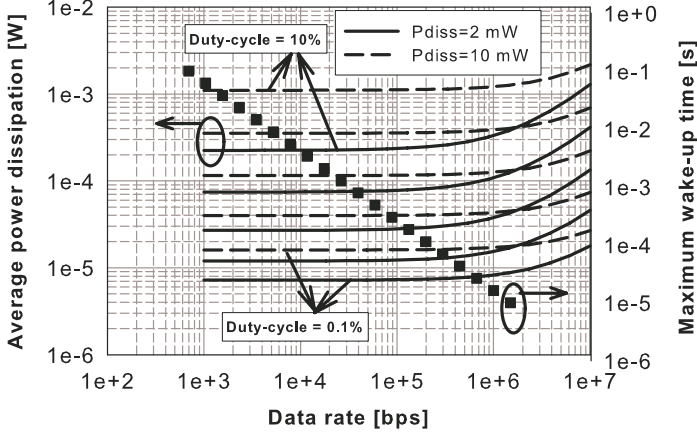


Fig. 2.8 Average transmitter power consumption and maximum wake up time for different value of duty cycle, as a function of the data rate ($L_{\text{pack}} = 1000$ bits, $NF = 10$ dB, $\frac{E_b}{N_0} = 20$ dB, carrier frequency = 915 MHz, $P_{\text{idle}} = 10 \mu\text{W}$) [30]

2.2.2 Constant Time Between Two Consecutive Transmissions

Some applications may require a fixed time between two consecutive transmissions. In this case the time “ T ” is constant while the duty cycle “ d ” varies decreasing by increasing the data rate D . For example, for a temperature sensing inside an apartment a fixed interval between two consecutive transmissions may be sufficient. Equation (2.10) can be rewritten now with the time difference between two consecutive transmissions as a variable instead of the duty cycle:

$$P_d = K \frac{L_{\text{pack}}}{T} + P_{\text{diss}} \left(\frac{L_{\text{pack}}}{D \cdot T} + \frac{T_{\text{wu}}}{T} \right) + P_{\text{idle}} - P_{\text{idle}} \frac{L_{\text{pack}}}{D \cdot T} - P_{\text{idle}} \frac{T_{\text{wu}}}{T} \quad (2.11)$$

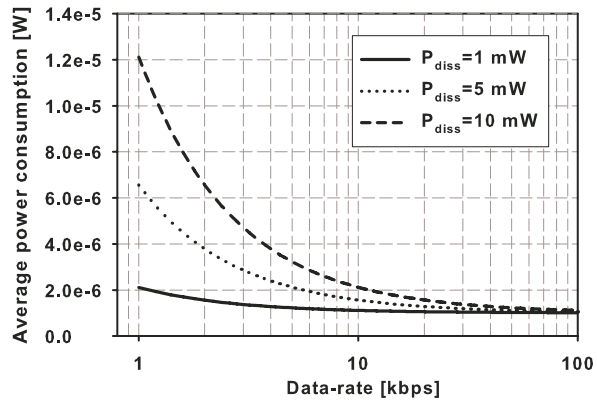
and for T sufficiently large it can be approximated as follows:

$$P_d \approx K \frac{L_{\text{pack}}}{T} + P_{\text{diss}} \left(\frac{L_{\text{pack}}}{D \cdot T} + \frac{T_{\text{wu}}}{T} \right) + P_{\text{idle}} \quad (2.12)$$

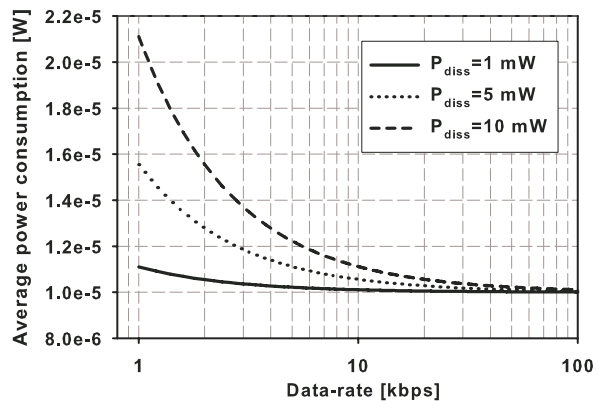
The simulation results are given in Fig. 2.9 as a function of the data rate for a given (constant) value for T and T_{wu} . As it can be seen from Fig. 2.9, increasing the data rate will make the average power consumption smaller and smaller. A limit is dictated by the idle power. This is easily understandable by looking at (2.12). When the idle power is $1 \mu\text{W}$, a data rate between 1 and 10 kbps is sufficient to not spoil the average power consumption for P_{diss} ranging between 1 and 10 mW. In this way, in the worst condition ($P_{\text{diss}} = 10$ mW) the average power consumption is determined in equal parts by the idle power and the power used during transmissions. When the idle power increases to $10 \mu\text{W}$, then the data-rate can be relaxed down to around 1 kbps in all the cases.

For high data-rate the pre-PA power consumption can become a strong function of the data-rate and the term P_{diss} cannot be any longer considered constant but

Fig. 2.9 Average power dissipation versus data rate ($L_{\text{pack}} = 1000$ bits, $T_{\text{wu}} = 500 \mu\text{s}$, $NF = 10$ dB, $\frac{E_b}{N_0} = 20$ dB, carrier frequency = 915 MHz, $T = 300$ s) for $1 \mu\text{W}$ and $10 \mu\text{W}$ idle power dissipation



(a) $P_{\text{idle}} = 1 \mu\text{W}$



(b) $P_{\text{idle}} = 10 \mu\text{W}$

will be a function of the data-rate. In the newest technologies it is reasonable to say that below 100 kbps the transmitter baseband part will consume much less than the RF part and, therefore, the approximation can be considered valid. It should be noticed that in this discussion it has been neglected that, on the receiver side, the power consumption is a function of the data-rate for all the baseband domain. Nevertheless, looking at the receiver, it is a reasonable choice as optimum data-rate the lowest possible, which fulfill the aforementioned considerations. Indeed at the receiver side several analog blocks (like Voltage Gain Amplifier (VGA), filters and ADC) work at baseband frequency and their power consumption is proportional to the data-rate.

While it seems that an idle power in the μW range is too high, it should be noticed that most probably a wake-up type of radio will be used. Some circuitry will be kept on listening the channel for an incoming transmission especially in asynchronous networks. Therefore, a power budget for this circuitry between 1 and $10 \mu\text{W}$ is a good choice. The considerations regarding the synchronization time for the case of constant duty-cycle networks do apply also to this case. Moving towards lower

data rates helps in relaxing the synchronization constraints. From all the previous considerations it is possible to conclude that a data rate between 1 and 10 kbps is sufficient to optimize the average node power consumption.

2.3 Transmitter Architectures

Generally, the transmitter part is under-estimated in terms of complexity and number of possible system and circuit trade-offs with respect to the receiver part. Any choice on the transmitter side will have a great impact on the receiver specifications as well. The transmitter topology of choice is a result of various trade-offs impacting the maximum level of unwanted frequency components, efficiency of the PA, linearity and maximum power consumption. The transmitter performs some operations before radiating the signal toward the receiver, which can be summarized as follows:

- Modulation
- Up-conversion
- Power amplification

Modulation has been already discussed in the previous section. It has been pointed out that at low data rates the information bandwidth is not a problem. Therefore, the choice has to be directed towards a power efficient modulation scheme rather than a bandwidth efficient scheme. Consequently, a constant envelope modulation scheme like FSK is preferable.

The up-conversion is the action with which the baseband signal is shifted to the wanted frequency band (915 MHz or 2.4 GHz) before transmission. The power amplification needs to increase the power of the radiated signal in such a way that the radiated signal in the worst distance conditions (for example 10 meters indoor) reaches the receiver at a power level equal or above the sensitivity level.

Transmitter architectures can be grouped in three main categories:

- Direct conversion
- Two-step conversion
- Offset PLL

2.3.1 Direct Conversion

A direct conversion transmitter is plotted in Fig. 2.10. The modulated data is directly up-converted to the wanted band. Therefore, the oscillator is running at the carrier

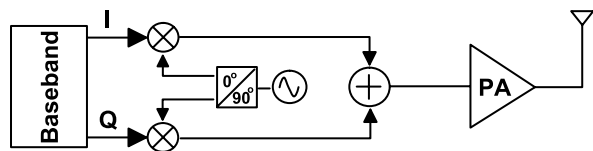


Fig. 2.10 Direct conversion transmitter

frequency. In this type of transmitters can be also included the direct modulated VCO type of transmitters. In this case the baseband data directly modulates the VCO frequency following the data stream. A straightforward implementation of this direct up-conversion scheme employs an FSK modulation type. Direct VCO modulated transmitters merge the modulation and the up-conversion phases into one single operation.

The simplicity of this architectures makes it attractive for a high level of integration, which means lower costs and lower power consumption.⁷ Unfortunately, a big drawback of this architecture is a phenomenon called “Oscillator pulling”. If the output of the PA and the oscillator frequency are very close to each other in frequency, the oscillator is heavily disturbed by the noise coupling back from the PA output. Even a noise level 40 dB below the oscillator level can cause an enormous amount of disturbance on the oscillator output. For this reason generally the two frequencies must be far apart and mostly uncorrelated. Several solution are possible to accomplish this result:

- The up-conversion signal is obtained by an integer division of the oscillator signal.
- The up-conversion signal is obtained by a non-integer division of the oscillator signal.
- The up-conversion signal is obtained by mixing two non divisible oscillator frequencies.

The first option is the most simple one to implement at hardware level. The PA output and the oscillator output are far apart, but the two frequencies are still correlated. This can give still some pulling especially at high noise level. The second option makes the two signals far apart and well uncorrelated. Unfortunately a non integer division requires complex hardware. For example, a multiplication by two followed by a divide-by-three stage can accomplish this result. The last option gives the best results but at the expense of two oscillators, one mixer and one BPF, required to suppress the unwanted harmonics generated from the mixing process.

The most power efficient solution, therefore, is the first option if a good degree of isolation can be guaranteed between the PA output and the oscillator in order to avoid a high level of noise injection in the oscillator.

2.3.2 Two-Step Conversion

A block diagram of a two-step conversion transmitter including the frequency spectrum at various points in the chain is shown in Fig. 2.11. This architecture eliminates the problem of VCO pulling by splitting the up-conversion into two phases.

⁷External components require in general matching to 50 Ω . This means that the stage preceding the external component must be able to drive a 50 Ω impedance. Such a low impedance level will cost a high current consumption.

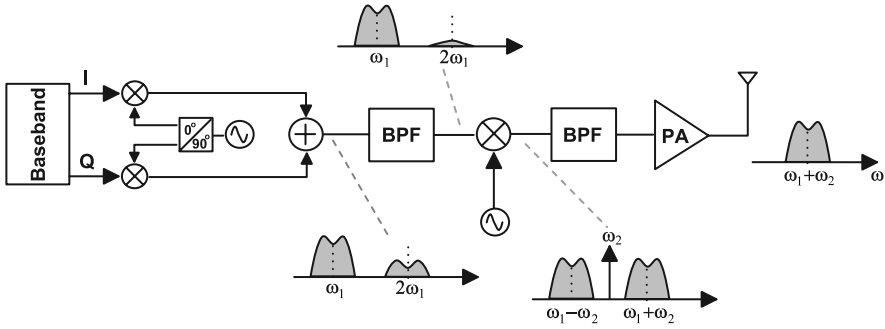


Fig. 2.11 Two-step conversion transmitter (adapted from [32])

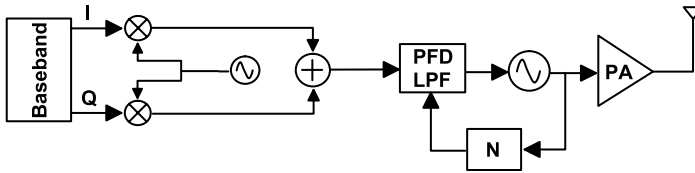


Fig. 2.12 Offset PLL architecture

An advantage of this architecture is that the quadrature up-conversion to the so called intermediate frequency is performed at a lower frequency compared to the direct up-conversion scheme. This greatly improves the matching of the I - Q signals. The main drawbacks of this architecture are first the increased amount of hardware needed and second the BPF before the PA. This BPF needs to attenuate the unwanted sideband more than 40 dB. Therefore, it generally requires an expensive and power hungry off-chip component.

2.3.3 Offset PLL

The schematic block diagram of an offset PLL based transmitter is depicted in Fig. 2.12. The main restriction of this kind of transmitter architecture is that it can be used only with constant envelope type of modulated signal. In this architecture, the PLL acts as a narrowband filter, rejecting all the high frequency noise coming from external sources. Therefore, the high frequency noise is generally determined by the VCO noise but this also happens in all the previously mentioned transmitter architectures. Unfortunately, the divider in the PLL chain, which is used to ease the design of the PFD can cause some severe drawbacks. Any phase modulation at the input of the PLL is amplified by a factor N in amplitude at the output of the PLL. The same happens to any change in the frequency of the baseband VCO. This means that if more channels are used, the baseband synthesizer must have a very fine tuning mechanism, which in general means longer settling times and higher average

power consumption. Therefore, this technique tends to be quite cumbersome to use especially in SS systems.

2.4 Receiver Architectures

The receiving part of a transceiver generally can consume a large amount of power if not optimized for a power constrained environment. Several trade-offs are also present in the choice of a suitable architecture for the receiver. The level of integration is an important aspect, which reduces the costs by eliminating bulky external components while also achieving a decrease in the peak power consumption. The level of achievable integration depends in general by a combination of receiver specifications and receiver architecture. Three main receiver architectures are generally used in modern wireless communication:

- Zero-IF
- Super-heterodyne (or heterodyne)
- Low-IF

2.4.1 Zero-IF

The zero-IF architecture has always been considered very suitable for integration. A schematic block diagram is depicted in Fig. 2.13. Unfortunately several issues can make cumbersome its implementation at transistor level. The first major drawback is the so called DC offset. The signal is down-converted to baseband in a single step. Therefore, any DC component is within the signal bandwidth. Now the self-mixing effect of the oscillator frequency via capacitive coupling to the LNA input creates a DC component at the input which amplified by the LNA can reach the millivolt level. If the gain after the LNA is big enough, the DC component so generated can saturate the receiver. The effect is that the weak signal (which is still in the hundreds of microvolt range) cannot be detected properly.

A simple way to eliminate this problem is again to choose correctly the modulation type in such a way that almost no signal energy is placed at DC. In this way, a simple HPF with a corner frequency of few kilohertz can eliminate the offset. Most of the modulation formats, unfortunately, contain much energy around DC. On the other hand, an FSK modulated signal with a large modulation index ($m > 1$) has almost no energy placed at DC. This is another point in favor of architectures that

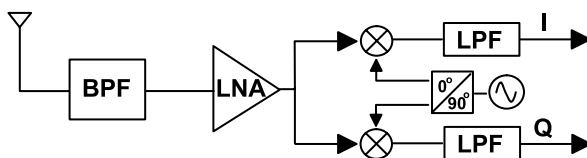


Fig. 2.13 Zero-IF architecture

can easily implement such kind of modulation format within a frequency diversity scheme like it happens for FHSS systems.

Second order distortion is also problematic in zero-IF receivers. A second order distortion in the LNA translates two closely spaced strong interferers to a lower frequency before the mixing process. This component passes through the mixer with finite attenuation due to imperfections in the mixer I - Q matching or Local Oscillator (LO) duty-cycle imperfections. This can again pose a problem to the receiving stage saturating it.

I - Q mismatch between the two quadrature mixers due to the presence of parasitics and due to the high operating frequency is also a big issue. Careful layout is generally needed, especially at high frequencies.

Lastly, a problem which depends on the technology used can come from the noise. For example CMOS transistors are affected by the so called flicker noise. The signal after the LNA is still generally quite low⁸ and therefore, very sensitive to noise. Flicker noise is dominant at low frequencies, which is exactly the frequency region where this architectures tends to directly translate the wanted signal to. Looking at [33] it is possible to see that the flicker noise corner frequency increased from around 1 MHz at 1.2 μm channel length to more than 100 MHz for 30 nm channel length. This means that technology scaling, while helping in reducing the power consumption especially in the digital domain gives more and more drawbacks when an analog block needs to be designed. This is a point which must be greatly taken into account when conceiving an architecture for ultra-low power wireless nodes in CMOS technology.

2.4.2 Super-heterodyne

A schematic block diagram of a super-heterodyne receiver architecture is depicted in Fig. 2.14. The heterodyne principle first down-converts the signal to an intermediate frequency called IF, and after a BPF and a further signal amplification it down-converts the IF signal to baseband. In the case of digital modulation the I and Q signal components are generated in the latest down-conversion stage.

This architecture alleviates some common drawbacks seen in the zero-IF architecture. For example the DC offset coming from the first two stages is filtered out by the BPF, while the one of the last stage is negligible thanks to the high gain of the first two stages. Because the I - Q signals are generated only in the last stage where the frequency is lower, the I - Q mismatch can be easily controlled and reduced to a very low level. Finally, this architecture has a very good selectivity achieved by splitting the filtering among different stages at progressively lower frequency.

Despite all these advantages, the heterodyne architecture suffers of some major drawbacks as well. The biggest problem is the image rejection problem. This issue

⁸Considering a required sensitivity level of -76.5 dBm at 2.4 GHz, then the signal at the output of an LNA with gain equal to 15 dB is generally smaller than 1 mV rms.

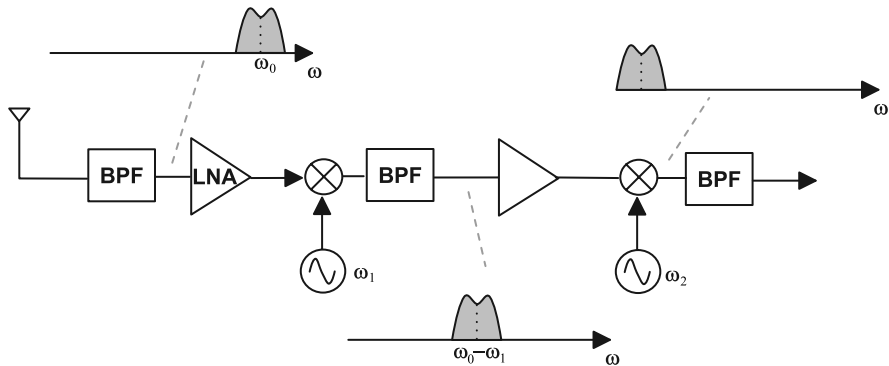
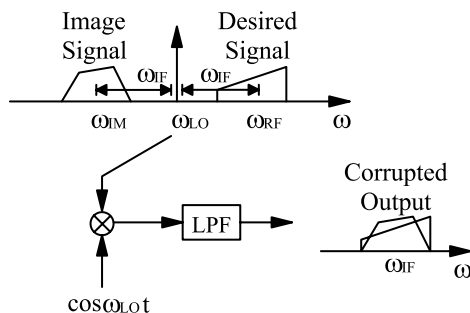


Fig. 2.14 Super-heterodyne architecture

Fig. 2.15 Image frequency problem



can be easily understood noting that all the signals at a distance equal to ω_{IF} from the LO frequency will be down-converted to the same IF frequency. This is illustrated in Fig. 2.15. To suppress the image frequency, an image reject filter is often used. The choice of the IF is not trivial and it entails a trade-off between sensitivity and selectivity. Indeed, to reduce the image noise, the IF frequency has to be chosen as large as possible. On the other hand, this will increase the constraints on the band selection filter coming after the first mixer. The higher the IF frequency, the greater the required Q for a given attenuation.

Generally, the image reject filter precedes the mixer and it is realized using external components. This translates in a worse transceiver form factor and it requires the LNA to drive a $50\ \Omega$ impedance. This, in turn, translates in higher power consumption due to a larger required bias current.

2.4.3 Low-IF

The low-IF architecture is closely related to the zero-IF topology and it tries to minimize the major drawbacks of the zero-IF topology by operating near the DC but not at DC. In a low-IF topology, the problem of the image suppression, which

burdens the heterodyne receiver, can be shifted to the IF stage. In this way, given the lower operating frequency the high-frequency BPFs can be integrated given the lower required Q .

On the other hand, this topology tends to shift the demanding specifications to the ADC preceding the Digital Signal Processor (DSP) in which the final demodulation is performed. This can be seen as a reasonable way to cope with the power reduction problem. Indeed, the ADC converter, is a mixed signal block in which lots of functions are anyhow performed in the digital domain. Therefore, it is more prone to power scaling with technology advances than a purely analog block.

Concluding, the low-IF topology is a good alternative to the zero-IF topology especially when problems like flicker noise and DC offset become a show stopper for further power reduction.

2.5 Conclusions

Though ultra-low power wireless nodes have very tight constraints in terms of power consumption, they should allow for a robust wireless link in the harsh indoor environment. Moreover, because it is necessary to offer to the end user a very low cost solution, ISM bands are generally used because they are license free. This option, however, presents the inconvenience of a very interferer crowded environment for the wireless radio.

It has been shown, in this chapter, that to cope with those strong non-idealities, while keeping the power consumption very low, a combination of modulation format, transmitter and receiver architectures and wideband techniques can be used. It has been proven, indeed, that spread spectrum techniques are an optimal choice to have a robust wireless link. Among several SS techniques, it has been shown that an FHSS system can offer a very robust wireless link while having the potential to be low power especially if it is combined with an FSK modulation. The possibility to trade between hopping rate and transmitted power is an unique advantage of this system, which allows for a not negligible reduction of the transmitted power without affecting the reliability of the link.

The choice of the data-rate affects also the average power consumption of the wireless node. It has been shown that increasing the data-rate can help in reducing the average power consumption. On the other hand, it has been demonstrated that it is useless to increase the data-rate above a certain level dictated by the idle power, because at this point the average power consumption is always dominated by the idle power. Increasing the data-rate above this threshold level, however, costs an increase in the receiver power consumption because several baseband blocks have a bandwidth, which is a function of the data-rate.

On the transmitter side, it has been proven that a single up-conversion scheme is the most appropriate choice to minimize the average power consumption of the wireless node. On the receiver side two options are foreseen. The best option in terms of integrability and power consumption is a zero-IF architecture. However, especially in the newest CMOS technology, flicker noise and other second order

effects, can cause severe difficulties in its physical implementation. For this reason, a low-IF topology, though it can be more power hungry, results in a good choice when the zero-IF becomes too problematic to be implemented at transistor level.

Concluding, this chapter has found the optimal system level choices for an ultra low power wireless node for wireless sensor area networks: an FHSS system with Binary Frequency Shift Keying (BFSK) modulation and a data-rate below 10 kbps, a single up-conversion transmitter architecture and a zero-IF (or a low-IF) receiver architecture. The following chapter focuses on FHSS systems and the limitations existing in the state-of-the-art solutions in terms of power consumption.

Chapter 3

FHSS Systems: State-of-the-Art and Power Trade-offs

In the previous chapter it has been shown that an FHSS system constitutes a good compromise between link robustness and power consumption. The simplicity with which an FSK modulation technique can be applied to this system, makes it very attractive for an ultra low power wireless node. In fact it has been shown that FSK modulation with non-coherent reception is very simple from a hardware point of view, very robust in a interferer crowded scenario and that it is a power efficient modulation format. It is of course not efficient from a bandwidth point of view, but given the low data rates, this does not constitute a real problem. Moreover, the use of a large modulation index helps in relaxing some very common implementation problems of a zero-IF receiver like flicker noise and DC offset. The possibility to use a zero-IF receiver is an important advantage because of its high level of integration. For these reasons, this chapter analyzes more in detail the FHSS architecture focusing on synchronization techniques, state-of-the-art FHSS systems, and FH synthesizer techniques. Finally, a power model for the two most common FH synthesizer architectures is described and a lower bound to the minimum power consumption for a given set of specifications is derived.

3.1 Synchronization

In general some synchronization is required between the transmitter node and the receiver node of a communication link before data can be sent and received reliably. Synchronization can take place in the frequency and/or in the time domain. Different kinds of synchronization domains can be identified:

- Carrier synchronization
- Code synchronization
- Symbol synchronization
- Frame synchronization
- Network synchronization

Carrier synchronization is the procedure by which the receiver replicates the transmitter carrier with the same frequency and generally, but not necessarily, with

the same phase. This is the first necessary step in the synchronization procedure.

To work properly, a FHSS system and more in general any SS system, requires that the locally generated PN sequence is synchronized with the transmitter PN sequence. This procedure is called code synchronization.

In order to make a decision regarding a transmitted symbol the receiver must know the epochs at which the symbol starts and ends. The estimation of these time instants is called symbol synchronization.

When the transmitted signal is highly structured, periodic timing is required between different frames containing a certain number of symbols. This process is called frame synchronization.

Finally, at the top of the hierarchy there is the network synchronization. It consists in methods and techniques to distribute and create a common timing reference to a number of nodes constituting the network.

The following paragraphs mainly focus on the code synchronization. In general, code synchronization takes much more time than carrier or symbol synchronization. Moreover, given the low complexity in the structure of the transmitted signal, frame synchronization is not an issue. Finally, if an asynchronous network is used, network synchronization is not required.¹

The code acquisition process is generally carried out in two steps:

- Code acquisition
- Code tracking

Code tracking is generally carried out by a closed loop system like a PLL and is generally much faster than code acquisition. Therefore, the following paragraphs deal with code acquisition techniques.

The uncertainty in the synchronization of these PNCs, derives from various sources. The time uncertainty is generally due to the following reasons:

- Propagation delay
- Shifts between TX and RX reference clocks
- Phase difference in the PN sequences between TX and RX

The main source of frequency uncertainty is the Doppler shift coming from the fact that the two nodes (TX and RX) have a non-zero relative velocity. Given the fact that WPANs can be considered mostly stationary, the frequency uncertainty is negligible and only a time uncertainty is considered.

Given the received signal and the locally generated code replica, the receiver needs to find the correct position in time of the locally generated replica code, which aligns with the received code. Each relative position between the two codes is called a cell. The uncertainty region is defined as the total number of cells to be searched.

¹Having an asynchronous networks will reduce the on time of the receiver saving therefore power. In a synchronous network all the nodes have to regularly wake up to synchronize their timing reference. This is not optimal from a power point of view. An asynchronous network can therefore be more efficient in terms of power consumption than a synchronous network.

The procedure used to search the uncertainty region is called the search strategy. The search strategies can be grouped into three categories [34]:

- Stepped serial search
- Matched filter acquisition
- Two-level acquisition

3.1.1 Stepped Serial Search

The schematic block diagram of the stepped serial acquisition system is shown in Fig. 3.1. In this algorithm the uncertainty cells are analyzed in a serial way starting from an initial epoch. For each cell the algorithm needs to perform a correlation between the locally generated PNC and the received PNC. In a stepped serial search, this operation is performed by the active correlator shown in Fig. 3.1 on a chip-by-chip basis. For an FHSS system, for analogy with a DSSS system, a chip is defined as the time spent by the system on a given frequency bin. This requires a time equal to MT_h , where M is the number of hops per integration interval, and T_h is the hopping time.

The output of the active correlator drives a comparator, which compares the correlator output with a pre assigned threshold. If the comparator threshold is not exceeded, then the phase of the locally generated PNC is shifted to the next cell. This process is repeated until the threshold is exceeded and an acquisition signal is obtained.

The most important parameter for an acquisition algorithm is its speed. A brief synchronization is important in order to reduce the wake up time of the wireless node. With brief in this book it is meant that the synchronization time must be small compared to the transmission time in order to not degrade the average power consumption.

As any kind of searching algorithm, there is always a non-zero probability of a false alarm. This means that even though the algorithm detects a synchro cell, in reality synchronization might not be achieved. In general this false alarm probability needs somehow to be minimized. It is supposed that the tracking phase can start after three hits are detected on the given cell. It can be proven [34] that the aver-

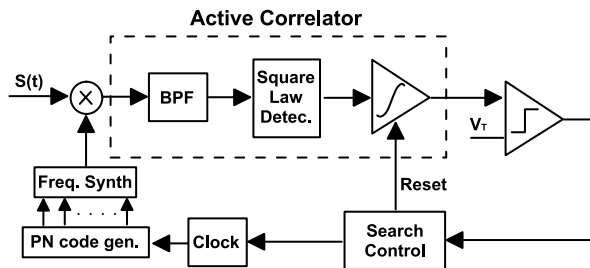


Fig. 3.1 Stepped serial search scheme adapted from [34]

age acquisition time for a stepped serial search algorithm is given by the following relation:

$$T_{\text{acq}} = (nN_c + 2)MT_h \quad (3.1)$$

where n is the number of cells examined per chip, and N_c is the maximum expected delay, expressed in number of cells, between the locally generated and the desired signal code phases.

Indeed, if the algorithm examines n cells for each chip and the maximum delay between the locally generated code and the received code is N_c , then the uncertainty region has $n \times N_c$ cells. For each of these cells a correlation on M chips needs to be performed serially by the active correlator. This requires as previously stated MT_h seconds. Moreover, on the synchro cell the active correlation will be performed two additional times, which justifies the factor 2 in (3.1).

3.1.2 Matched Filter Acquisition

The acquisition time can be decreased by replacing the active correlator with a matched filter correlator often called a passive correlator. A number of M different frequencies is correlated at the same time greatly improving both the reliability and the acquisition speed. A schematic block diagram of this acquisition scheme is shown in Fig. 3.2.

For each cell the time required to correlate the received PNC with the locally generated PNC is T_h and the number of cells in the uncertainty region is nN_c . When the synchro cell is detected a time equal to MT_h is required in order to allow the

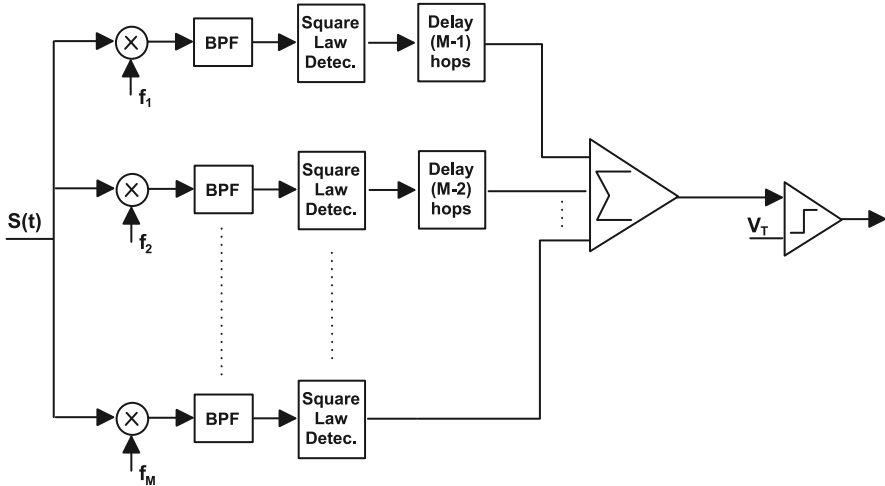


Fig. 3.2 Passive correlator adapted from [34]

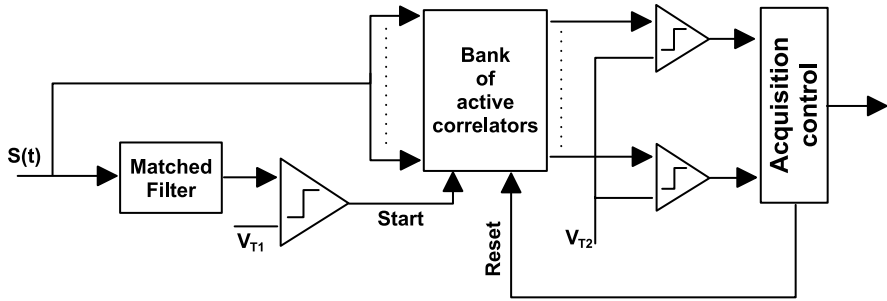


Fig. 3.3 Two-step acquisition scheme adapted from [34]

M correlators to output their result. Therefore, the average acquisition time is the following [34]:

$$T_{\text{acq}} = (nN_c + M)T_h \quad (3.2)$$

In this case it is supposed that a single synchro event is required to start the tracking phase.

3.1.3 Two-Level Acquisition

The two-level acquisition scheme combines the properties of the serial search and those of the matched filter acquisition. Its schematic block diagram is depicted in Fig. 3.3.

The matched filter makes a correlation on M frequencies at the same time. When the threshold V_{T1} is exceeded, an active correlator starts to serially correlate over K hops. Any time one of those correlator outputs exceeds the threshold V_{T2} the system is considered synchronized. If none of these correlator outputs exceed the threshold, the procedure is repeated.

The uncertainty region contains again of nN_c cells. When the synchro cell is detected, the system needs to wait M hops in order to have all the outputs of the matched filter settled. At this point an extra K hops are required in order to allow the active correlator to perform a correlation serially on the K hops.

Therefore, the average acquisition time in this case is [34]:

$$T_{\text{acq}} = (nN_c + M + K) T_h \quad (3.3)$$

3.1.4 Acquisition Methods Comparison

The serial search algorithm is very cheap in terms of required computational power but it is quite slow. A parallel search strategy (matched filter based) is fast, but it is

Table 3.1 Maximum hopping time for an acquisition time of 5 ms

Acquisition algorithm	T_h [μ s]
Stepped serial search	4.3
Matched filter acquisition	65.8
Two-level acquisition	58.1

expensive both in chip area (and therefore, higher costs) and computational power. The two-level acquisition method combines the two previous algorithms so it has a medium hardware complexity. In a power constrained environment the best option would be to use a serial search algorithm. However, as it will be shown, this requires quite a fast hopping time, which is unrealistic for state-of-the-art synthesizers. Further in this book, a new synthesizer is discussed that is fast enough that it allows the use of the low-power serial search algorithm.

Let us first investigate the maximum hopping time for the three synchronization algorithms. Looking at Fig. 2.8 in a data-rate between 1 and 10 kbps the maximum wake-up time ranges between 10 ms and 100 ms. Lets suppose that 50% of this time is used for synchronization purposes and the remaining time is used to correctly set all the transceiver blocks to the correct DC point. This means that the maximum synchronization time ranges between 5 ms and 50 ms. When 5 ms synchronization time is considered, and the synchronization time is assumed to be dominated by the acquisition time, the required hopping (dwell) time for the three different algorithms is shown in Table 3.1 for $n = 1$, $M = 20$, $K = 10$ and for a number of channels N_c equal to 56. Next we will discuss how the parameter's values are chosen.

The value of n is chosen according to the fact that the acquisition needs to put the two PNCs in phase at least within $T_h/2$.

The value of M can be derived from the graphs in [34]. The value of M must be chosen in order to minimize the miss probability. A reasonable value of 1% has been used here. For $M = 20$ the required E_b/N_0 in a fading channel ranges between 7 and 8 dB for the stepped serial search algorithm and the matched filter acquisition algorithm. The two-level acquisition algorithm performs quite well in a benign environment [34] (when no fading is present) but unfortunately the miss probability in a hostile environment can be as high as 20% for $E_b/N_0 = 7$ dB. This forces the use of a much higher E_b/N_0 during synchronization, which can require to transmit the synchronization pattern at a power level considerably higher than required. For this reason the two-step acquisition is not very suitable for an indoor wireless link. Nevertheless, for completeness, the comparison between all these three algorithms is carried out in this section.

The K value used in the two-level acquisition algorithm is chosen to be 10 as in [34].

The number of channels N_c is derived from considerations about the operating bands. The book focuses on the ISM bands to reduce the costs for the end user. The ISM bands available are shown in Table 3.2. The 433.05–434.79 MHz band is a very narrow band and therefore, not very suitable for an FH system if a reasonable frequency diversity has to be achieved. The same consideration applies to the European 868–870 MHz band. The 5725–5875 MHz band is large but it is at quite high

Table 3.2 ISM bands in Europe and in USA

ISM bands Europe [MHz]	ISM bands USA [MHz]
433.05–434.79	–
868–870	902–928
2400–2483.5	2400–2483.5
5725–5875	5725–5875

frequency. The attenuation tends to increase with the frequency. For this reason this band has not been considered in this book. Concluding the two most suitable ISM bands are the 902–928 MHz available only in USA and the worldwide available 2400–2483.5 MHz band.

Looking at FCC rules in those bands it can be seen that while in the 2400–2483.5 MHz band 25 hopping channels are the required minimum number of channels for FHSS, this limit under certain conditions can shift up to 50 channels in the 902–928 MHz band. To overcome this restriction and to be able to use the wireless transceiver in both bands with little or none hardware modification, a value larger than 50 has been chosen. In the previous example 56 channels are used but later in the book a system with 64 channels is also proposed. The number of channels is not a big issue as soon as it is larger than 50 and not so large to exceed the available bandwidth or to considerably increase the average acquisition time.

Given those parameters, from Table 3.1 it can be seen that the stepped serial search method is an order of magnitude slower than the matched filter acquisition algorithm. This forces to increase the hopping rate by roughly a factor ten in order to have the same acquisition time. This drawback constitutes the main stop point in using a common available FH synthesizer together with a stepped serial search algorithm: it is proven later in this chapter that state-of-the-art FH synthesizers can become extremely power hungry under settling-time constraints. Further in this book it will be shown that a new, very fast synthesizer, can allow the use of the (very low-power) serial search synchronization algorithm.

3.2 State-of-the-Art FHSS Systems

The FHSS technique has found its largest application area in Bluetooth devices but it is not widely used in ultra-low power nodes. Moreover it is used mostly in the 902–928 MHz ISM band but not very often in the 2400–2483.5 MHz worldwide available ISM band. In this band the DSSS technique is generally used in standards like ZigBee. In Table 3.3 some FHSS transceivers with their TX and RX power consumptions are shown.

From Table 3.3 it is clear that the TX power consumption strongly depends on the radiated power. It is possible to divide the overall transmitter power consumption in a PA power and in a pre-PA power such that:

$$P_{TX} = P_{\text{pre-PA}} + P_{PA} \quad (3.4)$$

Table 3.3 Power consumption of some commercially available FHSS products

Product	Freq. band [MHz]	Cur. consum. TX/RX [mA]	Supply voltage [V]
TRC103	902–928	16/3.5 @ +1 dBm	2.1–3.6
CC1101	902–928	15.6/13.1 @ –6 dBm	1.8–3.6
CC2500	2400–2483.5	17.3/15 @ –6 dBm	1.8–3.6
MICRF500	902–928	14/12 @ –6 dBm	2.5–3.4
TinyOne Classic	902–928	80/35 @ 14 dBm	3–5
Mica2	902–928	27/10 @ 5 dBm	2.7–3.3
ADF7020	902–928	19.1/19 @ 0 dBm	2.3–3.6
Smartmesh-XT M1030	902–928	28/14 @ 4 dBm	2.7–3.6

Supposing that it is possible to redesign the PA for the same efficiency value but for a different output power, then the rms current scales approximately linearly with the output power neglecting losses in the matching network. Indeed, given a power efficiency η , the output power and the bias current are related by the following equation:

$$I_{\text{bias}} = \frac{P_{\text{out}}}{V_{\text{supply}}\eta} \quad (3.5)$$

where V_{supply} is the supply voltage and P_{out} is the output power. Losses are quite high if a high output power is required, but it can be neglected at low power levels. Therefore, it is possible to normalize the TX current consumption to the transmitted power to be able to better compare the results.

For a normalized –6 dBm output power, a 50 Ω radiating resistance and a 60% power efficiency the bias current has to be around 0.42 mA. The bias current required for output power other than –6 dB can be calculated applying (3.5). The difference between this current and the one calculated for a –6 dBm output power can be subtracted from the overall TX current consumption in order to normalize all the products to a common output power of –6 dBm. The results for these commercial products are plotted in Fig. 3.4.

Research work on FHSS systems looks quite limited and restricted mainly to Bluetooth applications. The results of these research projects are summarized in Table 3.4.²

As it can be seen both from Table 3.4 and Fig. 3.4, the peak power consumption of state-of-the-art FHSS transceivers is still well above the requirement for an ultra-low power wireless node. It is important to find what the bottleneck is in the power consumption for FHSS systems. For that reason (3.4) is used in order to distinguish between the power used in the PA and the pre-PA power used in general for the signal conditioning and the up-conversion. Equation (3.5) is used to calculate the

²As supply voltage the geometrical mean point $(V_{\text{max}} + V_{\text{min}})/2$ between the maximum (V_{max}) and the minimum (V_{min}) allowed supply voltage has been considered.

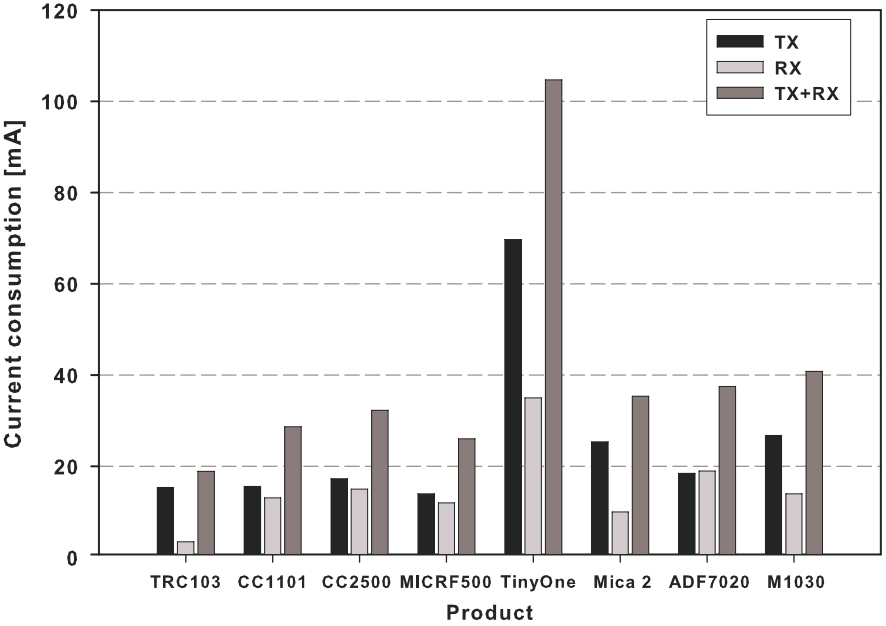


Fig. 3.4 Power consumption of some commercial FHSS products normalized to -6 dBm output power

Table 3.4 Power consumption of some FHSS ICs published in research papers

Ref.	Freq. band [MHz]	Cur. consum. TX/RX [mA]	Supply voltage [V]	Technology [μ m]
[35]	902–928	100/120 @ 13 dBm	3	1
[36]	2400–2483.5	26/34 @ 0 dBm	1.7–1.9	0.18
[37]	2400–2483.5	45/50 @ 0 dBm	2.5	0.25
[38]	2400–2483.5	40/60 @ 0 dBm	1.8	0.18
[39]	2400–2483.5	6.4/7 @ 0 dBm	2.5	0.25
[40]	2400–2483.5	25/33 @ 0 dBm	3	0.18
[41]	2400–2483.5	35/30 @ 0 dBm	2.5–3	0.18

PA current consumption for a power efficiency of 60% and a radiating resistance of $50\ \Omega$. This correction of PA power to retrieve the pre-PA power is done for several FHSS systems and the results are shown in Table 3.5. As shown, most of the power is wasted in the pre-PA part of the transceiver. This situation is emphasized especially at lower transmitted power (<0 dBm) where rarely the efficiency exceeds a few percents. One of the most power hungry blocks in the pre-PA part is the hopping frequency synthesizer. Therefore, the guideline toward power optimization of a wireless node starts from analyzing current architectures for frequency hopping synthesizers. Then, some bounding on the achievable minimum power consump-

Table 3.5 PA and pre-PA Power consumptions of FHSS ICs that are commercial or published in research papers

Ref.	Power consum. PA [mW]	Power consum. pre-PA [mW]	Efficiency [%]
TRC103	2.11	43.6	4.6
CC1101	0.43	41.6	1
CC2500	0.43	46.2	0.9
MICRF500	0.41	41	1
TinyOne Classic	42	278	13
Mica2	5.28	74.3	6.6
ADF7020	1.65	54.6	2.9
Smartmesh-XT M1030	4.19	84.1	4.7
[35]	33.3	266.7	11.1
[36]	1.7	45.1	3.6
[37]	1.7	110.8	1.5
[38]	1.7	70.3	2.4
[39]	1.7	14.3	10.6
[40]	1.7	73.3	2.3
[41]	1.7	94.6	1.8

tion under certain specifications is found for commonly used FH architectures. This study puts the basis toward a new FH synthesizer architecture able to reduce the synthesizer power consumption well below those limits. In this way an increase in the transmitter efficiency can be reached and this will also be beneficial for the receiver part decreasing its peak power consumption as well because the RX needs the same synthesizer.

3.3 FH Synthesizer Architectures

The core block of an FHSS system is the hopping synthesizer. While FHSS methods seem attractive for robust wireless applications, challenges in their implementation arise from requirements for operation at ultra-low power over a broad operating band. These requirements for agile and accurate frequency hopping require fast settling behavior for frequency selection as well as a high degree of accuracy in the frequency synthesis.

Hopping frequency synthesizers are conventionally based on a PLL with a digitally controlled variable divider. Using a fractional- N divider allows to synthesize frequency bins closely spaced to each other. Another way to synthesize accurately the frequency bins, is to use a mixed-signal approach, in which the accuracy and the fast settling are realized through the use of a DDFS and a DAC is used to translate the discrete time periodic wave coming out from the DDFS in a continuous

waveform with specified spectral characteristics. This arrangement requires a fixed-frequency LO, which, in general, can be easily integrated.

In the following sections the PLL-based hopping frequency synthesizer and the DDFS–DAC frequency synthesizer will be discussed more in detail.

3.4 Specifications for Ultra-low-power Frequency-Hopping Synthesizers

FCC rules demand for at least 25 hopping channels (20-dB channel bandwidth greater than 250 kHz) in the 902–928 MHz band and 15 channels in the 2.4 GHz ISM band. To be able to fulfill the FCC rules in both the ISM bands a 52 channel system has been considered here.

The spacing between two adjacent channels has been set to 0.5 MHz. The choice of such a large spacing between the adjacent channels even for a data-rate, which is supposed to be below 10 kbps is dictated by the necessity to have a certain frequency diversity in the indoor environment. Using the statistical approach, the coherence bandwidth is defined as the bandwidth over which the fading statistics are correlated to better than 90 percent. Clearly, if two frequencies do not fall within the coherence bandwidth, they will fade independently. The delay spread in a channel can be used as a figure of merit to calculate the coherence bandwidth. In typical indoor environments, the delay spread is around 50 ns at 2.4 GHz [42]. The coherence bandwidth for a 90% correlation is given by the following relation

$$BW_c \simeq \frac{1}{50\sigma_\tau} \quad (3.6)$$

where BW_c is the coherence bandwidth and σ_τ is the delay spread. For example, in the 2.4 GHz environment the coherence bandwidth is around 400 kHz. Hence, the choice of a 0.5 MHz channel spacing makes the channel to fade more independently.

The level of unwanted spurs is an important specification of the synthesizer. During transmission of a certain frequency it is necessary to keep the unwanted spurs below a certain level. They will generally pollute other channels causing co-channel interference to other nodes communicating on those channels.

In [43] it is shown that 30 dB rejection of the image signal is largely sufficient to avoid the deterioration of the BER of the transmission. Considering an extra margin of about 10 dB to take partially into account the fading condition on the wanted signal, a 40 dB required rejection has been considered. Moreover, given the fact that for the up-conversion a SSB scheme is used, due to process spread, the image will not be rejected perfectly. A common figure of merit for harmonic rejection is about 40 dB and in general higher degrees of rejection require complex calibration techniques. Given an extra implementation margin of 6 dB, the overall required SFDR for the synthesizer can be set equal to 46 dB.

3.4.1 PLL Based

The classical PLL system is generally known as the integer- N PLL. Its schematic block diagram is depicted in Fig. 3.5. The relation between the reference frequency and the synthesized output frequency is the following:

$$f_{\text{VCO}} = N \times f_r \quad (3.7)$$

The accuracy of the synthesized frequency is equal to the accuracy of the reference frequency. Therefore, if a 1 ppm accuracy reference signal is used to synthesize a 2.4 GHz output frequency, then the output frequency accuracy will be ± 2.4 kHz. At 915 MHz this accuracy will be ± 915 Hz. Supposing that in the 2.4–2.4835 GHz band the separation between adjacent channels is in the order of 500 kHz, then the required N is about 4800. If in the 902–928 MHz this separation is 150 kHz then $N = 6100$.

Therefore, in general N is a quite large number, which has severe drawbacks on the phase noise at the PLL output. Indeed, a frequency multiplication by a factor N raises the signal's phase noise by $20 \log_{10}(N)$ dB. Therefore, the noise floor of the phase detector raises up by 73.6 and 75.7 dB in the two previous cases. This means that the N value cannot be very large, implying a strict trade-off between channel spacing and noise performances.

The phase detector generates a high level of noise, which needs to be suppressed by a loop filter. The noise modulates the VCO input and generates tones at $\pm f_r$ and harmonics around the output frequency. If f_r is small (small inter-channel spacing) the loop-filter must have a narrow bandwidth. This translates in long switching times between channels. When the system is in the lock-in range and in the case of a second order PLL, this time can be approximated by the following expression:

$$T_{\text{settling}} \propto \frac{1}{\delta \omega_n} \quad (3.8)$$

where δ is called the damping factor and ω_n is the system natural frequency. Concluding, there is a trade-off between switching speed and spur suppression which can increase the power consumption of the system for very high switching speeds.

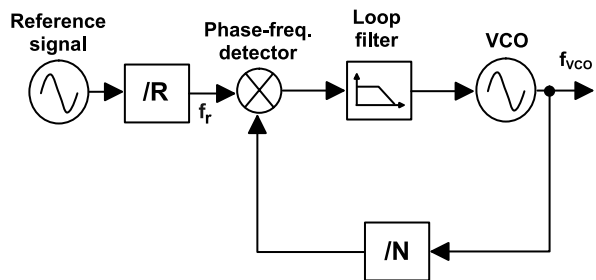
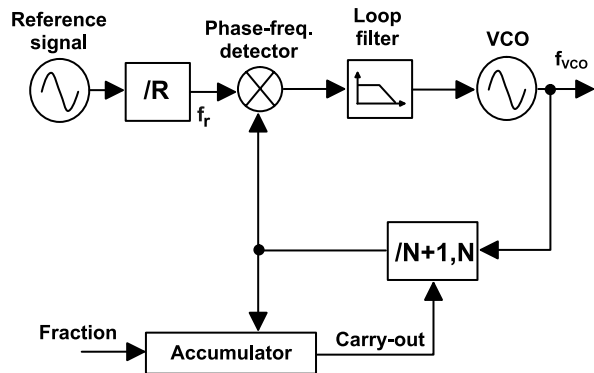


Fig. 3.5 Integer- N PLL block diagram

Fig. 3.6 Fractional- N PLL block diagram



To get around this problem a fractional- N type of architecture can be used. The schematic block diagram is depicted in Fig. 3.6. The output frequency in a fractional- N architecture is given by the following relation:

$$f_{\text{VCO}} = \left(N + \frac{K}{F} \right) \times f_r \quad (3.9)$$

where F is the fractional resolution with respect to the reference signal. For example if the fractional resolution is 8 then the reference frequency can be increased 8 times and the N value is decreased by the same amount. For the two previous cases, the reference signal can now be 4 MHz in the 2.4–2.4835 GHz band and 1.2 MHz in the 902–928 MHz band. The N value is decreased to 763 and 600 respectively. Theoretically, the phase noise is reduced by 18 dB but also the speed is affected because the loop bandwidth can be enlarged.

The principle behind the generation of a fractional division is averaging. Indeed, equation (3.9) should be seen as an average value of the VCO output frequency. This is obtained by changing the feedback divider dynamically between the values N and $N + 1$. If out of F cycles K times the division value is $N + 1$ and $F - K$ times is N , then the average output frequency equals the value given by (3.9).

Unfortunately, for narrow spaced channels, the traditional loop filter is not able to filter out all the fractional spurs generated by the instantaneous change of the divider ratio during lock condition. These spurs in the absence of a filter can be 7 dB higher than the carrier frequency. This can completely erase all the benefits of this kind of architecture. Therefore, a spur compensation circuitry is required. This avoids to lower the loop bandwidth and therefore, avoids the degradation in the settling time at the costs of a higher power consumption. The compensation can reduce the spur levels by 40 dB and together with the loop filter a 60 dB attenuation in the spur levels can be achieved.

One last drawback in this architecture is an increment of the $1/f$ noise coming from the PFD. This can be a big limitation with technology scaling. Some recent works on PLL based synthesizers are listed in Table 3.6. It appears clear that state-of-the-art PLL based synthesizers are still too power hungry for an ultra-low power node. This will be confirmed in the following section by a model, which connects together the PLL settling time and its predicted power consumption.

Table 3.6 Power consumption of some published PLL based frequency synthesizers

Ref.	Curr. consump. [mA]	Settl. time [μ s]	V_{supply}	Technology [μ m]
[44]	22	–	2.5	0.25
[45]	17	70	2.5	0.18
[46]	8.94	<25	1	0.18
[47]	6.15	200	0.8	0.18
[48]	14	650	2.7	0.35
[49]	9.3	20	1.8	0.18
[50]	12.4	<5	3.3	0.5
[51]	10.8	<67	1.8	0.18

3.4.2 DDFS Based

A second approach to the frequency synthesis is based on the DDFS topology combined with a DAC. Several architectures have been developed for the direct-digital frequency synthesis. They can be divided into two large groups:

- ROM based
- ROM-less

The ROM-based DDFSs use a ROM to convert the phase information into an amplitude information. This ROM can be sometimes very large and, therefore, very power hungry. For this reason several improvements to this architecture have been developed in order to reduce the ROM size. Some of these modified architectures are the following [52]:

- Modified Sunderland architecture
- Modified Nicholas architecture
- Taylor series expansion
- Cordic algorithm
- Quarter wave symmetry

All of these improved architectures can shrink the ROM considerably saving, therefore, a large amount of power especially when the ROM is large. However, as it will be proven in Sect. 3.6.5, in DDFSs designed for ultra-low power transmitters the ROM contributes only less than 15% to the overall power consumption. For this reason these improved architectures will not be crucial here, and thus will not be considered further in this paper.

ROM-less architectures manage to remove completely the ROM. The phase to amplitude conversion is performed using interpolation algorithms or non-linear DACs [52]. These ROM-less architectures achieve good performances when the synthesized frequency is very low compared to the clock frequency. In the proposed case the maximum synthesized frequency is around 13 MHz. The clock frequency must be kept low in order to have a low power clock generation. As it will be clear in Sect. 3.6.5, the clock frequency is around 90 MHz and therefore comparable with

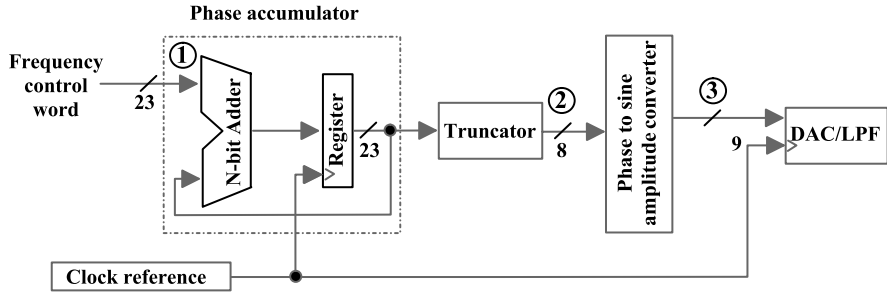


Fig. 3.7 Basic direct digital frequency synthesizer, with bits required in the various stages

the maximum synthesized frequency. Furthermore, considering that the ROM is not the dominant contributor to the overall power consumption (see Sect. 3.6.5), these ROM-less architectures will not be discussed further in this article and a simple ROM based DDFS is considered throughout the following sections.

A schematic block diagram of a ROM-based DDFS is shown in Fig. 3.7. The system has two inputs:

- Clock reference f_{clk}
- Frequency Control Word (FCW)

The phase accumulator integrates the value from the FCW on every clock cycle, producing a ramp with a slope proportional to the FCW and therefore, to the synthesized frequency. If N is the width (in bits) of the accumulator, then the ramp is given by:

$$f_{\text{out}} = f_{\text{clk}} \times \frac{FCW}{2^N} \quad (3.10)$$

At a glance it is possible to recognize that increasing the clock frequency will reduce the filter complexity (and therefore, most probably the filter power consumption) but will increase the power consumption of the digital part and probably of the DAC.

The phase accumulator output is converted to an approximated sine amplitude by the phase-to-sine amplitude converter. In the simplest case, this converter is a ROM. A DAC and a Low-Pass Filter (LPF) are used to convert the sinusoid samples into an analog waveform.

Therefore, in any DDFS based synthesizer three parameters need to be sized based on the required system specifications:

- Accumulator's number of bits (point 1 in Fig. 3.7)
- ROM word-length (point 2 in Fig. 3.7)
- DAC number of bits (point 3 in Fig. 3.7)

Let's first look at point 1 in Fig. 3.7. The number of bits in the phase accumulator depends on the frequency accuracy required. The frequency resolution is given by the following relation

$$f_{\text{res}} = \frac{f_{\text{clk}}}{2^N} \quad (3.11)$$

Therefore, the minimum number of bits in the phase accumulator is found by inverting the previous relation and approximating the result to the greater integer.

Let's now look at point 2 in Fig. 3.7. The address word is generally truncated in order to reduce the ROM size and therefore the power consumption. This means that the number of bits in the address word sets the system phase resolution. The phase quantization produces spurs at the DDFS output. The level of those spurs depends on how much the address word bits are truncated.

Finally we will have a look at point 3 in Fig. 3.7. Each address in the ROM contains (in an equivalent digital format) the amplitude value of the sinewave for the relevant phase (e.g. if the address word is 8-bit long then in the ROM there are 256 possible amplitude values, one for each of the possible phases).³ The number of bits in each memory cell used to encode the amplitude information for a given sine phase sets the amplitude quantization. It equals the number of bits in the DAC. The finite number of bits in the DAC and its Integral Non-Linearity (INL) are sources of distortion at the DDFS output as well.

There is a main difference between a PLL based hopping synthesizer and a DDFS based hopping synthesizer. The PLL generally can synthesize the required channels already in the wanted band, while the DDFS synthesizes the channels at baseband. The DDFS is a sampled system and therefore, it is limited by the Nyquist theorem, which states that the sampling frequency must be at least twice the maximum signal frequency. This means that if the 902–928 MHz band needs to be synthesized from the DDFS, the required sampling frequency should be 1.856 GHz while the 2.4–2.4835 GHz band requires a sampling frequency of 4.967 GHz.

Even if this sampling frequency can be achieved in the most advanced technologies it would require a very large power consumption. Therefore, the DDFS uses an analog up-conversion stage, that can be avoided if a PLL architecture is used. The full synthesizer including the up-conversion stage looks like in Fig. 3.8. The architecture shown in Fig. 3.8 uses also a SSB up-conversion scheme to further reduce by a factor two the maximum frequency the DDFS needs to synthesize to cover the full-bandwidth [53]. For example, to cover the 26 MHz of the 902–928 MHz bandwidth, the DDFS is required to synthesize up to 13 MHz. In the 2.4–2.4835 GHz band this value goes up to 41.75 MHz.

Given the fact that the DDFS is a sampled system, a common rule of thumb is that the reference frequency is restricted by the following relation:

$$f_{\text{clk}} \geq \frac{f_{\text{max}}}{0.4} \quad (3.12)$$

where f_{max} is the maximum signal component. For example, in the 902–928 MHz band, the maximum signal component is at 13 MHz frequency. Therefore, the clock frequency has to be bigger than 32.5 MHz.

In Table 3.7, performances of some recently published DDFS based synthesizers are summarized. Each work has been compared in terms of power consumption per

³Here the possibility to reduce the ROM size by using the quarter wave symmetry of the sine function has not been considered for clarity.

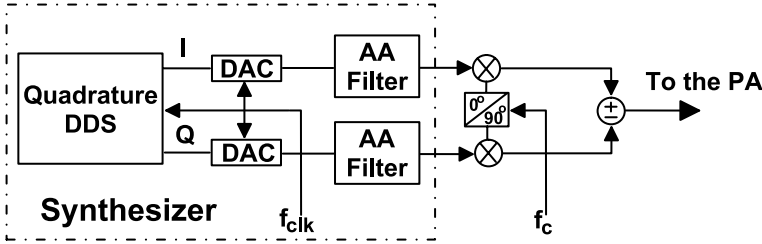


Fig. 3.8 Complete DDFS based hopping synthesizer including the up-conversion stage

Table 3.7 Power consumption of some published DDFS based frequency synthesizers

Ref.	P/MHz [$\mu\text{W}/\text{MHz}$]	P [mW] ^a	SFDR	DAC	Technology [μm]
[54]	400	13/41.6	50 @ 10 MHz	yes	0.5
[55]	317	10.3/32.97	58 @ 9.375 MHz	no	0.8
[56]	80	2.6/8.32	42.1 @ 1.56 MHz	yes	0.5
[57]	30	0.98/3.12	60	no	0.18
[58]	217	7.05/22.57	55 @ 18.75 MHz	yes	0.35
[59]	400	13/41.6	110	yes	0.25
[60]	72	2.34/7.49	80 @ 54 MHz	no	0.25
[61]	667	21.67/69.37	60 @ 8 MHz	yes	0.35

^aPower consumption is recalculated for the 902–928 MHz and for the 2.4–2.4835 GHz bands respectively

MHz. This power (see first column in Table 3.7) is derived by simply dividing the overall DDFS power consumption by its operating frequency. This value is used to scale the power consumption to a different operating frequency (which is set by the clock frequency) more appropriate for the cases considered in the book. In this book two possible frequency bands have been considered: the 902–928 MHz band and the 2.4–2.4835 GHz band. For these two bands the minimum required clock frequency is 32.5 MHz and 104 MHz.⁴ The scaled power consumption for the two frequency bands under consideration is shown in the second column of Table 3.7.

From Table 3.7 it is possible to see that the power consumption of this architecture is still very high. The work of [56] is able to reduce the power consumption at lower levels but this comes at the expense of a poor SFDR even at low frequencies. The work in [57], while achieving a very low power consumption, does not include the DAC. This, as it will be proven in the following sections, is the most power hungry block in low-end DDFSs. The remaining works listed in Table 3.7 deal with high-end DDFSs characterized by a large required SFDR. For that rea-

⁴The clock frequency is set by the frequency range to cover and by the Nyquist theorem as previously stated in this chapter.

son their power consumption is very high and therefore, not compatible with an ultra-low power implementation.

The synthesizer architecture is very important for the power optimization of the overall system. Therefore, it is mandatory to derive the minimum required specifications for the synthesizer and to create a power model that, starting from those specifications, can set a lower bound to the synthesizer power consumption. This operation is performed in the next two sections for both types of synthesizers.

3.5 PLL Power Estimation Model

A PLL power model has been developed at the Pennsylvania State University [16]. In this book the same model is applied to relate the PLL power consumption with its settling time. In this way the trade-offs between power consumption and synchronization time can be highlighted. The numerical analysis is carried out for both the 902–928 MHz band and the 2.4–2.4835 GHz band. The analysis is restricted to a charge-pump based integer-N type PLL. The combination of a PFD and a charge pump has some advantages over other structures. Mainly, the capture range is only limited by the VCO tuning range and in first approximation the phase error is zero under lock conditions.

The PLL, as mentioned in Sect. 3.4.1, is composed by different sub-blocks arranged in a closed loop system:

- VCO
- Frequency divider
- PFD
- Loop filter

The PLL behavior strongly depends on the frequency step that is applied to it or conversely how much the divider value is changed. If the frequency step takes the system beyond the lock-in range, then a pull-in process starts which is generally referred as a capture process. Therefore, the total system settling time is given by the sum of these two components:

$$t_{\text{settling}} = t_{\text{lock-in}} + t_{\text{capture}} \quad (3.13)$$

3.5.1 VCO

The VCO is characterized by its gain and by the required phase-noise at a certain distance from the carrier. The required gain depends on the frequency range needed to be covered. Generally a $\pm 10\%$ of the center frequency is required as a tuning range in order to compensate for the process spread. The tunability is obtained by changing for example the capacitance in an LC-tank based oscillator or the single-stage

time constant in a ring oscillator. The required VCO gain can be easily calculated by using the following equation:

$$K_{\text{VCO}} = \frac{2 \times \Delta_{\text{spread}} f_{\text{center}} + BW}{\Delta V} \quad (3.14)$$

where Δ_{spread} is the spread on the synthesized center frequency (i.e. 10% = 0.1), f_{center} is the band center frequency (i.e. 915 MHz or 2.4415 GHz), BW is the bandwidth within a certain ISM band (i.e. 26 MHz in the 902–928 MHz band) and ΔV is the maximum VCO control voltage variation. Supposing a 1 V supply voltage the required gain for the 902–928 MHz and the 2.4–2.4835 GHz bands is respectively 261 MHz/V and 725 MHz/V.⁵

The work in [16] considers as oscillator a ring oscillator. In general, LC-tank based oscillators are less power hungry than ring types given a certain phase-noise specification. Though this is true, it should be noticed that phase-noise specifications cannot be very tight in order to not harm the power consumption suggesting that the power consumption will be determined by the open-loop gain rather than by the phase-noise. Furthermore, Q factors of integrated coils, are not so high especially in CMOS technology. Therefore, though an LC-tank based oscillator can have a lower power consumption than a ring based oscillator, the second one has been considered. This does not mitigate the main idea that this section wants to communicate, which is the existing connection between settling time and power consumption in any PLL based synthesizer.

The model adopted for the effective capacitance in [16] for a differential VCO design is the following:

$$C_{\text{VCO}} = 2n^2 k (W_n (C_{\text{gate}} + C_{\text{drain},n}) + 2W_p C_{\text{drain},n}) \quad (3.15)$$

where k ($V_{\text{sw}} = k V_{\text{supply}}$) defines the voltage swing of each cell, W_n and W_p are the normalized width of the NMOS and the PMOS, n is the number of stages in the VCO, C_{gate} is the gate capacitance per unit area and $C_{\text{drain},n}$ and $C_{\text{drain},p}$ are the drain diffusion capacitance per unit area of an NMOS and a PMOS respectively. This model supposes that the single delay cell in the ring oscillator has the topology depicted in Fig. 3.9. The dependence of the VCO effective capacitance with the squared value of the number of stages in the VCO comes from the fact that, in a differential cell (see Fig. 3.9), there is always one branch, which conducts current. Therefore, the average current is proportional to the number of VCO stages, while the VCO frequency is inverse proportional to the number of stages.

The VCO effective capacitance is the only required value to estimate the VCO power consumption taking into account that it acts as a digital block and therefore, its power can be expressed as follows:

$$P_{\text{VCO}} = C_{\text{VCO}} f_{\text{capture,average}} V_{\text{supply}}^2 \quad (3.16)$$

⁵The VCO control voltage range has been restricted to 0.8 V. This is quite common because for example varactors cannot be controlled over the entire supply voltage range but only on a portion of it. The same restrictions apply to a ring oscillator when the bias current has to be changed in order to change the per-stage time constant.

Fig. 3.9 Schematic of a differential cell used in differential ring oscillators

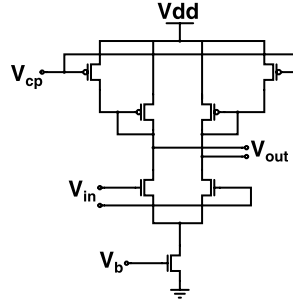
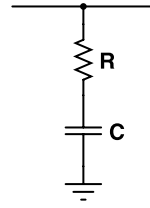


Fig. 3.10 Simple first-order PLL loop filter



where $f_{\text{capture,average}}$ is the mean value between the start and the stop frequencies during the capture process and V_{supply} is the supply voltage.

3.5.2 Loop Filter

The loop filter affects the dynamic of the whole synthesizer. This analysis is restricted to a second-order continuous time PLL and therefore, as a loop filter a simple RC filter is used where the R and the C are connected in series as depicted in Fig. 3.10. The filter is passive and therefore, it does not contribute to the system power consumption but it affects only its dynamic behavior.

3.5.3 Charge Pump

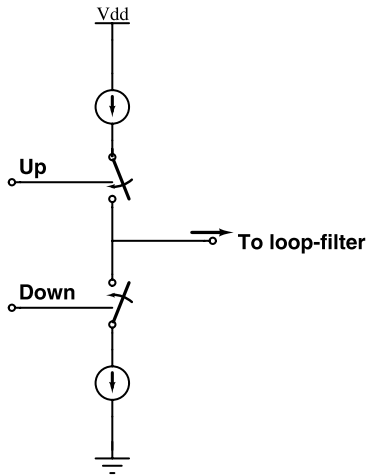
The charge pump is characterized by its current I_{CH} . This current is delivered to the loop filter in burst and its direction (sourced current or sunk current) is determined by the PFD output. A simplified block diagram of a charge pump is depicted in Fig. 3.11.

The charge pump power consumption can be estimated using the following relation:

$$P_{\text{CH}} = \left(\frac{\Delta V^2 \cdot C}{t_{\text{capture}}} \right) + \left(\frac{\Delta V \cdot C}{t_{\text{capture}}} \right)^2 R = \left(\frac{\Delta V \cdot I_{\text{CH}}}{2} \right) + \left(\frac{I_{\text{CH}}^2 \cdot R}{4} \right) \quad (3.17)$$

where ΔV is the total voltage variation on the capacitor C during the capture process and t_{capture} is the capture time.

Fig. 3.11 Simplified block diagram of a charge pump



3.5.4 PFD and Frequency Divider

As in the case of the VCO, also the PFD and the frequency dividers can be characterized by their effective capacitance. For the PFD:

$$C_{\text{PFD}} = N_{\text{PFD}} G_{\text{PFD}} C_{\text{tech}} \quad (3.18)$$

where N_{PFD} is the number of minimum size transistors in the PFD, G_{PFD} is the average gate activity and $C_{\text{tech}} = C_{\text{gate}}(W_n + W_p) + W_n C_{\text{drain},n} + W_p C_{\text{drain},p}$.

The same model can be applied to the frequency dividers where G_{PFD} and N_{PFD} are substituted by G_{div} and N_{div} respectively.

3.5.5 Complete PLL Power Model

In this section the relation between PLL power consumption and settling time is derived starting from the power consumption of the single PLL building blocks disclosed in the previous sections. The PLL lock-in time can be approximated by the following relation:

$$t_{\text{lock-in}} = \frac{1}{\delta \omega_n} \quad (3.19)$$

where the damping factor δ equals

$$\delta = \frac{R}{2} \sqrt{\frac{1}{N} K_{\text{VCO}} I_{\text{CH}} C} \quad (3.20)$$

and the natural frequency ω_n equals:

$$\omega_n = \sqrt{\frac{K_{\text{VCO}} I_{\text{CH}}}{C N}} \quad (3.21)$$

Table 3.8 N_{mean} values for the two ISM bands discussed in the book for $f_{\text{ref}} = 0.5$ MHz

Band [MHz]	N_{max}	N_{min}	N_{mean}
902–928	1856	1804	1830
2400–2483.5	4800	4967	4883

Another important parameter is the level of spurs (β) the system can tolerate. Given the filter shown in Fig. 3.10 and considering dominant the leakage dominated spurs⁶ the required capacitance can be easily calculated using the following relation:

$$\beta_{\text{leak}} = \frac{K_{\text{VCO}} I_{\text{leak}}}{f_{\text{ref}}^2 \cdot C} \quad (3.22)$$

where I_{leak} is the charge pump leakage current and f_{ref} is the reference clock signal. The N value can be chosen as the geometric mean between the minimum N_{min} and maximum N_{max} required N :

$$N_{\text{mean}} = \sqrt{N_{\text{max}} \cdot N_{\text{min}}} \quad (3.23)$$

These values can be calculated given the required bandwidth to be covered and the reference clock signal. For the reference clock signal as an example a value of 0.5 MHz can be chosen.⁷ The N_{mean} for the two ISM bands are summarized in Table 3.8. The capture time can be calculated by the following equation:

$$t_{\text{capture}} = |\Delta f| \frac{2C}{K_{\text{VCO}} I_{\text{CH}}} \quad (3.24)$$

where Δf is the frequency jump after a change in the N value. From (3.13), (3.19), (3.20) and (3.24) the settling time can be expressed as follows when δ is assumed to be equal to $\frac{1}{\sqrt{2}}$:

$$t_{\text{settling}} = \sqrt{\frac{2N_{\text{mean}}}{K_{\text{VCO}}}} \times \sqrt{\frac{C}{I_{\text{CH}}}} + \frac{2|\Delta f|}{K_{\text{VCO}}} \times \frac{C}{I_{\text{CH}}} \quad (3.25)$$

This equation can be solved with respect to C/I_{CH} . The value of C can be derived from (3.22) given the required SFDR of the synthesizer. From C and the solution of (3.25) the charge-pump current can be calculated. Finally the R value can be easily calculated from (3.20).

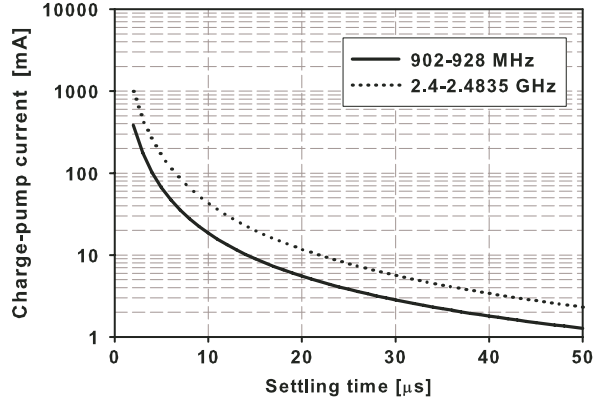
As an example the charge-pump current in the 902–928 MHz band and in the 2.4–2.483.5 GHz band is calculated. The required capacitance values for a 46 dB SFDR and a 5 nA leakage current are 87.2 nF and 247 nF respectively.⁸ In the two

⁶The mismatch dominated spurs can be made negligible by careful layout and design and by using compensating circuitry.

⁷This value while looking arbitrary has been chosen in order to have a fair comparison with the architectures proposed in the next chapters.

⁸Those capacitance values are too large to be integrated. Therefore, an active filter will be probably required. This increases the overall power consumption above the value calculated using a simple passive filter. Another possibility is to have external capacitances, which will increase cost and form factor of the wireless node.

Fig. 3.12 Charge pump current consumption versus the required settling time



cases it has been supposed that the bandwidth to be covered is equal to 26 MHz, which corresponds to 52 channels if a 0.5 MHz inter-channel spacing is considered. This fulfills the FCC rules in both the 902–928 MHz band and the 2.4–2.4835 GHz band. The required current in the charge pump as a function of the settling time is shown in Fig. 3.12.

It should be noted that, at very small settling time, the reference frequency f_{ref} needs to be changed in order to assure the stability of the loop. It is reasonable to say that the reference frequency has to be larger than at least 6 times the loop bandwidth to assure the loop stability. Now the loop bandwidth can be approximated by the following equation:

$$P_{LL\text{BW}} = \frac{K_{VCO} R I_{CH}}{2\pi N} \quad (3.26)$$

For a ratio between the loop bandwidth and the reference frequency larger than 6 the settling time has to be larger than 4.2 μs in the 902–928 MHz band and larger than 4 μs for the 2400–2483.5 MHz band. When the required settling time goes below these threshold values f_{ref} must be increased. This means that a fractional- N synthesizer is required to allow a frequency step of 0.5 MHz. Concluding, the overall power consumption will increase due to the fact that more hardware is required in a fractional- N loop compared to the integer- N loop here disclosed.

The PLL power consumption can now easily be derived with the following equation:

$$\begin{aligned} P_{PLL} &= P_{PFD} + P_{VCO} + P_{div} \\ &= \frac{1}{t_{\text{settling}}} (P_{PLL, \text{lock-in}} t_{\text{lock-in}} + P_{PLL, \text{capture}} t_{\text{capture}}) \end{aligned} \quad (3.27)$$

The PLL power consumption during the lock-in process can be evaluated using the following relation:

$$\begin{aligned} P_{PLL, \text{lock-in}} &= (N_{PFD} G_{PFD} + N_{\text{mean}} \times (2n^2 k \cdot 8 + N_{div} G_{div})) \\ &\quad \times C_{\text{tech}} V_{\text{supply}}^2 f_{\text{ref}} \end{aligned} \quad (3.28)$$

In the work in [16] an extra term is considered to take into account also the power consumption of the bias circuitry. This term has been neglected here for the purpose of what this section wants to prove.

During the capture process the power consumption can be approximated by the following relation:

$$P_{\text{PLL,capture}} = (C_{\text{PFD}} + C_{\text{VCO}} + C_{\text{div}}) V_{\text{supply}}^2 f_{\text{capture,ave}} + P_{\text{CH}} \quad (3.29)$$

where C_{div} is the effective frequency divider capacitance and C_{div} is the effective phase-frequency detector capacitance. The overall power consumption versus the settling time for the two ISM bands considered in this book is shown in Fig. 3.13. The higher power consumption in the 2.4–2.4835 GHz band comes from the higher required VCO gain to cover the process spread variation in the VCO center frequency. Indeed the spurious level increases by 6 dB at every doubling of the VCO gain. Therefore, if the spurious has to be below 46 dB, the integrating capacitance must increase proportionally. A larger capacitance requires a larger charge-pump current in order to guarantee a certain settling time (see (3.21) and (3.24)).

Therefore, a way to reduce the power consumption consists in reducing the VCO gain and the required division ratio N . A possible way to reduce these two values consists in synthesizing the channel bins close to baseband and then to up-convert those frequencies to the wanted frequency band by using a mixer. The concept is illustrated in Fig. 3.14. The center frequency cannot be chosen too low otherwise the required VCO tuning range will be too big. Supposing that a $\pm 15\%$ tuning range can be achieved and a 26 MHz bandwidth is required, then a center frequency of 87 MHz is required. Furthermore, any mixing process generates an upper and a lower sideband. This means that one sideband must be filtered out by an high-frequency BPF (as shown in Fig. 3.14) or by using a quadrature up-conversion scheme with sideband rejection. Supposing that the lower sideband is removed, the free running frequencies for the VCO in the two ISM bands is 828 MHz in the 902–928 MHz band and 2354.75 MHz in the 2.4–2.4835 GHz band.

A PLL with a center frequency of 87 MHz and 26 MHz tuning range, when a 10% process spread is taken into account on the center frequency requires at

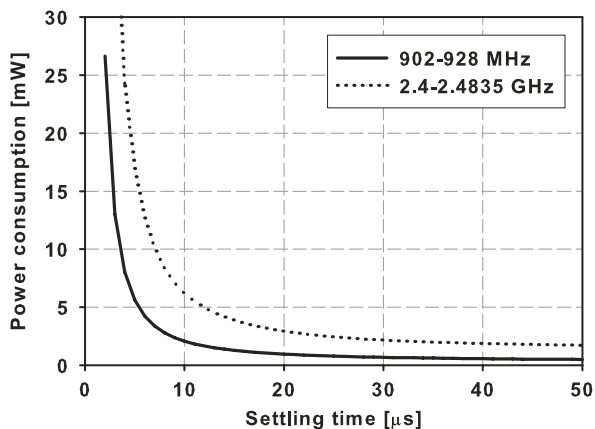


Fig. 3.13 PLL power consumption estimation versus the required settling time

Fig. 3.14 Modified transmitter architecture using a baseband PLL followed by an up-conversion mixer

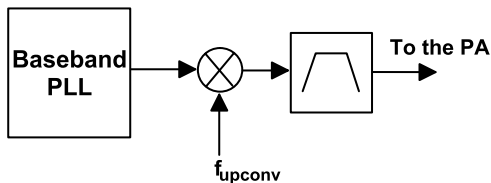
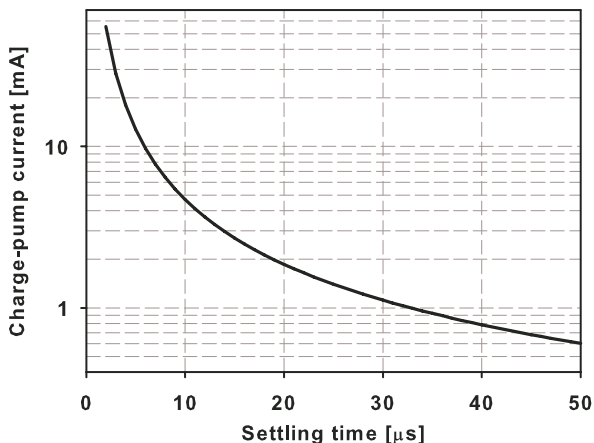
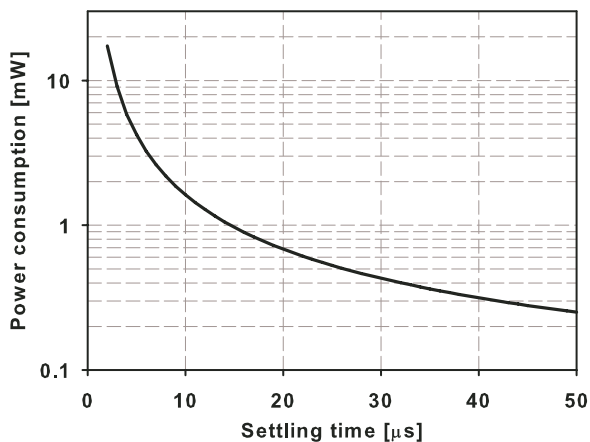


Fig. 3.15 Power performances of a baseband PLL for frequency hopping



(a) Charge-pump current consumption versus settling time



(b) Power consumption of the baseband PLL versus settling time

least a VCO gain of 54.25 MHz/V given a $\Delta V = 0.8$ V according to (3.14). Supposing a 0.5 MHz reference signal, $N_{\min} = 160$, $N_{\max} = 212$ and $N_{\text{mean}} = 184$. With these values the power consumption and the required charge-pump current can be estimated as a function of the required settling time. The result is shown in Fig. 3.15. Though the required power consumption is reduced compared to the previously analyzed high-frequency PLL, it should be noticed that an up-conversion

stage is required in this configuration. Furthermore, either a high frequency BPF or a quadrature up-conversion is required to remove the unwanted frequency components coming from the mixing process. This will again translate in an increased power consumption.

From this section it can be concluded that though several configurations or topologies can be chosen in order to increase the locking speed without affecting too much the power consumption, there will always be a connection between these two quantities. This connection comes from the fact that any PLL system is a closed loop system that uses a-posteriori information about the signal. This condition creates a link between speed and power consumption.

As stated in Chap. 2, the capability of an FHSS system to hop very fast brings several benefits including the possibility to lower the transmitted power trading it off with the frequency diversity. Now this should not come at expenses of an increased power consumption in the synthesizer. Therefore, any kind of PLL based system poses a hard limit to the maximum hopping speed at which the system can hop. This limits does not only affect the capability to reduce the transmitted power but also the speed in the synchronization process. The end result is an increase in the wake-up time of the node and therefore, an increase in the average power consumption.

3.6 DDFS Power Estimation Model

A DDFS synthesizer is a feed forward system, which bases its operation on an a-priori information on the wanted output. This eliminates the connection between the settling time and the power consumption, which was the main bottleneck in a PLL based hopping synthesizer architecture. This, as already mentioned in Sect. 3.4.2 comes at the expense of a high power consumption. Furthermore, most of the works published in literature and available on the market tend to address synthesizers with a very high SFDR. Such a high SFDR is not required in the applications foreseen for ultra-low power nodes given the fact that duty-cycling and frequency diversity make the probability of collision very low. Therefore, the power model, which is addressed in this section, targets mainly DDFS synthesizers with a low required SFDR. As it will be clear later in this section, with low SFDR, we intend an SFDR lower than 64 dB.

3.6.1 DDFS Specifications for Frequency-Hopping Synthesizers

The number of channels is set to 52 and therefore fulfills the FCC rules in both the ISM bands under discussion. Therefore, the maximum signal frequency at the DDFS output equals 13 MHz.

The resolution is set in order to not limit the choice of the FSK demodulator topology. The most stringent requirements in terms of frequency offset come

Table 3.9 DDFS requirements

	Spec.	Unit
Maximum frequency	13	MHz
Minimum frequency	0.5	MHz
Resolution	5	Hz
SFDR	46	dB

from the correlation based demodulator [62]. For example for a data-rate of 3 kbps⁹ a 10 ppm accuracy is a suitable choice (worst case offset is 260 Hz). This translates in a 5 Hz resolution when a 0.5 MHz frequency is synthesized.

Table 3.9 summarizes the DDFS specifications for an ultra-low-power DDFS based synthesizer for low data-rate applications. Considering (3.11) and remembering that the sampling function in the DAC creates an image of the wanted frequency at the output, the first trade off can be easily foreseen. Increasing the clock frequency will require a larger accumulator working at a higher frequency (more power) but it will relax the specification on the AA-filter necessary to attenuate the image frequency. Anyhow the minimum number of bits can be derived for the minimum clock frequency requirement (32.5 MHz). This brings to an accumulator with at least 23 bits. The final clock frequency may be different. Indeed while the digital back-end and the DAC would require a low operating frequency to reduce their power consumption, the filter would require a high operating frequency in order to reduce its order. Reasonably an optimum exists which is different from the minimum theoretical operating frequency of 32.5 MHz.

The output of the phase accumulator addresses a ROM in which the amplitudes of a sine wave are stored. In practice the ROM limits the usable length of an input phase-word.¹⁰ Therefore, the phase-word length is larger than the ROM address word length and simply a truncation is used on the phase word. This is one of the main sources of spurious tones in the output spectrum of a DDFS.

The harmonic content at the output of the DDFS will be thus determined by three main sources:

- Images of the fundamental frequency due to sampling nature of the system
- Spurs from the phase truncation and DAC quantization error¹¹

⁹Smaller data-rate are possible but they will require a larger number of bits in the phase accumulator.

¹⁰A 23 bit word for the phase information is a moderate choice in DDFS applications. If every phase step is mapped in an amplitude step with, for example, 8 bits precision than a 64 Mbit ROM would be required. This can be reduced to 16 Mbit if the quarter-wave symmetry of the sine is exploited but it is still a too large value for a power constrained system.

¹¹The quantization error is generally treated as white noise. In a DDFS this approximation is valid only in the limiting case when the numerical repetition period of the ROM quantization error is long. In our case, the maximum output frequency is comparable with the clock frequency (see Sect. 3.6.5) and therefore, this approximation is not valid anymore. In the case of a short numerical period of the ROM quantization error, a number of spurs will appear around the fundamental as an effect of the amplitude quantization in the DAC.

- Harmonics of the fundamental frequency due to DAC INL

The image filter mainly affects the filter requirements in terms of filter order. A DAC has at the output a SH functionality. The effect of this block is to shape the frequency components of the wanted signal in a $\frac{\sin(x)}{(x)}$ fashion where $x = \frac{\pi(f_{\text{CLK}} - f_{\text{max}})}{f_{\text{CLK}}}$. Given the fact that the system will synthesize one frequency at a time this does not constitute a problem. This characteristic can also be used to relax the filter specifications introducing some oversampling. Therefore, in the worst case condition of a 40 MHz reference clock the first image frequency will be attenuated by roughly 6.5 dB.¹² The attenuation required by the AA-filter in the worst case condition is therefore around 40 dB.

Phase truncation plays a big role in setting the SFDR of the DDFS synthesizer. The maximum spur at the output of a DDFS system in case of a truncation larger than 4 bits is upper-bounded by $-6.02P$ dBc [63] where P is the number of bits in the truncated phase words. To not spoil the SFDR of the system a 46 dB figure is required. This translates in at least an 8 bits phase word address for the ROM (phase-to-amplitude converter). Therefore, with an 8 bits phase word an SFDR of around 48 dB is achieved.

The number of bits in a DAC is determined by its required spurious performance. Spurs in the DAC come from both the quantization error (see footnote 3) and the DAC INL. The level of the highest spur due to phase truncation is 48 dB. To reasonably derive a specification for the DAC, the approach in [64] is used. If the sum of all the harmonics due to phase truncation is considered as “noise”, then a SNR can be defined when only phase truncation is considered. Now, given the large truncation used for the phase word (from 23 bits to 8 bits), the following expression can be used:

$$\frac{S}{N_{\text{phase}}} = 6.02P - 5.17 \quad (3.30)$$

where S is the signal power and N_{phase} is the noise-like power due to spurs coming from the phase truncation. Therefore, the SNR of the truncated phase signal is around 43 dB. The DAC should contribute negligibly to the overall noise¹³ and therefore, the required SNR for the DAC has to be larger than 49 dB. From [64] it can be proven that because the quantization errors are not evenly distributed in one period because of the shortness of the period (e.g. when the maximum frequency is synthesized), the SNR at the DAC output is given by the following relation:

$$\frac{S}{N_{\text{DAC}}} = -3.01 + 6.02N_{\text{DAC}} \quad (3.31)$$

This translates in N_{DAC} of 9 bits.

¹²The first image frequency is in the worst case at $f_{\text{CLK}} - f_{\text{max}}$ where f_{max} is the maximum signal frequency.

¹³It is important to highlight that this noise does not show a flat continuous spectrum, but that it is made up by a large number of spurs due to the phase truncation and quantization error in the DAC.

Lastly, from the SFDR specification the required maximum INL is readily calculated by the following equation [65]

$$INL = \frac{\frac{1}{SFDR} - 2^{-1.5N_{DAC}}}{2^{-N_{DAC}}} \quad (3.32)$$

where the INL is expressed in LSB. Given a 46 dB required SFDR the maximum allowed INL is 2.5 LSB.

The power consumption of the DDFS will be derived in the next chapters considering the power consumption of each block at different clock frequencies. Then the frequency which minimizes the overall power consumption will be chosen as the optimal solution. Given the fact that the choice of the reference clock frequency will heavily affect the filter design, the choice will depend mainly on the filter specifications and thus on the image frequency.

In this sense the reference clock frequency is chosen in such a way that the first image frequency in the worst case condition is placed one, two or three octaves far from the maximum wanted signal. For $f_{CLK} \geq 10k \times f_{max}$ the filter can be implemented completely in a passive way and have virtually zero power consumption. Following this procedure, the following analysis will be carried out at $f_{CLK} = k \times f_{max}$ where $k = 3, 5, 9$.¹⁴

3.6.2 AA-filter Power Consumption

The first step consists in evaluating the required filter order for a given attenuation. Supposing a 6 dB extra margin on the minimum required attenuation of 40 dB the filter must attenuate the image frequency by around 46 dB. Each pole attenuates the unwanted image frequency by 6 dB/octave. Furthermore, the amplifier at the DAC output will have a bandwidth roughly equal to the maximum signal frequency. Therefore, it will act as a first order filter for the image frequency. Table 3.10 summarizes the required number of poles depending on the position of the first image frequency.

The estimation of power consumption in LPFs can be derived following the work in [66]. We will shortly go to the set of equations useful to derive the filter's power consumption.

Table 3.10 Number of poles in the filter per octave increment in the distance between wanted signal and image frequency

	Number of poles	Clock frequency [MHz]
One octave	7	39
Two octaves	3	65
Three octaves	2	117

¹⁴The position of the image frequency is supposed to be $2f_{max}$ or $4f_{max}$ or $8f_{max}$. Now considering that the image frequency is at $f_{CLK} - f_{max}$, then f_{CLK} has to be 3, 5 or 9 times f_{max} .

The basic parameters required for a correct estimation are the noise excess factor (ξ), and the required Dynamic Range (DR). The first step is to evaluate the required DR. The noise level at the DAC has been set to 52 dB below the signal level in the bandwidth of interest. The AA-filter must negligibly contribute to the overall noise. Therefore, the noise of the filter in the same bandwidth must be at least 6 dB below the noise contribution coming from the DAC. This means that a 58 dB DR for the AA-filter is required.

With these inputs and starting from the state-space representation of the wanted filter, the following equation holds:

$$DR = V_{\max\text{-rms}}^2 \frac{\frac{1}{2\pi} \int_{-\infty}^{\infty} |H(j\omega)|^2 d\omega}{\sum_{i=1}^n S_{n,i(\omega)} w_{ii} \max_j k_{jj}} \quad (3.33)$$

Where, w_{ii} and k_{jj} are the diagonal elements of the matrices W and K which can be obtained by solving the following Lyapunov equation sets:

$$A^T W + W A = -C^T C \quad (3.34)$$

$$A K + K A^T = -B B^T \quad (3.35)$$

where the upper index T indicates the transpose matrix and A, B, C, D are matrices in the state space representation.¹⁵ $H(j\omega)$ is the normalized filter transfer function.¹⁶ $V_{\max\text{-rms}}$ is the maximum rms voltage the integrators in the filter can handle without introducing a large non-linearity. If a 1 V supply is considered and a margin of 200 mV is taken then $V_{\max\text{-rms}} = \frac{0.3}{\sqrt{2}}$. Finally, $S_{n,i(\omega)}$ are the input-referred noise spectra of the integrators and they are equal to

$$S_{n,i(\omega)} = \frac{2kT\xi}{C_i} \left(|b_i| + \sum_{j=1}^n |a_{ij}| \right) \quad (3.36)$$

where k is the Boltzmann constant, T is the temperature in Kelvin degrees C_i is the capacitance at the output of each integrator, b_i are the elements of the matrix B and a_{ij} the elements of the matrix A .

Supposing all the capacitances equal to a value C , the only unknown in (3.36) is C , which can be derived for various filter orders. In Table 3.11 the coefficients for the Butterworth filters are shown for the three different filter orders derived previously in this section. The required capacitor values in the three cases are also given in the same table. The excess noise factor has been assumed to be equal to 7 [67].

At this point each integrator must be able to output the wanted signal in the worst-case condition and given the capacitance value required to achieve the DR. Considering the signal a pure sinusoid at the maximum frequency of 13 MHz we

¹⁵ D is zero for any practical single-input single-output time invariant system.

¹⁶ The normalized filter has a bandwidth equal to one. If the bandwidth is ω_c , the matrices A and B should be multiplied by ω_c .

Table 3.11 Butterworth normalized polynomial coefficients, overall filter capacitance and predicted power consumption

Filt. Order	c_0	c_1	c_2	c_3	c_4	c_5	c_6	c_7	Cap. [pF]	Pow. cons. [mW] ^a
2	1	1.41	1	–	–	–	–	–	1.85	0.09
3	1	2	2	1	–	–	–	–	4.2	0.31
7	1	4.49	10.09	14.59	14.59	10.09	4.49	1	172	29.5

^aAt $V_{\text{supply}} = 1 \text{ V}$ and $V_{\text{max}} = 0.3 \text{ V}$

can write the following equations:

$$I = C \frac{dV}{dt} \quad (3.37)$$

$$V = V_{\text{max}} \sin(\omega_{\text{max}} t) \quad (3.38)$$

where V_{max} is the sinusoid peak voltage (0.3 V can be a suitable choice). From (3.37) and (3.38) we obtain that to prevent slew-rate limitations and, therefore, to keep the distortion low the bias current has to be equal to the peak current delivered to the integrating capacitance and therefore

$$I_{\text{BIAS}} = 2\pi C f_{\text{max}} V_{\text{max}} \quad (3.39)$$

With these values the predicted current consumption for the three different filters is summarized in Table 3.11.

3.6.3 Phase Accumulator and ROM Power Consumption Estimation

The phase accumulator power estimation can be quite cumbersome due to the fact that the dynamic power consumption mainly depends from the statistics of the signal applied to the various logic depths. For this reason, instead of developing an involved model for the power consumption, the block has been simulated in a CMOS 90 nm process and the results are given in Table 3.12. As it can be seen, the power consumption is low compared to the AA-filter power consumption. On the other hand the power consumption can be easily scaled up or down for different technologies and frequencies by using the following equation:

$$P_{\text{acc}} = f_{\text{clk}} C_T V_{\text{supply}}^2 \quad (3.40)$$

where V_{supply} is the supply voltage. Therefore, the power scales linearly with frequency and also with the overall block capacitance C_T . Now if the ratio between a more advanced technology minimum length and the simulated 90 nm technology is α then the overall capacitance will vary as roughly α^2 for a constant aspect ratio.

The phase word as mentioned in Sect. 3.6 is the truncated 8-bit word while the amplitude information must be coded at least with 9 bits. The size of the ROM is

Table 3.12 Phase accumulator power consumption estimation

Clock frequency [MHz]	Power consumption [μ W]
39	35
65	52
117	83

therefore, about $(2^8 \times 9)$ 2304 bits and a $2^4 \times 2^8$ memory matrix is sufficient. Let us consider $n - k = 4$ rows and $k = 8$ columns.¹⁷

A simplified block diagram for a ROM is given in Fig. 3.16. Following the work in [68] the ROM power consumption can be evaluated by considering all the different components involved in a single phase to amplitude conversion. The power consumed by the 2^k memory cells on a row during one pre-charge or evaluation can be approximated by the following expression

$$P_{\text{memcells}} = \frac{2^k}{2} (c_{\text{int}} l_{\text{column}} + 2^{n-k} C_{\text{tr}}) V_{\text{supply}} V_{\text{swing}} \quad (3.41)$$

where P_{memcells} is the approximated power consumption of the ROM memory core, c_{int} is the capacitance of a unit wire length with minimum width, C_{tr} is the minimum size gate capacitance, and V_{swing} is the voltage swing of each memory cell. Defining the memory cell as $d_m \times d_m$ square, the column interconnection length of the memory matrix is $l_{\text{column}} = 2^{n-k} d_m$.

Now considering as an example the 90 nm CMOS technology, the parameters shown in Table 3.13 are available. With these values $P_{\text{memcell}} = 718$ fW/operation. Concluding, the power consumption for the three different operating frequencies is given in Table 3.14.

The row decoding part power consumption is negligible and therefore, it will not be considered. The power from row-driving is given by [68]

$$P_{\text{rd}} = 0.5 \{ 2^{k+1} C_{\text{tr}} + 2(n-k) C_{\text{tr}} + c_{\text{int}} [8(n-k) W_{\text{int}} + l_{\text{row}}] \} V_{\text{supply}}^2 \quad (3.42)$$

while the column select power consumption is given by

$$P_{\text{cs}} = \frac{1.3}{2} \left(\sum_{i=1}^{k-1} 2^{k-i} C_{\text{tr}} + k c_{\text{int}} l_{\text{column}} \right) V_{\text{supply}}^2 \quad (3.43)$$

where $l_{\text{row}} = 2^k d_m$. The power consumption summary is given in Table 3.15.

The last contribution mainly comes from the sensing amplifier. This contribution is equal to $V_{\text{supply}} I_{\text{sense}}$ and it is frequency independent. The sense amplifier consumes approximately 90 μ W as can be derived from Table 3.13. The overall ROM and accumulator power consumption for the three different operating frequencies is given in Table 3.16.¹⁸ At this point it is possible to make an important observation.

¹⁷Especially for bigger ROMs several optimization techniques can be used to reduce the ROM power consumption.

¹⁸By using simple logic functions it is possible to use the same ROM for both sine and cosine generation. For this reason the power consumption of a single ROM as been considered even if quadrature generation is required.

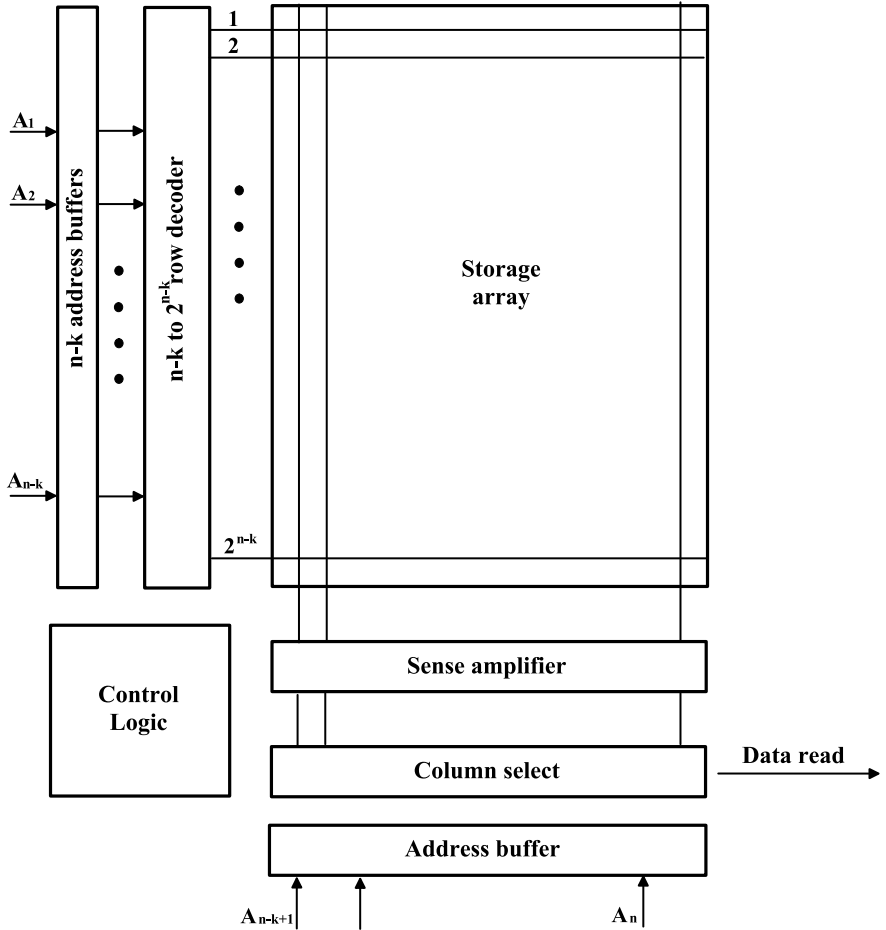


Fig. 3.16 General memory structure

Table 3.13 90 nm CMOS Technology Parameters

Parameter	Value	Unit
C_{int}	90 ^a	aF/ μm
C_{tr}	220	aF
d_m	1.45	μm
V_{dd}	1	V
V_{swing}	1	V
W_{int}	0.14	μm
I_{sens}	90	μA

^aMinimum interconnection width equal to 0.14 μm

Table 3.14 Memory-cell evaluation power consumption

Clock frequency [MHz]	Power consumption [μ W]
39	28
65	46.6
117	84

Table 3.15 Row-decoding and column select global power consumption

Clock frequency [MHz]	Power consumption [μ W]
39	4.8
65	7.9
117	14.2

Table 3.16 Overall phase accumulator and ROM power consumption

Clock frequency [MHz]	Power consumption [μ W]
39	145.4
65	184.5
117	259.4

It is common to believe that in a DDFS the ROM consumes most of the power. This is true for a high demanding DDFS while for a power constrained DDFS the situation is different. Indeed, the highest power consumption for the digital DDFS back-end is around 300 μ W and therefore, far from the commonly reported figure of some tens of mW of high-end DDFSs.

3.6.4 DAC Power Consumption Estimation

The last step consists in predicting the DAC power consumption with the reference clock frequency. The architecture choice is driven by reduction of the power consumption and the capability to handle the high speed operations. The most common DAC architectures are the following

- Oversampled $\Sigma\Delta$
- R - $2R$
- Charge redistribution
- Current steering

The oversampled DACs are mainly used to achieve high resolution at low frequencies (<10 MHz). Furthermore, they involve analog operations and can be very complex at high frequencies. The current steering architecture is used at lower resolution and higher speed. This architecture has as a main drawback the higher power consumption, which makes it not suitable for ultra-low power applications. Therefore, these two architectures will not be discussed further in this book.

For a relative high frequency low power DAC, a very efficient architecture is the R - $2R$ or the charge-redistribution topologies. In these topologies most of the current will be used (see Sect. 3.6.4) only to drive the following stage (the AA-filter) and not for the conversion itself.

R - $2R$ DAC

There are two ways in which the R - $2R$ ladder network may be used as a DAC, known respectively as the voltage mode and the current mode. They are shown in Figs. 3.17(a) and (b). In the voltage mode R - $2R$ ladder DAC, the “rungs” of the ladder are switched between V_{ref} and ground, and the output is taken at the end of the ladder. The voltage output is an advantage of this mode (because the gain is decoupled from the impedance level of the ladder), as is the constant output impedance, which eases the stabilization of any amplifier connected to the output node. Additionally, the switches shift the arms of the ladder between a low impedance V_{ref} connection and ground, so capacitive glitch currents tend not to flow in the load.

The main advantage of the current mode topology is its speed because the input of the buffer is at virtual ground. The normal connection of a current-mode ladder network output is to an OPAMP configured as current-to-voltage (I/V) converter, but stabilization of this OPAMP is complicated by the DAC output impedance variation with digital code. Of course capacitive glitches are larger for the current-mode topology with respect to the voltage-mode. Glitches create an unwanted harmonic distortion during waveform generation and therefore, a voltage-mode topology is used. Consequently also the AA-filter will be supposed to work in voltage-mode.

In an R - $2R$ DAC there are three main sources of power consumption:

- R - $2R$ resistive chain
- Output buffer
- Switch drivers

The power consumption of the resistive ladder can be derived by using the following formula [69]

$$P = \sum_{i=1}^n d_i I_i V_{\text{ref}} \quad (3.44)$$

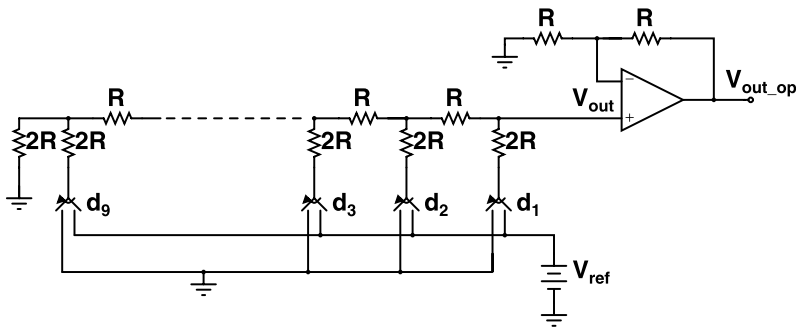
where V_{ref} is the DAC reference voltage, d_i ($i = 1, 2, \dots, n$) is the n -bit binary number to be converted into the analog voltage and I_i is the current flowing through each leg of the resistive chain and it is given by the following equations

$$I_1 = \frac{d_1 V_{\text{ref}} - V_{\text{out}}}{2R} \quad (3.45)$$

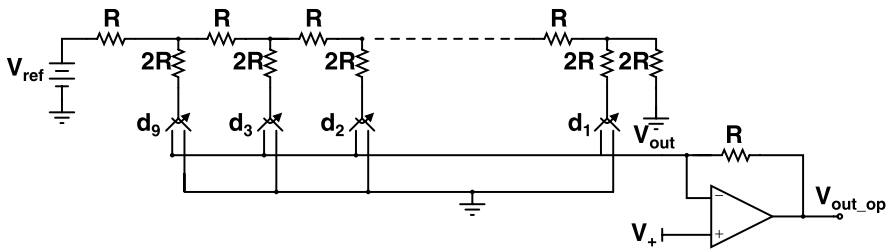
$$I_{i|(i=2,3,\dots,n+1)} = \frac{d_i V_{\text{ref}} - [V_{\text{out}} - R \sum_{k=1}^{i-1} (i-k) I_k]}{2R} \quad (3.46)$$

where $d_{n+1} = 0$ and V_{out} is the output voltage of the DAC and it is equal to

$$V_{\text{out}} = (d_1 2^{-1} + d_2 2^{-2} + \dots + d_n 2^{-n}) \times V_{\text{ref}} \quad (3.47)$$



(a) Voltage mode R-2R DAC



(b) Current mode R-2R DAC

Fig. 3.17 R-2R DAC topologies

The power consumption of the R -2 R ladder is a function of the input code. Now, supposing that the probability that a certain code is addressed is equal for all the codes, then the rms power dissipation can be calculated using the following equation:

$$P_{\text{rms}} = \frac{1}{N_{\text{code}}} \times \sum_{j=1}^{N_{\text{code}}} P_j \quad (3.48)$$

where N_{code} is the total number of codes in the DAC (e.g. for a 9 bit DAC it equals $2^9 - 1$).

As can be seen from (3.45) and (3.46) the power consumption depends on the resistance value R . At each node of the R -2 R ladder in Fig. 3.17 there will be a parasitic capacitance due to interconnections and junctions. Furthermore, the input of the operational amplifier is not at virtual ground as it would have been in the case of a current-mode R -2 R topology. This means that between the time at which the digital word is set and the time at which the converted voltage will be available at the input of the output buffer there will be a delay caused by the buffer input capacitance. This delay should not exceed 50% of the time available for a single conversion, which depends on the DDFS operating frequency.

The maximum delay time and maximum value of R in the three different cases are shown in Table 3.17. The values have been derived via simulations. The resis-

Table 3.17 Maximum delay time and R value for the R - $2R$ ladder

Op. freq. [MHz]	Max Delay [ns]	$R_{\max}$ [k Ω]	(Power Cons. [μ W])
39	12.8	22	10.1
65	7.7	13	17.1
117	4.3	7.5	29.6

tance value has been derived through simulation of the same ladder including the parasitics at the internal nodes. In the same table the power consumption of the R - $2R$ ladder in the three possible cases is also shown. V_{ref} has been considered equal to 0.5 V.¹⁹

Another limit to the maximum value of resistance in the ladder comes from noise considerations. Indeed the noise added by the chain equivalent resistance and the feedback resistances must fulfill the required SNR specification. In Sect. 3.6.1 a 9-bit DAC has been chosen. This, from (3.31), translates in an SNR equal to 51 dB. Supposing that the noise contribution of the OPAMP and the one of resistive ladder have to be 6 dB below the quantization noise, their SNR must be 60 dB.

Supposing to have the noise from resistances 60 dB below the signal level which has a peak value V_{peak} the following equations hold:

$$\frac{V_{\text{peak}}^2}{12kTRf_{\max}} \geq 10^6 \quad (3.49)$$

where k is the Boltzmann constant and T the temperature in Kelvin degrees. From the previous equation the resistance value has to be smaller than 83 k Ω (with $V_{\text{peak}} = 0.25$ V) which is in line with the values in Table 3.17.

The output buffer of the DAC needs to drive the input impedance of the following stage (the AA-filter). The filter for the three different cases (different image position) can be designed by using standard software tools. These tools require generally a typical resistor value. Of course the value of this resistor will set the required capacitance and ultimately the input capacitance of the filter. The maximum resistor value can be derived from noise considerations. The filter must contribute negligibly to the overall noise. Therefore, having set earlier the DAC SNR to 51 dB (the system is spurs dominated) and because no gain is achieved up to the AA-filter input it is desirable to have the noise floor of the filter well below the required SNR. Therefore, an SNR of 58 dB is a good choice. Now from [70]

$$\frac{V_{\text{rms,sig}}^2}{V_n^2} = \frac{V_{\text{rms,sig}}^2}{4kT \times nR\xi \times BW} \quad (3.50)$$

where $V_{\text{rms,sig}}^2$ is the rms signal value, V_n^2 is the noise power, n the filter number of poles, R the typical resistor value, ξ the excess noise factor and BW the filter bandwidth. The maximum resistance value and filter input capacitance for the three different frequencies are given in Table 3.18.

¹⁹The non-inverting configuration of the OPAMP has a gain equal to two.

Table 3.18 Typical Filter resistance, filter input capacitance and DAC buffer current consumption versus frequency

Clock freq. [MHz]	Max filt. res. [kΩ]	AA-Filt. input cap. [pF]	Buffer Power cons. [μW]
39	14	0.76	318
65	33	0.32	121
117	49	0.16	60

Table 3.19 90 nm CMOS technology parameters

Parameter	Value	Unit
C'_{ox}	7.8	$\frac{\text{fF}}{\mu\text{m}^2}$
μ_n	200	$\frac{\text{cm}^2}{\text{V}\cdot\text{s}}$
Early Voltage V_A	2.75	V

For the OPAMP in the output buffer a gain of 100 is required to avoid non-idealities coming from the amplifier finite gain. All the parameters required for the evaluation of the buffer power consumption for the CMOS 90 nm technology are listed in Table 3.19. To reduce the current consumption, it is reasonable to drive the transistors in weak inversion. Given the required voltage gain A_{V0} , the minimum channel length to achieve this gain with a single stage amplifier is given by the following equation [71]

$$L = \frac{n\phi_t A_{V0}}{V_A} \frac{1 + \sqrt{1 + i_d}}{2} \quad (3.51)$$

where ϕ_t is the thermal voltage, roughly equal to 25 mV, n is the slope factor ($n \simeq 1.6$ in weak inversion), V_A is the Early voltage and i_d is the inversion coefficient (between 0.1 and 1 in weak inversion). Considering an inversion coefficient equal to 0.5 the required transistor channel length for a gain of 100 is equal to 1.6 μm.

The relation between current and gate-source voltage in a CMOS transistor is given by the following relation

$$I_{DS} = \frac{1}{2} \mu_n C'_{ox} \frac{W}{L} (V_{GS} - V_T)^2 \quad (3.52)$$

where μ_n is the electron mobility, C'_{ox} is the gate capacitance per unit area, W is the transistor width, L the channel length and $V_{GS} - V_T$ is the overdrive voltage.

Now it can be proven [72] that the bandwidth of an amplifier can be expressed as a function of the bias current by the following equation

$$BW = \frac{3 \cdot I_{DS}}{2\pi \cdot (V_{GS} - V_T) \cdot C'_{ox} W L} \quad (3.53)$$

The inversion coefficient i_d [73] can be also expressed as a function of the bias current with the following equation

$$i_d = \frac{I_{DS}}{2n \cdot \mu_n C'_{ox} \frac{W}{L} \cdot \phi_t^2} \quad (3.54)$$

Now from (3.52), (3.53) and (3.54) the maximum achievable bandwidth can be expressed as a function of the channel length

$$BW = \frac{3}{2\pi} \sqrt{ni_d} \frac{\mu_n \phi_t}{L^2} \quad (3.55)$$

From (3.55) substituting all the known values and a channel length of 1.6 μm , a maximum achievable bandwidth of 81 MHz is obtained. This proves that a single stage amplifier can achieve at the same time the required speed and gain given the required bandwidth of 13 MHz.

The minimum bias current to achieve the required specifications is given by the following relation [71]

$$I_{DS} = n\phi_t(2\pi GBC_L) \quad (3.56)$$

where GB is the gain-bandwidth product and C_L is the load capacitance. The buffer estimated power consumption for the three different frequencies is given in Table 3.18.

To avoid performance degradation in the DAC, the switch on-resistance has to be much lower than the $2R$ value in the R - $2R$ ladder. The actual required resistor matching can be derived by the following equation [74]

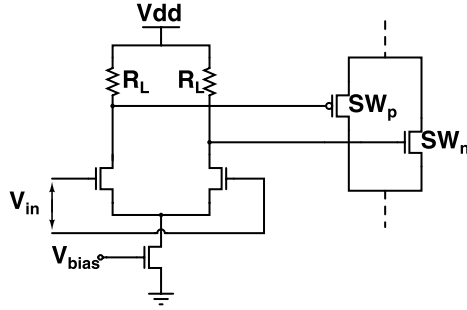
$$T \leq \frac{INL}{0.577 \times 2^N} \quad (3.57)$$

where T , as defined in [74], is the resistor tolerance. This contribution comes both from the actual resistor and the switch resistance. Given the fact that they are stochastically independent they sum in a quadratic form and if we suppose their contribution equal, then $\Delta R_{TOT} = \sqrt{2}\Delta R_{SW}$ where ΔR_{TOT} and ΔR_{SW} are the maximum allowed variance in a single ladder branch and the variance of the switch resistance respectively. In this way supposing a $\pm 30\%$ mismatch of the transistor on-resistance the maximum resistance value can be calculated for the LSB switch for the three different operating frequencies (see Table 3.20).

Table 3.20 Switch parameters

Frequency [MHz]	R_{on}^{max} [Ω]	$W \cdot L$ [μm^2] ^a	Gate capacitance (C_{SW}) [fF]
39	437	0.86	10.9
65	259	1.36	18.4
117	148	2.43	32.8

^aAll switches are considered minimum length

Fig. 3.18 Switch driver schematic**Table 3.21** Time constants, maximum driver load resistance LSB current and overall power consumption for different DAC operating frequencies

Freq. [MHz]	Max time const. ($\tau_{\max}$) [ns]	$R_L^{\max}$ [k Ω]	I_B^{LSB} [μA]	P_{TOT} [μW] ^a
39	2.56	117	7.7	69.3
65	1.54	42	21.4	192.6
117	0.85	13	69.6	626.4

^aAt $V_{\text{supply}} = 1$ V

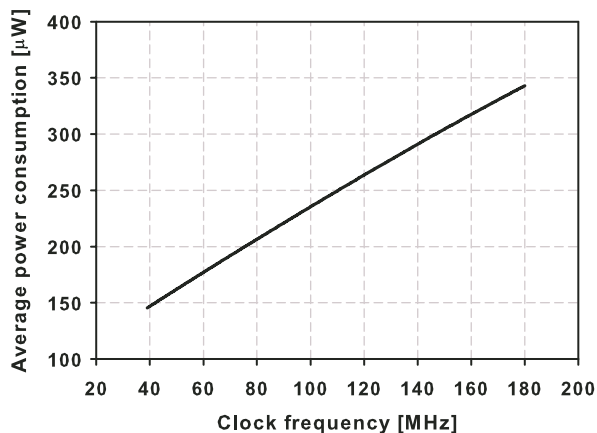
A very common transistor level implementation of a CMOS switch together with the driving circuit is depicted in Fig. 3.18. The driver should be able to switch the gate transistors on and off very rapidly. Therefore, the transition edges have to be very sharp. This means that the time constant at the gate of the nMOS transistor (SW_n) and pMOS transistor (SW_p) must be very small compared to the DAC sample time. Considering a ratio of 10 between the DAC sample time and the switching time for the LSB, Table 3.21 summarizes the maximum RC time constant at the switch gate and therefore, the maximum load resistance R_L for the various DAC operating frequencies.

At this point the driver must be able to switch its voltage outputs between V_{supply} (when in one of the branches no current is flowing) and $2 \times (V_{\text{GS}} - V_{\text{TH}}) \simeq 0.1$ V to allow the current source transistor in Fig. 3.18 to work in saturation. Considering a 1 V power supply the driver has a peak to peak voltage swing of about 0.9 V single-ended. Therefore, the required DC current for the LSB switch driver is simply given by $I_B^{\text{LSB}} = \frac{V_{\text{swing}}}{R_L^{\max}}$ and the values for the three different frequencies are given in Table 3.21. The time constant at the input of the switch is roughly given by $R_L \times 2C_{\text{SW}}$ ²⁰ where C_{SW} is given in Table 3.20 for the three different frequencies.

The total current consumption for all drivers is $9 \times I_B^{\text{LSB}}$. The drivers current consumption for the three different DAC operating frequencies is given in Table 3.21.

²⁰The factor two takes into account that we need two switches per branch because each branch can be switched either towards V_{ref} or ground.

Fig. 3.19 Digital back-end power consumption versus reference clock frequency



Charge Redistribution DAC

The resistive ladder can be replaced by a capacitive ladder in which the resistance R is replaced by a capacitance $2C$ and the resistance $2R$ by a capacitance C . This architecture has a big advantage in terms of both area and power consumption over the common weighted capacitance architecture as shown in [75]. This topology has as main advantage over the R - $2R$ architecture the fact that no DC current is flowing. Indeed, the energy consumption only takes place during charging or discharging cycles. On the other hand, capacitive dividers are sensitive to stray capacitances which can heavily affect the accuracy.

In [75] it has been shown that on Silicon substrate because of the bottom plate parasitic capacitance it is impossible to realize even a 4 bits C - $2C$ DAC. Therefore, a simpler weighted sum architecture has to be used but as stated before the power consumption becomes comparable with a common R - $2R$ architecture without any other advantage.

3.6.5 Power Dissipation of the Whole DDFS

From Fig. 3.8 the overall DDFS power consumption for a quadrature output is given by the following relation:

$$P_{\text{DDFS}} = P_{\text{dig}} + 2 \times P_{\text{DAC}} + 2 \times P_{\text{AA}} \quad (3.58)$$

where P_{dig} is the digital back-end power consumption, P_{DAC} is the DAC power consumption and P_{AA} is the AA-filter power consumption. The power consumption of the digital blocks versus the operating frequency is shown in Fig. 3.19. As expected the overall power consumption grows roughly linearly with the reference clock frequency.

Fig. 3.20 DAC power consumption

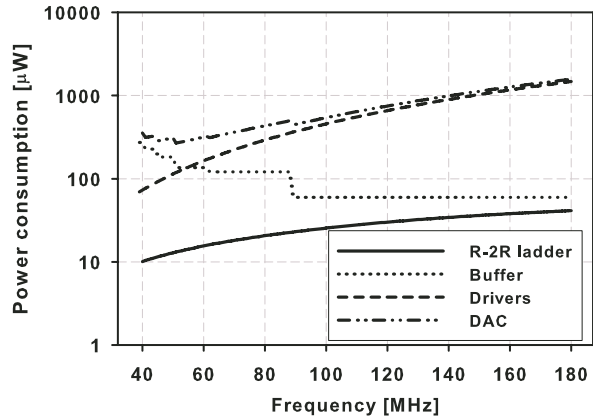
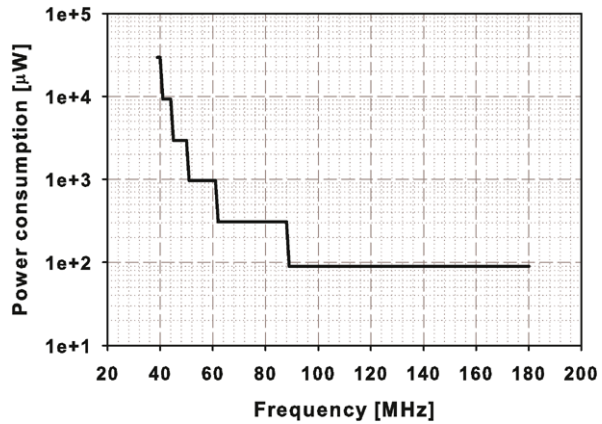


Fig. 3.21 AA-filter power consumption



The DAC power consumption is depicted together with the power consumption of its building blocks in Fig. 3.20. As it can be seen, the *R-2R* ladder power consumption is negligible. A minimum in the overall DDFS power consumption is located around 90 MHz (see Fig. 3.22). On the left side of the minimum the power is dominated by the output buffer. Indeed, lowering the reference frequency requires an increase in the filter order and consequently the filter input capacitance grows. As a consequence the output buffer bias current needs to increase. The power consumption of the output buffer has a typical staircase structure due to the fact that a given order for the filter it has as a load fulfills the requirements in terms of image attenuation for a certain range of frequencies. On the right side of the optimal reference frequency, the overall power consumption is dominated by the switch drivers.

The filter power consumption is plotted on a logarithmic scale in Fig. 3.21, while the overall power consumption is plotted in Fig. 3.22. A minimum is located at around 90 MHz. On the left side the filter will dominate the overall power consumption, while on the right side the drivers in the DAC are again the most power hungry blocks.

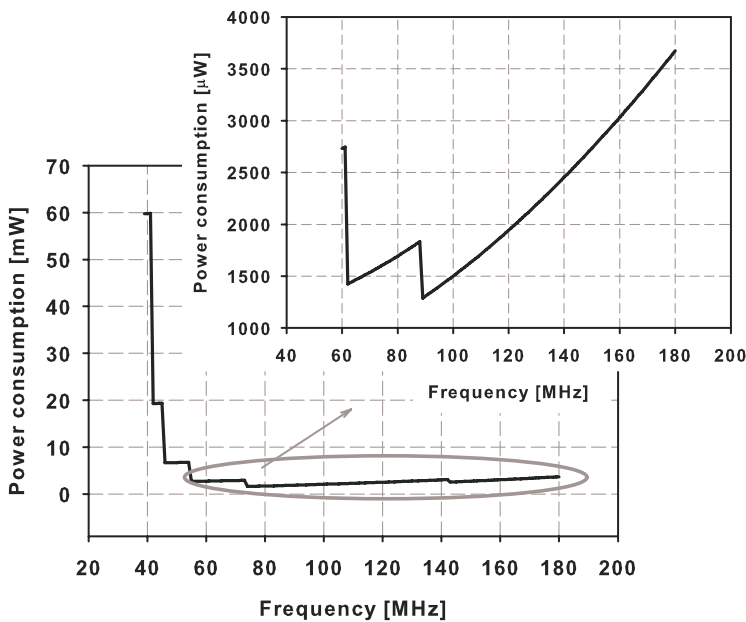
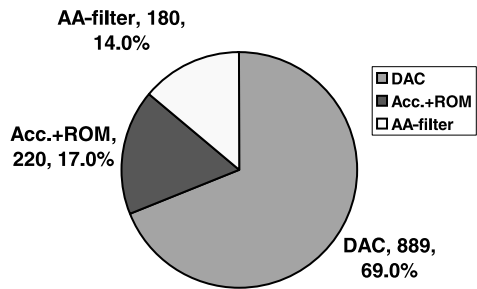


Fig. 3.22 DDFS power consumption

Fig. 3.23 Power consumption break-down by block (power in μW)



The minimum power consumption is around 1.3 mW for the complete DDFS system. Given the 90 MHz reference clock frequency, the power consumption break-down is given in Fig. 3.23. The most power hungry blocks are the DACs with almost 70% of the overall power consumption. Globally, the analog blocks account for more than 80% of the overall DDFS power consumption.

The attenuation of the image frequency due to the sinc function has been derived in Sect. 3.6 considering an initial reference frequency of 39 MHz. Now with the new reference frequency it is possible to recalculate this value to check if it is possible to reduce the number of bits or the filter order. If this would be the case the overall procedure needs to be repeated iteratively till convergence. The extra attenuation due to the sinc function with respect to the 39 MHz case is only 1.4 dB and therefore,

not enough to reduce the filter order below the optimal value of two. Concluding, the previously calculated values are still valid.

3.7 Summarizing Discussion

A PLL system has the advantage to directly synthesize the wanted channels at high frequency. This avoids the use of any kind of up-conversion stage, which generally requires a mixer and a BPF filter to remove the unwanted components or a quadrature up-conversion scheme. Unfortunately, a problem arises when a fine channel spacing is required together with a very short settling time. This generally translates in a quite high power consumption for a given SFDR. To relax the specifications and to achieve a larger loop bandwidth with a fine inter-channel spacing a fractional- N architecture can be used. Nevertheless, a more complicated architecture is required, fractional spurs can pollute the spectrum and still the power consumption increases by decreasing the settling time. In Sect. 3.1 it has been shown that from a power point of view a serial search algorithm should be used as a synchronization algorithm. This is true in a two-way link while in a one-way link the receiver is not power constrained and therefore, any kind of algorithm can be used.²¹ Looking at the two way link, and at the state-of-the-art PLL based synthesizers it is clear that it is very difficult to meet the required power target at a settling time below 5 μ s.

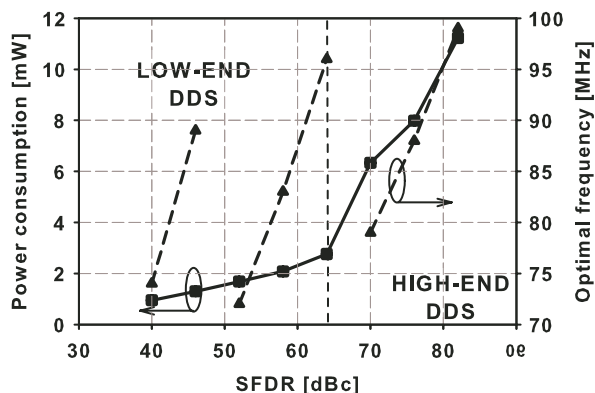
A possible way to overcome the problem is to synthesize the required channels at baseband. In this way the VCO gain can be relaxed and therefore, the power consumption can decrease because spurs decrease proportionally to the VCO gain. Still a connection exists between settling time and power consumption. Looking at the power model developed in Sect. 3.5 and at the state-of-the-art summarized in Table 3.6 it is clear that PLL synthesizers with a settling time below few microseconds come at the expense of a large power dissipation. This will drastically reduce the transmitter and transceiver efficiency.

The main reason why a short settling time requires a high power consumption comes from the fact that any closed loop system works on a-posteriori information. A way to remove this constraint and to fully use the capability of the FHSS system to trade power for hopping rate as explained in Sect. 2.1.5 is to use a-priori information. This can be done by removing the loop and by using a feed-forward approach. A DDFS achieves this result but it requires an up-conversion stage to reduce its power consumption. Therefore, all the required frequency bins are synthesized at baseband and then they are up-converted in quadrature.

Most of the DDFSs are designed for very high SFDR. Given the low probability of collision and the use of a frequency diversity scheme, a much lower

²¹In a one way link anyhow the pre-PA power consumption must be minimum and ideally zero. In this way the transmitter efficiency can be limited only by the PA efficiency. In the next chapter a TX architecture for a one-way link is proposed which reduces the pre-PA power consumption below 200 μ A.

Fig. 3.24 Minimum power consumption and optimal reference clock versus SFDR



SFDR can be used without affecting the reliability of the communication link. A model has been developed using a combination of basic blocks, which can be seen as an optimum trade-off between power dissipation and high level specifications.

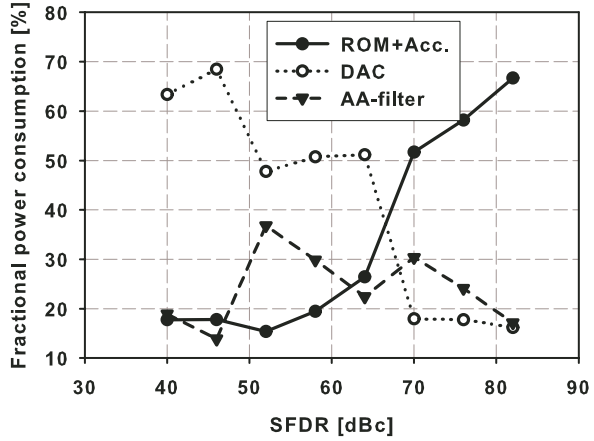
In Fig. 3.24 the minimum power consumption versus the SFDR is shown. In the same graph the optimal reference frequency is also shown. It can be seen that there are two trends. Below roughly 64 dB SFDR the system is dominated, in terms of power consumption, by the analog part (DAC+AA-filter). The power increases with a slope which is lower than in the case of an SFDR above 64 dB. This can be explained with the fact that the maximum operating frequency is always constrained to 13 MHz and the dynamic range of the filter has been kept constant (58 dB).²² The increment is due to an increase in the filter requirements and in the digital back-end. Above 64 dB of SFDR the digital back-end starts to consume most of the power due to the huge increment in the ROM size. Therefore, in the right area the power consumption is dominated by the digital back-end.²³ This area is the area of the high-end DDFSs.²⁴ In the area of high-end DDFSs the choices made in the previous section may be not valid any more and a better solution has to be found for power optimization. To further confirm these results, the power breakdown versus the SFDR is shown in Fig. 3.25. The crossing point of 64 dB is the point at which the power consumption of the digital back-end exceeds the power consumption of the analog components and sets the new trend in the overall power

²²Getting a high dynamic range from an analog block can cause a significant increase in the power consumption. Therefore, if the digital back-end is designed for higher SFDR (the cost in terms of power consumption to obtain this requirement is decreasing with technology scaling), it is possible to conceive a dedicated pre-distortion algorithm or a notch filter to cancel or reduce the harmonic distortion of the analog filter without using the brute force approach (increase the power consumption).

²³The phase accumulator has a constant power consumption. Its size depends only on the wanted resolution and reference clock frequency.

²⁴The DAC required number of bits in this area is above 12. Therefore, to still keep the same DAC topology a laser trimming of the resistances or a calibration loop are required.

Fig. 3.25 Power consumption break-down versus SFDR



consumption. As already known from literature at high SFDR the power consumption is dominated by the ROM power consumption and the previous graphs confirm this rule. What can be concluded for a DDFS designed for low-data rate applications is that most of the power is used in the analog blocks. This power is mainly used for driving purposes (switches or driving of a following stage like in the DAC). This means that this power does not scale down with technology and it is fundamentally bounded to the range of frequency to synthesize and to the system requirements (like the attenuation of the image frequency). Therefore, there are little margins of improvements in this kind of architecture. Moreover this is a theoretical bound, which does not take into account parasitic effects, required biasing circuitry and all the second order effects of a real transistor level design. Therefore, the actual power consumption can be higher than the predicted one by a factor of two or more when a physical realization is obtained. This means that the pre-PA power consumption still dominates the overall power consumption at low transmitted powers.

3.8 Conclusions

In the previous sections it has been shown that current architectures for FHSS synthesizers cannot achieve at the same time agility and very low power consumption. Indeed, the closed loop systems like PLLs, becomes very power hungry at very low settling time. DDFS systems can achieve very good agility but at the costs of higher complexity and, therefore, higher power consumption.

To be able to reduce the power consumption of the synthesizer so that the transmitter and therefore, the transceiver efficiency can be boosted considerably, a novel architectural approach in the design of synthesizers for ultra-low power FHSS systems is required. The new synthesizer should be able to achieve a fast settling time as a DDFS but at a much lower power consumption. This will highly improve the

efficiency in the transmitter for a one-way link and improve also the receiver efficiency in the two-way link. To achieve the target the architecture must exploit the benefits of each scenario and therefore, two novel synthesizer architectures are proposed in this book. The description of those architectures are the main topic of the next two chapters.

Chapter 4

A One-Way Link Transceiver Design

As stated in Chap. 1, the application area of WSNs can be subdivided in two smaller areas:

- Applications requiring only to transmit data
- Applications requiring to transmit and receive data

In the first area most of the applications fall, which are related to domotica and to a large number of applications that need to sense a physical parameter like the temperature, the pressure, the humidity etc. This chapter focuses on this class of applications.

A novel architecture is disclosed, which exploits the low requirements for these kinds of networks which, from now on, will be referred to as asymmetric networks. This architecture reduces the requirements of the wireless autonomous node (the transmitter) to a minimum, while shifting the complexity to the mains supplied node (the RG). The RG can be connected via a wireless or a wired connection to the outside world allowing the asymmetric network to be connected to remote sites.

4.1 General Guidelines for Transmitter Design

To reduce the costs for the user, the transmitter nodes and the RGs use the ISM bands, which are license free. The commonly used ISM bands are the 915 MHz band (available only in the US) and the worldwide available 2.4 GHz band. In the recent years a large number of new standards made the 2.4 GHz band highly overcrowded both in the number of users and in the power levels employed. Furthermore, the presence of continuous interferers, like the microwave oven makes this band a very harsh environment for ultra low-power networks. For this one-way link architecture, the 915 MHz band was chosen. However, the described techniques can be equally used in the 2.4 GHz band.

As previously mentioned in Sect. 2.1.5, a commonly used modulation format for FHSS systems, is the FSK modulation format. For the prototype proposed in

this book, a robust wideband BFSK modulation scheme was employed (modulation index larger than five).

High level of integration and reduction in the average power consumption can be achieved by employing a simple direct up-conversion scheme. Nevertheless, this architecture suffers from an important drawback: the disturbance coming from the PA output, being injected in the local oscillator. This drawback can be alleviated by shifting the output spectrum far from the LO frequency [76] or by shielding the VCO tank from the antenna by using a buffer. In this way it is possible to combine, in a single block, both the VCO required for the up-conversion and the PA required to transmit the information through the channel. Those two architectures (the VCO-divider based and the power-VCO based FHSS transmitter architectures) are disclosed more in detail in Sects. 4.2.4 and 4.2.5.

Prior to any data transmission, transmitter and receiver should synchronize their center frequencies and PNCs in time. Given the limited power budget on the transmitter side, this process has to be fast and mostly handled by the receiver. Therefore, the transmitter will first calibrate its center frequency in such a way that all the hopping channels lie inside the specified bandwidth (i.e. 915 MHz ISM band). The RG will then align its center frequency to the transmitter center frequency employing a dedicated algorithm [77]. This algorithm is disclosed more in detail later in this chapter. Data transmission will start after synchronization of the PNCs.

4.2 Transmitter Architecture

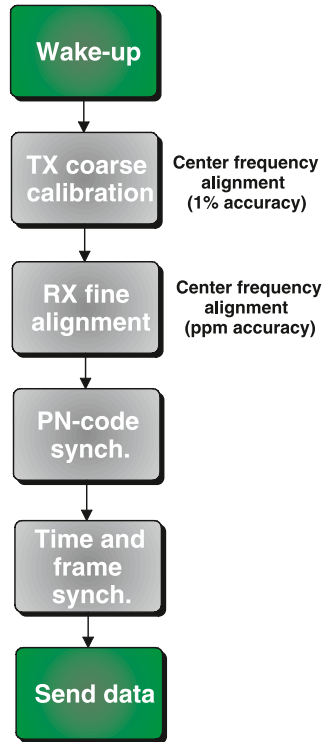
In Chap. 3 it has been shown that traditional FHSS systems are based on a Phase-Locked Loop (PLL) with a digitally controlled variable divider or on a Direct Digital Frequency Synthesizer (DDFS) and a Digital to Analog Converter (DAC) that is used to translate the discrete time periodic waveform from the DDFS in a continuous waveform with specified spectral characteristics. Both these methods require complex hardware and therefore, a high power consumption is expected when implemented on silicon.

In this chapter a new architecture concept is proposed, which requires simple digital techniques and low complexity, low frequency analog building blocks, based on frequency pre-distortion. In this architecture, two sources of error in the frequency synthesis exist:

- Synthesizer center frequency spread
- Non-linearity induced frequency errors

The first error is common to any TX and RX node and it is caused by the components process spread in silicon implementation. Before any communication starts, the local oscillators at the TX and RX sides must be aligned within part per million accuracy. The second error source is specific to the proposed architecture and it will be discussed more into detail in Sect. 4.2.1. Before any data transmission occurs a

Fig. 4.1 TX operation flowchart

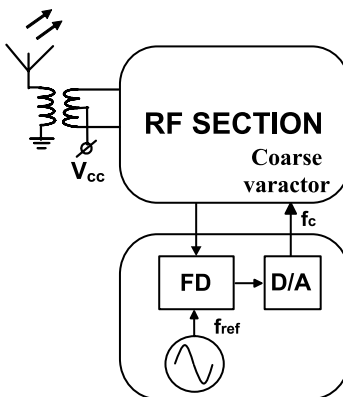


sequence of required operation must be performed between TX and RX in order to assure that a reliable communication link is established. The various operation steps are summarized in Fig. 4.1.

The node first wakes up. At this point the TX cannot transmit because there is a possibility that its center frequency is shifted from the ideal band center frequency. If the misalignment is too large, the TX would transmit outside the ISM band, which is not allowed by any communication regulation rules (like the FCC regulations for example). Therefore, the first step is to align the TX center frequency to the RX center frequency. Because no crystal is used at the TX side, the accuracy of this first alignment is within 1%. This is enough to fulfill the regulation rules, but not enough to start a reliable communication. For this reason, the RX needs to align its center frequency to the TX center frequency within the required ppm accuracy. At this point the PN synchronization starts (see Chap. 3, Sect. 3.1). After the PNCs are synchronized, time and frame synchronization must be achieved. This step is outside the scope of this book and therefore, it is not discussed further. Finally, data transmission can begin.

The proposed concept based on the operation principle described in Fig. 4.1 will be used in the two architectures described in Sects. 4.2.4 and 4.2.5.

Fig. 4.2 Initial calibration loop (FLL)



4.2.1 Concepts and Block Diagrams

The initial center frequency calibration at the TX side can be achieved in two ways:

- Stored in an EEPROM at factory
- Using a dedicated FLL

Using an EEPROM can save some power but it increases the cost of the IC because the coarse calibration needs to be performed on every IC. For this reason the second solution has been adopted here and the conceptual block diagram is shown in Fig. 4.2.¹

During the initial coarse center frequency alignment, the FLL in the base-band measures the VCO signal, and calibrates the output frequency via the coarse tune input f_c of the Front-End (FE), to ensure that the complete TX band falls within the ISM band. The tuning is achieved by changing the voltage applied to a varactor bank dedicated to coarse calibration. Making use of a pilot tone generated by the transmitter node and transmitted to the RG, an algorithm in the RG performs the fine calibration listed in Fig. 4.1. The algorithm can recover an offset of up to 8.2 MHz, with a precision of 7 kHz (<8 ppm at 915 MHz), in less than 300 μ s [77]. This corresponds to less than one bit overhead at 1 kbps and three bits overhead at 10 kbps. This algorithm is described in detail in Sect. 4.3.1.

The FLL and the coarse calibration DAC are realized on a PCB and make use of an ultra low-power micro-controller for their operation. The DAC consumes only 100 μ A during operation. At the end of the coarse center frequency calibration process the voltage applied to the coarse varactor must be preserved and kept constant during the whole synchronization process and data transmission. This result can be obtained in two ways:

- The coarse calibration word is digitally stored in a Random Access Memory (RAM)

¹It should be noticed that the extra hardware required for the FLL based coarse calibration respect to the EEPROM based factory calibration is the FD and the digital dividers (not shown in Fig. 4.2).

- The coarse calibration voltage is trapped on the capacitance of the coarse varactor

The first option, though feasible, requires the coarse calibration DAC to be on also after calibration. To save power, therefore, the second option has been used. Indeed, considering the large capacitance of the coarse varactor bank and keeping any leakage current low, it is possible to limit the variation of the voltage across the coarse varactor bank during the whole synchronization and data transmission process. Therefore, once this calibration is performed, the whole FLL is powered down, thus not contributing to the total power dissipation.

The reference signal used for the initial coarse center frequency alignment, is generated by the internal microprocessor oscillator and used to lock the transmitter center frequency to the 915 MHz band center frequency within 1% accuracy. This accuracy can be achieved either by factory alignment of the local oscillator or by using a factory calibrated on-chip oscillator of a standard low-power microprocessor. In this prototype the second solution was preferred.

Besides the center frequency calibration, another source of frequency error is present as described in Sect. 4.2. The details on how this error source is taken into account and on how correction is performed are described in Sect. 4.2.2. In this section we want to focus on a higher level concept and diagrams that describe how this error source is handled and what hardware is required.

A conceptual block diagram of the proposed architecture concept is depicted in Fig. 4.3. The incoming data, together with the desired hopping code, addresses a particular word cell in the ROM. The ROM has been split into two blocks depending on the data bit. In each memory cell the pre-distorted word is stored that drives the fine tuning DAC with a defined data bit. The DAC then directly drives the varactor array, changing the capacitance and therefore, the VCO oscillation frequency to the desired frequency bin around the band center frequency.

The flowchart in Fig. 4.4 describes how the architecture concept handles the frequency errors in the frequency bin synthesis. As previously described, the center frequency calibration process (divided in coarse and fine) is performed before data transmission. The coarse calibration is performed at the TX side while the fine is

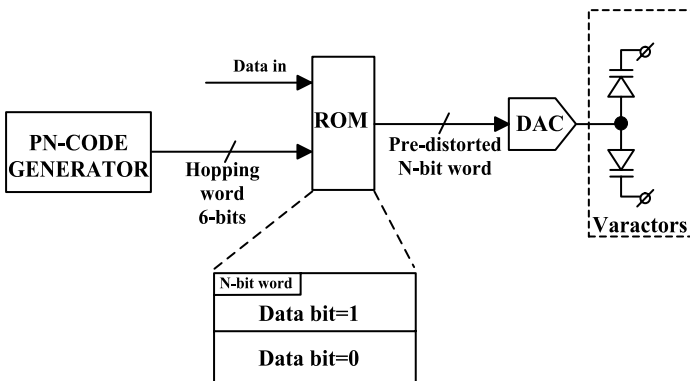


Fig. 4.3 Frequency pre-distortion conceptual block diagram

Fig. 4.4 Center frequency alignment and frequency bin synthesis flowchart

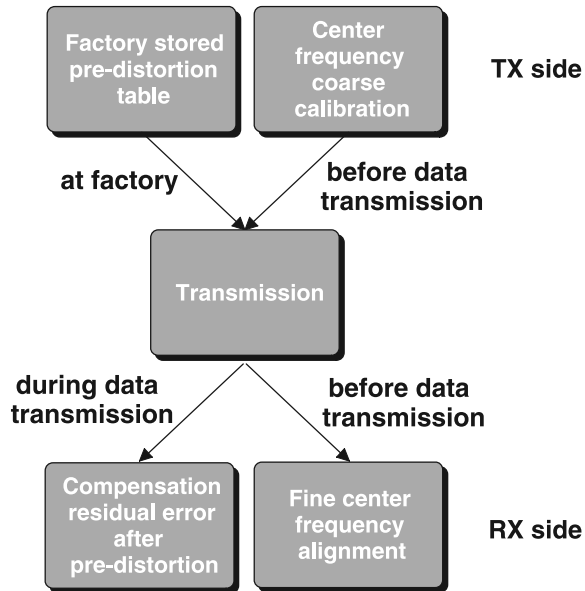
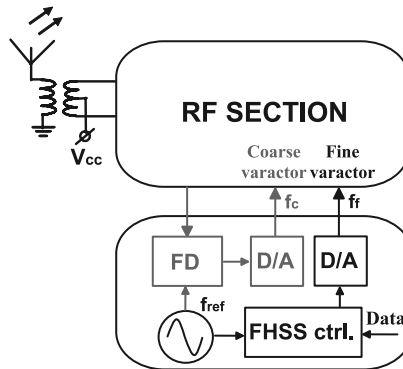


Fig. 4.5 FHSS synthesizer block diagram: the frequency pre-distortion path is included



performed at the RX side. The accurate synthesis of the frequency bins around the calibrated center frequency is also split between the TX and the RX. A coarse synthesis is performed at the TX side. This coarse synthesis is performed at factory level storing the mean values of pre-distorted words into a ROM as shown in the conceptual block diagram of Fig. 4.3. During data transmission a dedicated algorithm implemented at the RX side corrects for the residual frequency errors caused by PVT variations around the pre-distorted mean values stored at the TX side. This algorithm and its implementation are described in Sect. 4.3.2, while the pre-distortion algorithm used to derive the pre-distorted words stored in the TX ROM is described in Sect. 4.2.2.

The overall FHSS synthesizer block diagram is shown in Fig. 4.5. The grey loop, also shown in Fig. 4.2, is the initial coarse center frequency calibration loop. The

black path implements the frequency pre-distortion concept for the frequency bin synthesis. The FHSS controller contains the ROM to store the pre-distorted words and eventually the RAM to store the calibrated center frequency word.² A fine DAC is used to convert the pre-distorted digital words into an analog voltage that is applied to the fine varactor bank. This will change the overall capacitance allowing to change the synthesized frequency. This operation is described more into detail in the following sections.

4.2.2 Frequency Planning and Pre-distortion

The frequency of an LC type oscillator and the tank capacitance are related by the following well-known relation

$$f_{\text{osc}} = \frac{1}{2\pi\sqrt{LC}} \quad (4.1)$$

where L and C are the total capacitance and inductance of the tank respectively, and f_{osc} is the oscillation frequency. This, in practice, is realized by using a varactor diode, which has a capacitance that varies non-linearly with its reverse voltage. Therefore, by applying the correct voltages to the varactor diodes, it is possible to synthesize all the required frequency bins with minimum hardware complexity (virtually only a VCO).

In an FHSS system the various frequency bins are addressed in a pseudo-random fashion. Pseudo-random codes are generated in the digital domain, while the varactor diodes require an analog control voltage. Consequently, a DAC is required as interface between the digital world and the analog world. While passing from the digital world to the analog world several non-idealities can affect the precision with which the frequency bins are generated. The main sources of error in this conversion process are the following:

- DAC quantization error (deterministic)
- Varactor non-linearity (deterministic and stochastic)
- DAC INL (stochastic)
- Square-root relation between frequency and tank capacitance (deterministic)

All these non-idealities in the transmitter chain can be divided into two groups. One group has a deterministic behavior and it does not vary due to process spread. The square-root non-linear relation and the DAC quantization error fall into this category. The other group has a stochastic behavior. This means that it will depend on the process spread and therefore, a different behavior can be expected for different ICs.

²In the implemented case only a ROM is used because the calibrated coarse varactor voltage is stored on the coarse varactor capacitance and the coarse DAC is switched off after coarse calibration is performed.

Deterministic Errors

The non-linear relation between frequency and tank capacitance can be easily corrected by mapping the required frequencies to required capacitance values. From these values, knowing the varactor characteristic, a set of required voltages can be mapped and knowing the DAC specifications a set of digital words, which can be stored in a ROM, can be derived.

Therefore, neglecting for the moment the stochastic nature of the DAC linearity and of the C - V characteristic of the varactor, the problem can be simplified to the correct choice of DAC resolution to reduce the residual frequency error below a certain threshold.

The quantization error will produce a non-linear frequency error passing through the non-linear C - V characteristic of the varactor and through the square-root relation between frequency and tank capacitance.

Looking at the square-root relation, when the overall capacitance is the smallest (at higher frequencies), an error ($\Delta C_{\sqrt{LC}}$) on the capacitance due to quantization, will produce the largest error in the synthesized frequency. In this situation, the reverse voltage applied to the varactor is at its maximum value (for example -1.6 V). In this region, the varactor exhibits a highly linear behavior, contributing less to the overall frequency error. In these conditions, the frequency error is mainly caused by the quantization error passing through the square-root relation.

Close to the minimum reverse voltage (for example -0.2 V), the varactor characteristic is highly non-linear. In this case, the frequency error (ΔC_{var}) is mainly caused by the quantization error passing through the varactor non-linearity, while the square-root relation will minimally contribute to it.

In both cases a maximum capacitance error can be defined above which a frequency error larger than the required threshold is present at least in one of the synthesized frequency bins. These maximum errors are ($\Delta C_{\sqrt{LC}}$ and ΔC_{var}) defined as follows in the two cases

$$\Delta C_{\sqrt{LC}} = \frac{1}{[(2\pi(f_{\max} + \Delta f))^2 L]} - C_{\min} \quad (4.2)$$

$$\Delta C_{\text{var}} = \frac{1}{[2\pi(f_{\min} - \Delta f)]^2 L} - C_{\max} \quad (4.3)$$

where $f_{\max}$ is the highest channel center frequency, $f_{\min}$ is the lowest channel center frequency, Δf is the maximum allowed residual frequency error, $C_{\min}$ is the capacitance at the highest channel frequency, $C_{\max}$ is the capacitance at the lowest channel frequency and L is the LC-tank inductance value.

The maximum acceptable frequency error after pre-distortion can be derived on the basis of the following considerations. PNCs orthogonality has to be preserved. This means that the relative position of the channels along the frequency grid has to remain unchanged with respect to the ideal case. From this consideration, a maximum frequency error equal to the inter-channel spacing is allowed (for example 100 kHz). If a 100 kHz maximum channel frequency shift is considered, then there

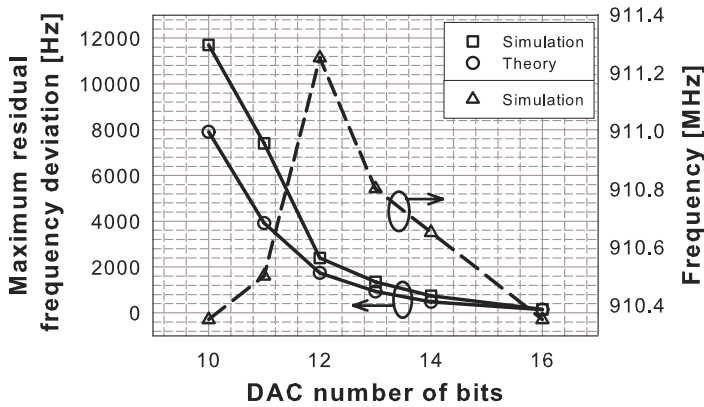


Fig. 4.6 Maximum uncorrected frequency error and its position in the frequency band versus DAC number of bits

would be the possibility that two channels become adjacent to each other (the inter-channel spacing becomes zero). In this situation, if the specification on the oscillator phase noise remains unchanged, the amount of noise leaking in the adjacent channels increases degrading the SNR in those channels.

In order not to degrade the BER in the adjacent channels considerably, a 0.5 dB maximum degradation on the phase noise was considered.³ Under this condition, a maximum uncorrected frequency error of 25 kHz can be tolerated [31].

Considering a 1.4 V swing on the varactor control voltage, an inductance value of 4.1 nH, and 64 channels placed around 915 MHz (ISM band), it can be found from (4.2), (4.3) that the largest error comes from the quantization error passing through the varactor non linear $C-V$ characteristic. This is shown in Fig. 4.6. In this picture three curves are presented. The two solid lines curves represent the calculated (via (4.2) and (4.3)) and simulated maximum residual frequency error after pre-distortion is applied versus the DAC resolution (DAC INL equal to zero). The dotted line curve represents the position in the frequency range at which the aforementioned frequency error occurs. As can be seen the largest frequency error due to the quantization error occurs at the lower portion of the frequency range.⁴

Given the previous considerations and looking at Fig. 4.6, in which all the non-linear sources are considered, it can be concluded that the frequency error coming from the DAC quantization error is mainly caused by the non-linear mapping of this error to the frequency domain via the $C-V$ characteristic of the varactor.

³The error probability of a non-coherent BFSK modulated signal is given by $\frac{1}{2}e^{-\frac{E_b}{2N_0}}$. Considering a 0.1% initial BER, a 0.5 dB degradation in the phase noise translates into a 0.5 dB degradation in the SNR at the demodulator input and therefore, into a BER close to 0.2%.

⁴Given 64 channels and a 150 kHz separation between adjacent channels, the minimum and maximum channel frequencies are 910.35 MHz and 919.65 MHz, respectively.

Given that the maximum residual frequency error has to be lower than 25 kHz, it can be concluded that a 10-bit DAC is sufficient to achieve the required specification.

Stochastic Errors

In the previous analysis the DAC was considered linear and the spread on the varactor capacitance was neglected. In the real case non-linearity and varactor spread will affect the overall residual frequency error.

Given a certain DAC, its non-linear behavior can be taken into account by applying a dedicated pre-distortion table. Unfortunately, the INL of each DAC will be different due to its statistical behavior. Therefore, this would require a different programmable look-up table per chip, which can be costly for a system that aims to be very cheap. As a result, it is necessary to fulfill the required specifications on the residual frequency error even when the DAC properties do change.

A DAC model was built in Simulink that also includes its non-linear behavior [31]. The results are shown in Fig. 4.7. Here both the quantization and the INL of the DAC are considered. On the x axis the maximum allowed DAC INL is shown in millivolt. For a 10 bit DAC, a 1 mV INL corresponds to roughly 1 LSB.⁵ On the y axis the maximum residual frequency deviation is shown. This error has to be, according to the derived specifications, below 25 kHz.

Roughly four regions can be recognized. The upper region (*Out of spec*) does not fulfill the maximum residual frequency offset requirement. Then, there are two regions namely *Difficult* and *Not worth*. The first region presents a difficult task for the designer due to harsh requirements in terms of maximum INL. The second one, while relaxing the INL requirements, necessitates the design of a higher resolution

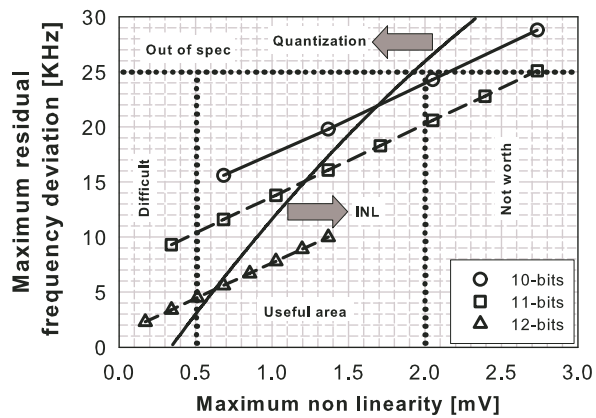


Fig. 4.7 Effect of DAC INL on the residual frequency error

⁵For an 11 bit DAC 1 mV corresponds to 2 LSB and for a 12 bit DAC it corresponds to 4 LSB.

DAC than in the *Useful area* to fulfill the maximum residual frequency error specification. Designing one more bit of resolution generally requires more area and more power. So the design of the DAC is near optimum inside the *Useful area* in Fig. 4.7.

Indeed, in this region the INL requirements are not too harsh and the residual frequency error due to the combined effect of both quantization error and DAC INL is smaller than the 25 kHz specification. This region is divided into two by a bold line. This line represents the points at which the contribution of the quantization error and of the INL to the residual maximum frequency error are equal. Therefore, the right part of this area is dominated by the INL error while the left part is dominated by the quantization error. Reducing the number of bits will reduce the chip area and in the end the costs and the power consumption of the DAC. As a result, given the low frequency operation of the DAC, a lower resolution DAC, which does not require an extremely small INL can be chosen.

Among different DAC specifications, which fulfill the maximum residual frequency error requirement, given the previous considerations and looking at Fig. 4.7, a 10-bit DAC with an INL between 1 and 1.5 mV can be chosen as a near optimum solution.

The last source of error is the variation in the C - V characteristic of the varactor due to process spread. As regards the DAC INL, this problem can be corrected by a dedicated pre-distortion look-up table. Though this is an effective solution it can be costly.

At this point we focus on the RG calibration of the residual frequency offset. This residual frequency offset is the sum of two contributions:

- DAC INL and quantization
- Varactor spread

The first source has been described, while the second source of residual frequency error will be described in the following part of this section.

Figure 4.8 shows the fine frequency characteristic of 20 IC samples, each calibrated to a common center frequency (915 MHz) via coarse tuning (namely f_c in Fig. 4.2). The channel bin tuning voltages are generated by the DAC, driven by the base-band microprocessor. The look-up table in the ROM was not changed and also the DAC is the same, therefore, the final effect is only due to the process spread on the varactor. Although the maximum frequency deviation of 20 ICs within the same batch were found to be no more than 220 kHz ($\sigma \approx 32$ kHz), inter-batch spreads would be larger.

In Fig. 4.9(a) the two extreme cases in Fig. 4.8 are plotted. The effect of the varactor spread (in the two aforementioned extreme cases) on the position in the band of the frequency bins for a given pre-distortion table is illustrated in Fig. 4.9(b). It can be seen that due to the difference in the varactor C - V characteristic, there is a frequency offset accumulating while moving from channel 1 to channel 32 or from channel -1 to channel -32 . The maximum error occurs in channel $+32$ or -32 , as can also be seen from the measured curves in Fig. 4.8. Given the monotonicity of the C - V varactor characteristic, the amount of frequency error between two adjacent channels due to the different C - V characteristics is negligible compared with

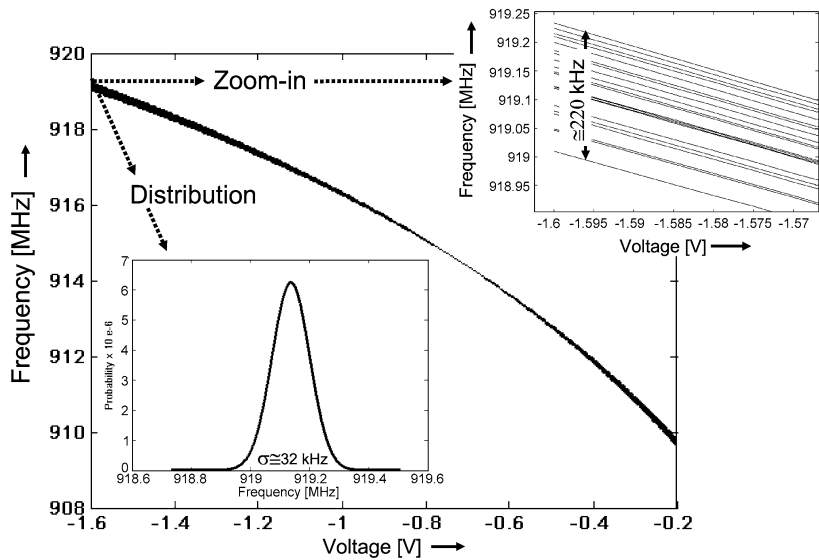
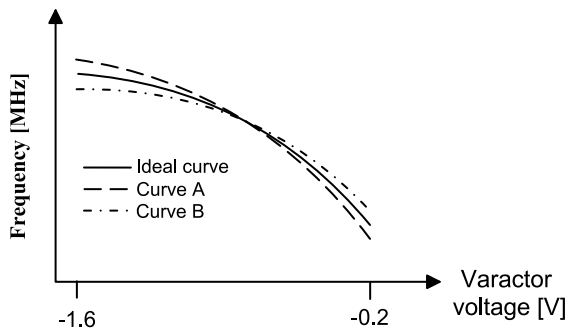
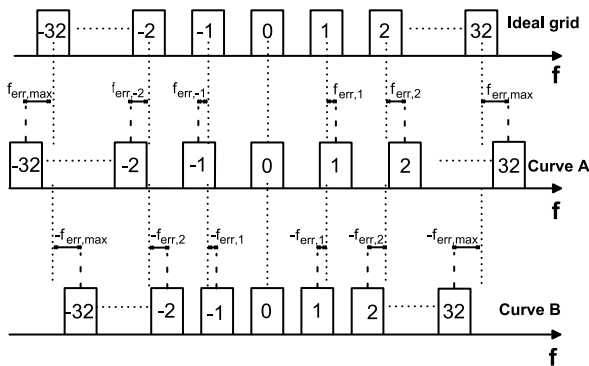


Fig. 4.8 Measured fine tuning characteristic for 20 IC samples (same DAC)

Fig. 4.9 Effect of the C - V varactor characteristic spread on the frequency synthesis



(a) The two extreme cases in Fig. 4.8



(b) Effect on the channel position in the frequency band

Fig. 4.10 Minimum inter-channel spacing when varactor spread is considered

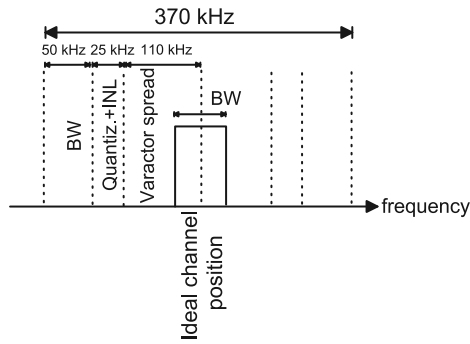
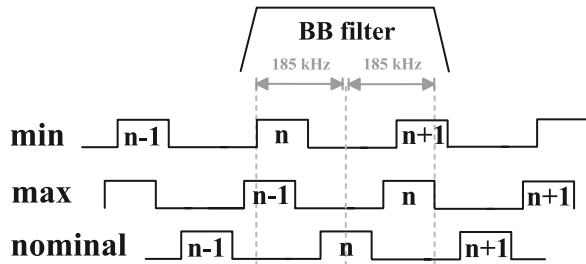


Fig. 4.11 Effect of residual frequency errors on the choice of the RG baseband filter bandwidth



the frequency error due to the quantization error. Unfortunately, assuming still that the content of the pre-distortion table is kept the same for all the ICs in a batch, this phenomenon will pose a problem on the receiver side. Due to its statistical dependence on the process, the absolute position of the last channel with respect to the ideal position will be known with a precision of ± 110 kHz (see Fig. 4.8). All the other channel positions will be known with a precision better than that.

At the receiver side the inter-channel spacing has to be larger than the sum of all spreads ($50 + 25 + 110 = 185$ kHz) and the channel bandwidth at the receiver side should be twice that value (2×185 kHz = 370 kHz) to cover the whole range where the transmitted band can be under spread conditions. This is clearly shown in Figs. 4.10 and 4.11. Indeed, the absolute position of the channel can spread 110 kHz (at 4 sigma) in both directions due to the varactor spread. Furthermore, the quantization error plus the INL of the DAC will add 25 kHz more uncertainty, in the worst case, in the channel position. Finally, the two channels should not overlap in the worst case, and therefore, the center frequency of the adjacent channel has to be at least the bandwidth apart (in this case 50 kHz). Therefore, the distance between two adjacent bins should be 185 kHz and the baseband filter of the RG has to span two times 185 kHz, which makes 370 kHz.

If all the 64 channels are utilized, the occupied bandwidth will be around 12 MHz. Given that no crystal has to be used, this will pose some stricter requirements on the reference frequency accuracy. Indeed, the accuracy has to become better than 0.75% in this case, while in the implemented case it can be relaxed to 1%. On the other hand, FCC rules demand only a minimum of 25 hopping channels for

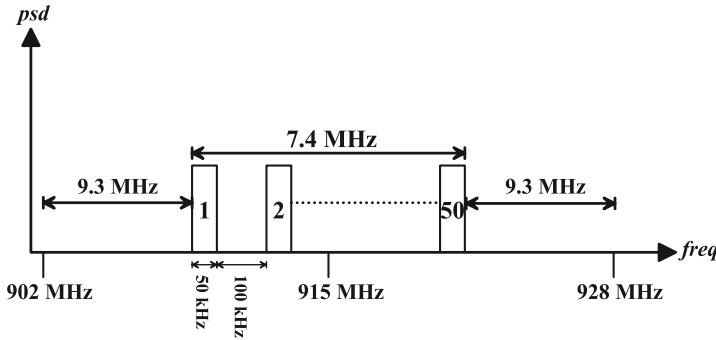


Fig. 4.12 Proposed channel allocation scheme for the 902–928 MHz ISM band

power levels below 0.25 W (which is generally the case for ultra low-power wireless nodes). Therefore, another possibility is to still keep the 1% accuracy for the reference frequency but to reduce the number of hopping channels.

The implemented prototype has been designed in order to use a 1% accuracy reference signal. Therefore, the inter-channel spacing has been slightly reduced to 150 kHz and while the synthesizer is able to generate 64 different channels, only 50 have been used. The implemented frequency allocation is depicted in Fig. 4.12.

4.2.3 Transmitter Specifications

This section focuses on the transmitter part of the wireless link. The transmitter part in a one-way link scenario is the power constrained portion of the wireless link. The receiving part is the RG, which is mains supplied. Therefore, the available power budget for the receiving part can be virtually considered infinite. Given these considerations, the transmitter architecture is conceived according to the following guidelines:

- The transmitter must fulfill the FCC rules before any data transmission is performed
- A minimum amount of processing must be carried out from the transmitter
- Complexity must be pushed to the receiver part of the link

Some important requirements must be satisfied in order for the receiver to reliably receive the incoming data. Important parameters are the following:

- Transmitted power
- Synthesizer phase noise
- Synthesizer coarse and fine tuning ranges

The transmitted power relates to the received power via the losses due to propagation and to the required SNR at the demodulator via the receiver noise figure NF . The

synthesizer phase noise is of primary importance in order to avoid the corruption of the incoming data. Indeed, supposing that the receiver oscillator can be made very low noise, the phase noise leaking into adjacent channels from an unwanted transmitter can degrade the SNR of the wanted channel. The coarse tuning range must be specified in order to obey the FCC rules while the fine tuning range depends on the number of hopping channels as well as from the inter-channel spacing. The next part of this section will address these requirements more into detail. The design addresses the 902–928 MHz band but all the considerations can be extended to the 2400–2483.5 MHz band.

Transmitted Power

In Sect. 1.4.3 it has been shown that the propagation losses in the 902–928 MHz band, considering 10 meters distance between the transmitter node and the receiver node, equal about 62 dB. The receiver sensitivity is a key point in calculating the required transmitted power. A sensitivity of -90 dBm is not difficult to achieve for a receiver with ideally unlimited power budget. From (1.5), considering -90 dBm receiver sensitivity and taking a 3 dB extra margin, the required transmitted power must be larger than -25 dBm.

This derivation does not take into account the possibility that fading occurs due to reflection paths. Fading conditions are quite prone to occur in an indoor environment and can severely affect the reliability of any kind of wireless link. As an example, we can calculate the required increase in the transmitted power when fading is considered, a BFSK modulation is used, and the raw BER⁶ is set to 1%.

If a non-coherent FSK demodulation scheme is applied at the receiver side, then it is well known from communication theory that in a AWGN environment:

$$P_{e,\text{NCFSK}} = \frac{1}{2} e^{(-\frac{E_b}{2N_0})} \quad (4.4)$$

where $P_{e,\text{NCFSK}}$ is the error probability. Using (4.4) we obtain an $\frac{E_b}{N_0}$ equal to 9 dB.

When fading is considered, it can be proven [10] that the probability of error for a non-coherent BFSK becomes

$$P_{e,\text{NCFSK}} = \frac{1}{2 + \Gamma} \quad (4.5)$$

where

$$\Gamma = \frac{E_b}{N_0} \alpha^2 \quad (4.6)$$

⁶With raw BER, in this book, we mean the BER when no error correction code is applied.

where α is the gain of the channel with Rayleigh distribution. The term Γ represents the average value of the normalized SNR. Therefore, to have a raw BER of 1%, a $\Gamma = 20$ dB is required. This translates in more than 10 dB increase in the transmitted power in order to cope with the fading conditions. From a power point of view this is not efficient. The solution adopted in this book shifts the problem from a mere brute force approach (increasing the power) to a more cumbersome approach, which consists of the following three principles:

- Fading is frequency dependent and can be counteracted by a hopping scheme
- Hopping rate can be traded for transmitted power
- Forward Error Correction (FEC) can be adopted in order to improve the BER

Using frequency hopping techniques assures that changing frequency helps to “avoid” deep faded channels. When the indoor environment is very hostile the same bit can be sent on different channels in order to increase the probability of correct reception if some of those channels present heavy fading conditions. Finally, Reed-Solomon correction code can be applied in order to improve the BER.

When Reed-Solomon error correction code is applied starting from a 10^{-2} raw BER, it can be proven [78] that a 10^{-6} BER can be achieved both in a fading channel as well as in the presence of a partial band jamming. Looking at a target application like a temperature sensor, one transmission of a few hundred bits packet every ten minutes gives a good Quality of Service (QoS). Now, supposing that the errors during a transmission are uniformly distributed, the packet error rate can be calculated to be 0.03% for a FEC BER of 10^{-6} and a 300 bits packet length.⁷ This translates to an error about every 20 days, which is good for this application especially if an acknowledge signal is also used to further reduce the PER.

Synthesizer Phase Noise

The synthesizer phase noise acts as a noise source. Let us suppose that an unwanted transmitter is transmitting on channel N and it has a noisy synthesizer. Let us suppose that the wanted transmitter is transmitting on a channel different from channel N (for example channel $N + 1$). Finally, we assume that the local oscillator on the receiver side is ideal (noiseless).

The noise leaking in channel $N + 1$ due to a communication on channel N can cause a degradation of the SNR for the wanted signal if it is not sufficiently low. Sufficiently low means that its level has to be at least the noise figure plus the required SNR at the demodulator input below the signal on the wanted channel. Furthermore, the wanted and the unwanted signals can be attenuated differently due to different propagation losses and so will be the noise leaking into the wanted channel. This means that an extra protection factor is required in order to consider the

⁷Supposing the errors in a packet are uniformly distributed is a worst case scenario especially for burst based communications. The formula used to derive the PER starting from the channel BER is, in this simplified case, $PER = 1 - (1 - BER)^P$ where P is the packet length.

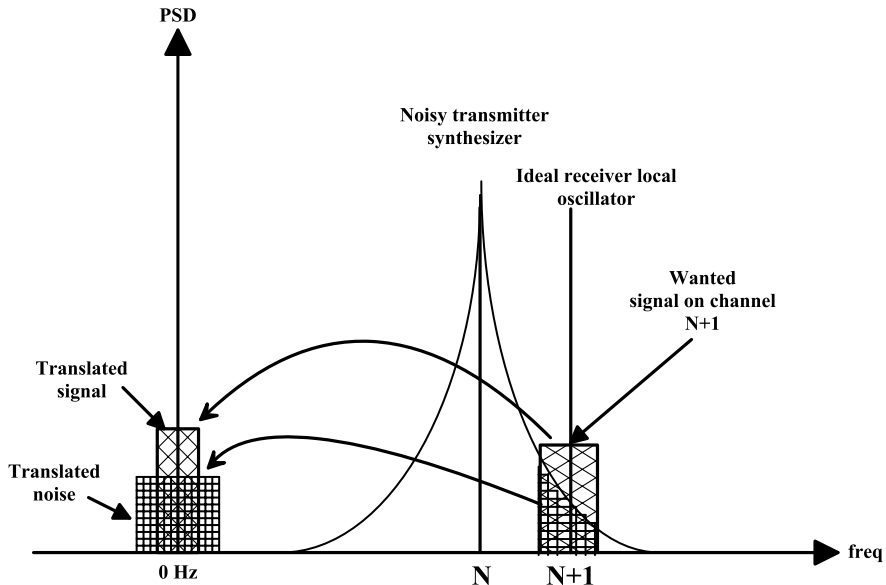


Fig. 4.13 Effect of phase noise produced by an unwanted communication on the reception of the wanted signal

different propagation losses for the two signals. The described example is depicted in Fig. 4.13 when a zero-IF receiver is considered.

In [79] it is shown that for the narrowband case, the phase noise slope can be neglected across the bandwidth. Therefore, in the wanted bandwidth a constant value of the PSD of the LO will be used (which can be reasonably chosen as the value in the middle of the band). Furthermore, to simplify the calculations, we will assume an ideal receiver and we will consider a 3 dB margin to account for implementation degradation. Concluding, the phase noise requirements for the LO in decibel units can be derived by the following equation:

$$L_{\Delta f} \leq P_{WS} - P_{US} - BW - SNR - NF - 3 \text{ dB} \quad (4.7)$$

where $L_{\Delta f}$ is the phase noise expressed in dBc/Hz at a certain frequency offset from the carrier, P_{WS} is the power of the wanted signal, P_{US} is the power of the unwanted signal, BW is the noise bandwidth and the 3 dB is the implementation margin. The term $P_{WS} - P_{US}$ is the protection factor, which takes into account the different propagation losses for the wanted and the unwanted signal.

Now a reasonable level for the wanted and unwanted signal should be chosen. Looking at Bluetooth specifications, the power level difference between the unwanted and the wanted signals has been chosen to be as shown in Table 4.1. In the Bluetooth standard, system simulations have proven that these protection ratios, result in a very little degradation in the throughput. From these specifications, it is

Table 4.1 Power difference between wanted and unwanted signals for Bluetooth

Channel	$P_{\text{unwanted}} - P_{\text{wanted}}$
Adjacent	0 dB
Alternate	30 dB
Third and beyond	40 dB

reasonable to suppose that, when the unwanted signal is in the adjacent channel, the wanted and unwanted signals can have the same power (this means that they can be at the same distance). The required signal bandwidth can be derived via the Carson's rule:

$$BW = 2\Delta f + f_m \quad (4.8)$$

where $f_m = \frac{2}{T_s}$ with $\frac{1}{T_s}$ the data rate and Δf the frequency deviation. For a data rate up to 10 kbps and $m = 10$, a 120 kHz noise bandwidth needs to be considered. In this case the allowed phase-noise from (4.7), given $NF = 10$ dB, is $L_{\Delta f} = -73$ dBc/Hz.

In the alternate channel, we suppose that the unwanted signal can have a power of 20 dB higher than the wanted one. From (4.4) and (4.5), given $n \in [2, 3]$ (see Sect. 1.4.3), this corresponds to a relative distance between wanted and unwanted transmitters between 5 and 10. In this case, $L_{\Delta f} = -93$ dBc/Hz.

Finally, in the third and beyond channels, it is reasonable to set the power of the unwanted signal 30 dB higher than the wanted one. This corresponds to a relative distance between wanted and unwanted transmitters between 10 and 30 and translates in $L_{\Delta f} = -103$ dBc/Hz.

The most demanding phase noise requirement is, therefore, at the third and beyond channels. The phase noise requirement will be at 3 MHz far from the carrier around -119 dBc/Hz, that is less demanding than what is found in typical Bluetooth systems. If resonators with a quality factor Q of 10 are available, then an ultra-low power VCO can be designed. Of course, it is possible to trade NF for phase noise (see (4.7)) to relax as much as possible the transmitter specifications. On the other hand, the lower the receiver NF the more stringent will be the requirements for the RG. This can raise the cost of the RG above a reasonable level.

Synthesizer Coarse and Fine Tuning Ranges

As previously mentioned in this section, an initial calibration is required in order to allow the transmitter to fulfill the FCC rules. Given the fact that the absolute values of capacitance and inductance in any oscillator spread due to process, the absolute value of the synthesized frequency spreads as well. This spread can be large enough to cause some of the channels to be, at transmitter power-up, outside the allocated frequency band (902–928 MHz in this particular case). This is not allowed by FCC rules and therefore, it must be corrected. An FLL loop is used in this particular example as discussed previously in Sect. 4.2.1.

Table 4.2 Transmitter specifications

	Value	Unit
Output power	> -25	dBm
Phase noise	< -103	dBc/Hz @ 450 kHz
Coarse tuning range	> 46	MHz
Fine tuning range	> 7.5	MHz

To derive the required coarse tuning range it is necessary to know the absolute spread of capacitances and inductances of the process. The ultra low-power transmitter prototype is realized in a bipolar process, which is produced on SOI wafers followed by a substrate transfer to glass [80]. In this process the inductances spread much less than the capacitances and therefore, we will assume the inductance value ideal and only the capacitance spread will be considered. The capacitances can spread as much as 10% of the nominal value. We will assume that the oscillator is based on a general LC tank topology. Therefore, the basic oscillation equation is (4.1).

Starting from (4.1) and assuming the inductance ideal and ΔC spread on the capacitance, the synthesized frequency is:

$$f_0 + \Delta f = f_0 \left(\frac{1}{\sqrt{1 - \frac{\Delta C}{C}}} \right) \quad (4.9)$$

where f_0 is the oscillator ideal center frequency and Δf is the frequency deviation due to the process spread on the capacitance. If we expand the term between brackets in (4.9) in McLaurin series, then the frequency spread due to ΔC spread on the capacitance can be approximated by the following equation:

$$\Delta f = \frac{f_0}{2} \frac{\Delta C}{C} \quad (4.10)$$

Considering a center frequency of 915 MHz and 10% spread of the capacitance the required coarse tuning range is larger than 46 MHz.

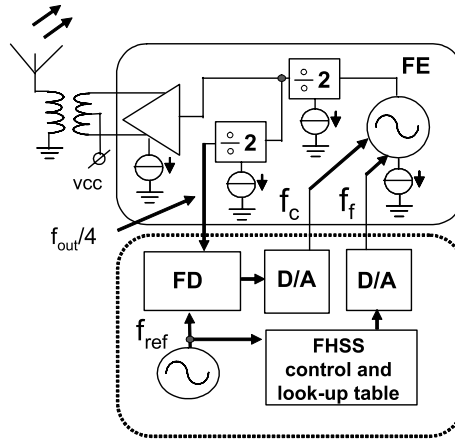
The fine tuning range can be easily derived. In the 902–928 MHz band 50 hopping channels are required. If the inter-channel spacing is assumed to be 150 kHz in this prototype, then the required fine tuning range is 7.5 MHz. The requirements for the transmitter are summarized in Table 4.2.

4.2.4 Oscillator-Divider Based Architecture

The schematic block diagram of this architecture is shown in Fig. 4.14.

The transmitter consists of a 1.8 GHz VCO, 2 dividers and 1 output stage. The FLL measures the $f_{\text{out}}/4$ signal from the VCO, and calibrates f_{out} via the coarse

Fig. 4.14 Proposed VCO-divider based frequency hopping transmitter architecture



tune input f_c , to ensure that the complete TX band falls within the ISM band. Once this calibration is performed, the FLL and related dividers are powered down, thus not contributing to the total power dissipation.

The fine tuning input of the VCO (f_f), is directly modulated by the base-band to allow a variable hopping frequency FHSS transmitter with 64 channels, employing wideband BFSK. The FHSS control logic block takes in the baseband digital data and performs digital FSK modulation and frequency hopping to generate, through a DAC, an analog input (f_f) for the VCO.

The system front-end consists of a resonator based LC-VCO, a divider and an output stage able to deliver -25 dBm power on a $50\ \Omega$ load. To minimize the oscillator pulling the VCO operates at 1.8 GHz and divided down to the TX frequency using a divide-by-two frequency divider.

Oscillator Design

The VCOs can be grouped into two categories:

- Non-resonant oscillator
 - Ring oscillator
 - Relaxation oscillator
- Resonant oscillator
 - LC tank oscillator
 - Crystal oscillator

Crystal oscillators cannot be integrated because the crystal cannot be integrated and they have a very poor tunability. Therefore, they will not be further discussed in this section. Summarizing, three possible topologies are available: two RC-type oscillators and one resonant oscillator.

Ring oscillators are not easily implementable in a bipolar technology. Indeed, while in a CMOS technology they have the advantage of low power consumption

at the expense of degraded phase noise at high operating frequencies, in a bipolar process this advantage disappears. The main reason is that any logic gate in a bipolar process consumes static power. This is not the case in CMOS technology where power is burned only during transitions. As an example of a bipolar ring oscillator the work in [81] is considered. In this work two ring oscillators have been designed. They have a different number of stages: two and four respectively. The oscillator frequency is 1.8 GHz, which roughly corresponds to the required oscillation frequency of the VCO used in the proposed architecture. The two stages oscillator outperforms the four stages in terms of current consumption and phase noise performances. At 7.5 mA current consumption it achieves -82 dBc/Hz at 100 kHz. Given the fact that the phase noise in the thermal noise region decreases by roughly 6 dB/octave at 450 kHz the phase noise is around -94 dBc/Hz. This value is still far from the required -103 dBc/Hz.

Let's now consider the Leeson equation, which relates the phase noise to the quality factor Q of the oscillator and to the power used in the oscillator:

$$L(f_m) = 10 \log_{10} \left(\frac{FkT}{P} \frac{1}{8Q_L^2} \left(\frac{f_0}{f_m} \right)^2 \right) \quad (4.11)$$

where F is the noise factor of the active device at the power level P , Q_L is the loaded quality factor, f_0 is the oscillator carrier frequency and f_m is the frequency offset from the carrier. It is clear that any doubling in the oscillator power brings only 3 dB better phase noise performance at a given offset frequency from the carrier. To reach the required -103 dBc/Hz at 450 kHz from the carrier 9 extra dB are required. The required power has to be, therefore, 8 times larger than what is measured in [81]. This example, clearly shows that a ring oscillator is not a reasonable choice at least when a bipolar process is considered.

In [82] it is shown that a relaxation oscillator, though in theory it can reach the same performance as a ring oscillator, in practice is further away from its thermodynamic limit than the ring oscillator topology. Though in the latest technologies, due to the channel excess noise, the difference between the practical phase noise and the theoretical thermodynamic limit in ring oscillator is increasing, they generally outperform the relaxation topology.

The last topology, the LC-tank based oscillator, is based on the principle of resonance. From (4.11), it can be seen that the phase noise improves with the square of the quality factor. This means that, for a fixed phase noise performance at a certain frequency offset from the carrier, the oscillator with the largest Q consumes the least amount of power. The drawback compared to the non-resonant oscillators is the silicon area. Indeed, LC-tank oscillators require an inductor, which can consume a large area when integrated on chip. Given the very high quality passives available in the SOA technology, and the fact that a differential inductor is available, the choice for an LC-tank base topology is the best trade-off in terms of power-area-performances.

Two different LC-tank based oscillator are generally used:

- Cross-coupled based LC oscillator
- Colpitts or Hartley based oscillator

In the oscillator-divider architecture described in this chapter, a cross-coupled based LC oscillator has been implemented. In Sect. 4.3.2 a power-VCO based architecture is disclosed, which uses a Colpitts based oscillator. In general, it is possible to state that a Colpitts oscillator outperforms, for a given bias current, the cross-coupled topology. Indeed, the signal amplitude at the oscillator output, in the two cases, is given by (4.12) and (4.13).

$$A_{\text{cross-coupled}} = \frac{2}{\pi} R_{\text{TANK}} I_{\text{bias}} \quad (4.12)$$

$$A_{\text{Colpitts}} = 2 R_{\text{TANK}} I_{\text{bias}} \quad (4.13)$$

where I_{bias} is the bias current and R_{TANK} is the tank resistance.

As it can be seen from these equations, for a given value of the tank resistance and for a given bias current, the signal amplitude is larger in a Colpitts oscillator compared to a cross-coupled oscillator. This means that the phase noise performance in the Colpitts oscillator is better than in a cross-coupled oscillator for a given bias current. In the same way, given a certain phase noise performance at a certain offset frequency from the carrier, the bias current of the Colpitts oscillator can be sized to be smaller than that in a cross-coupled oscillator.

The SOA technology has several inductances available. Some of them are also differential inductances with a center tap. These kinds of inductances are an optimal choice to reduce the area occupancy. As it will be shown later, the area occupied by an inductance exceeds by far the area occupied by all the active circuits in this architecture. In this technology, the substrate is removed and replaced for a high ohmic substrate such as glass. In this way, the parasitics of the active devices are reduced and at the same time it is possible to integrate passive elements with high Q s.

As an example, consider L values and Q factor curves simulated for some of the readily available coils, as shown in Figs. 4.15 and 4.16 respectively. The nomenclature of the coils is the following:

- D is the diameter of the coil in μm .
- N is the number of turns.
- W is the width of the turn(s) in μm .
- S is the spacing between the turns in μm .

Q factors peak between 4 GHz and 6 GHz and can reach values as large as 30 at 1.8 GHz. From (4.11) it is clear that the availability of high- Q passives is a strong argument in favor of a resonant based oscillator topology.

From the previous considerations, a differential coil is chosen, based on its high Q factor of 30 at 1.83 GHz. The diameter of this coil is 700 μm and it has an inductance value of 4.19 nH (a balanced coil with a higher Q factor of 38 at 1.83 GHz is available, but with a considerably larger diameter of 1000 μm). The Q factor of the tank is given by the parallel combination of the Q factor of the inductance and the Q factor of the capacitance. Given the fact that a certain tunability is required,

Fig. 4.15 Inductance of coils as a function of frequency

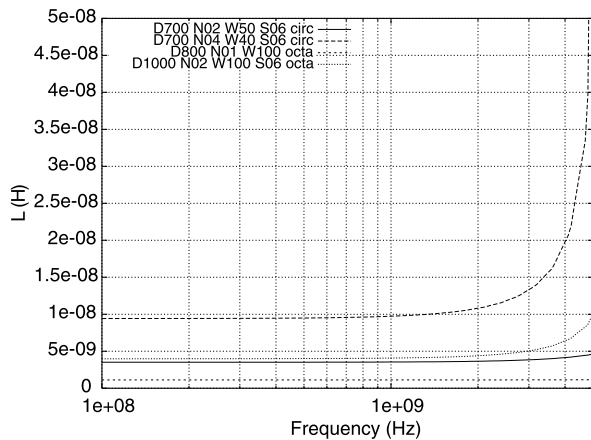
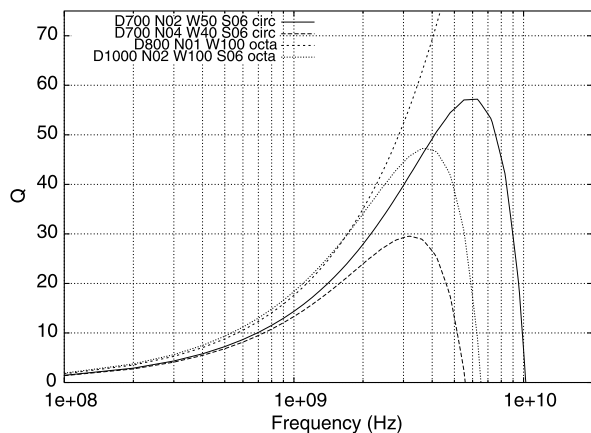


Fig. 4.16 Q factor of coils as a function of frequency



and considering a varactor as a voltage controlled capacitance, it is safe to assume, looking at Fig. 4.17, that the tank Q factor is dominated by the inductance and that the capacitance (a varactor in this example) has a virtual infinite Q compared to the inductance.

Two possible oscillator biasing techniques are analyzed in terms of area and power consumption. These topologies are depicted in Figs. 4.18(a) and (b). The AC coupled topology is generally preferred over the current mirror biased topology because it shows a better phase-noise performance and requires a lower supply voltage (only one transistor instead of two stacked transistors). Care must be taken so that the base-collector junction of the bipolar transistors does not become forward biased in order to avoid any performance degradation. If this happens, the real part of the input impedance is degraded decreasing the LC-tank Q (thus increasing the phase noise). Moreover, the transit time of the transistor will be increased. Therefore, the collector potential must always remain above the base voltage. One way to achieve this is by connecting the bases to the collectors via emitter followers. Un-

Fig. 4.17 Varactor Q factor versus forward bias voltage

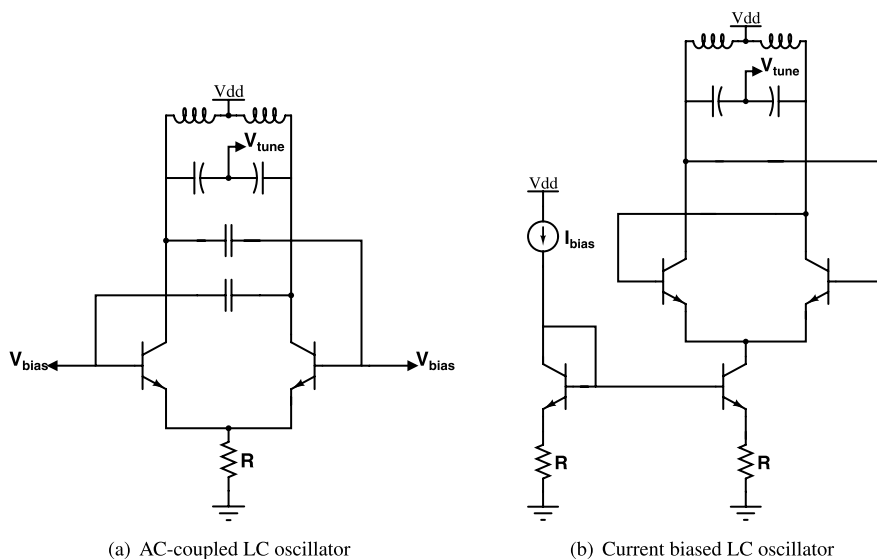
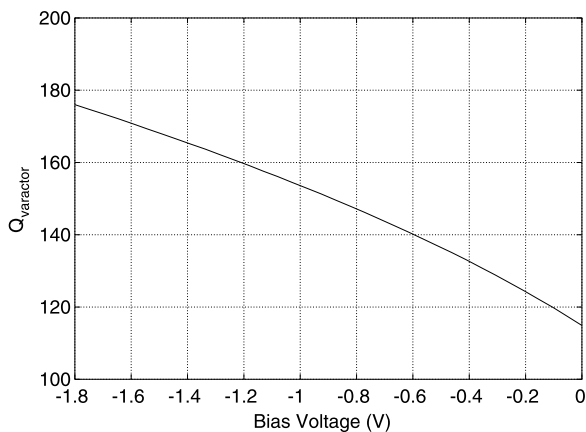


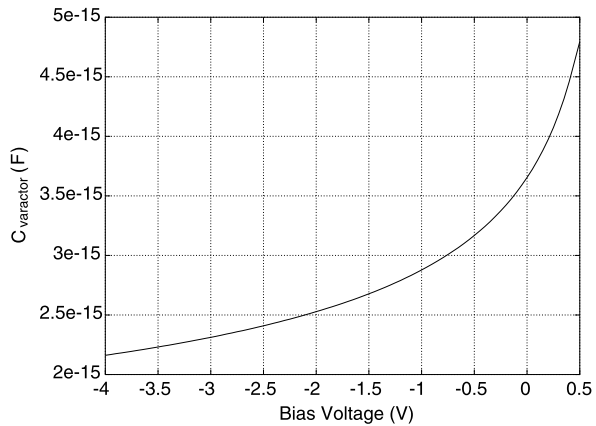
Fig. 4.18 Two possible biasing techniques for an LC cross-coupled oscillator

fortunately, this approach adds noise. A better method is by AC-coupling the device bases as shown in Fig. 4.18(a). A drawback of the AC-coupled topology is that it requires a biasing network for the LC cross-coupled differential pair. Both methods will be compared in power consumption and phase noise performance.

The oscillator needs to be tunable in two dimensions:

- Hopping bins synthesis
- Process spread compensation

In the SOA technology a unit varactor cell is provided. A supply voltage of 1.8 V has been chosen for this prototype. Considering a 0.2 V margin, the tuning range

Fig. 4.19 Simulated varactor capacitance curve

for the fine tuning varactor ranges between -0.2 V and -1.6 V. Figure 4.19 shows the simulation results of the available P - N varactor capacitance versus the forward bias voltage. The single cell varactor capacitance ranges from 2.644 fF to 3.422 fF.

The hopping bins synthesis requires a tuning range (fine tuning range) of approximately 7.5 MHz (see Table 4.2). Considering some margin, the system has been designed for a 10 MHz fine tuning range at 915 MHz. This translates in a 20 MHz fine tuning range at 1.83 GHz. The number of varactors $N_{\text{var},\text{fine}}$ needed in parallel to span the required 20 MHz fine tuning range, can be calculated from

$$N_{\text{var},\text{fine}} = \frac{C_{\text{max},\text{fine}} - C_{\text{min},\text{fine}}}{C_{\text{var},\text{max}} - C_{\text{var},\text{min}}} \quad (4.14)$$

Here, $C_{\text{max},\text{fine}}$ and $C_{\text{min},\text{fine}}$ are the respective maximum and minimum fine tuning capacitance values required and $C_{\text{var},\text{max}}$ and $C_{\text{var},\text{min}}$ are the maximum and minimum varactor capacitances as determined above. Considering the relation between oscillation frequency and the values of inductance and capacitance in the tank (4.14) can be written in the following form

$$N_{\text{var},\text{fine}} = \frac{\frac{1}{(2\pi f_{\text{min}})^2 L} - \frac{1}{(2\pi f_{\text{max}})^2 L}}{C_{\text{var},\text{max}} - C_{\text{var},\text{min}}} \quad (4.15)$$

with $f_{\text{max}} = 1.84$ GHz and $f_{\text{min}} = 1.82$ GHz the maximum and minimum frequency respectively, assuming the band is tuned symmetrically around the center frequency. Evaluating the expression for the provided values, results in

$$N_{\text{var},\text{fine}} \approx 51 \quad (4.16)$$

Fig. 4.20 Fine tuning curve of a single varactor cell in the SOA technology

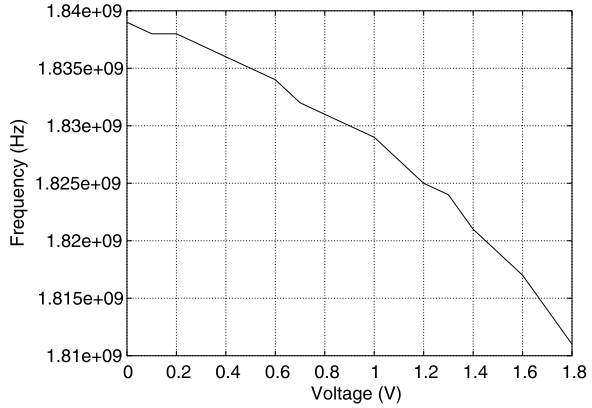


Figure 4.20 shows the fine tuning curve after simulating the VCO over the complete voltage range. Between 0.2 V and 1.6 V tuning voltage, the tuning range is 20 MHz, as was expected.⁸

The goal of the coarse tuning is to make sure that the complete frequency hopping spectrum, already calculated to be approximately 20 MHz, is completely inside the 1.804 – 1.856 GHz band. When the fine tuning varactor is biased at -0.9 V, meaning when it is tuned for the center frequency, the total fine tuning capacitance is $51 \cdot 2.927 \text{ fF} = 149.3 \text{ fF}$. Supposing the VCO is coarsely tuned around 1.83 GHz, then:

$$C_{\text{center,coarse}} = \frac{1}{(2\pi \cdot 1.83 \text{ GHz})^2 L} - 149.3 \text{ fF} \quad (4.17)$$

and equals 1.65 pF.

There are two possible ways to coarsely tune the VCO into the assigned working frequency band:

- Band switching
- Varactor tuning

Figure 4.21 shows a simplified model of an off-state and on-state typical band switch circuit implementation.

In the off state, the transistor switch acts as a parasitic capacitance C_Q in series with the switching capacitance C_{switch} . This parasitic capacitance should be small, as it reduces the effective tuning range, which can be expressed by

$$\frac{C_{\text{max}}}{C_{\text{min}}} = \frac{C_{\text{switch}}}{\frac{C_{\text{max}} \cdot C_Q}{C_{\text{max}} + C_Q}} \quad (4.18)$$

⁸Note that the tuning voltage equals $1.8 \text{ V} + V_{\text{var}}$, with V_{var} the voltage over the varactor, because, when the tuning voltage is 0 V, there is already 1.8 V at the cathode.

Fig. 4.21 Typical switched capacitor circuit implementation and a simplified model for off-state and on-state respectively

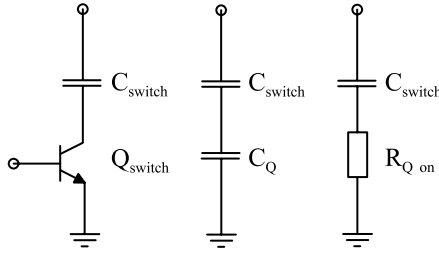


Table 4.3 Single band switch simulation performance

No. of transistors	C_Q (fF)	$\frac{C_{\max}}{C_{\min}}$	Q_b	Base current (μA)
1	3.57	17.8	2.9	3
5	17.9	4.3	14.5	15
10	35.7	2.68	29	30
20	71.4	1.84	58	60

In the on state, the transistor can be modeled as a resistor R_{Qon} , so the quality factor of the switch is of primary importance and is given by

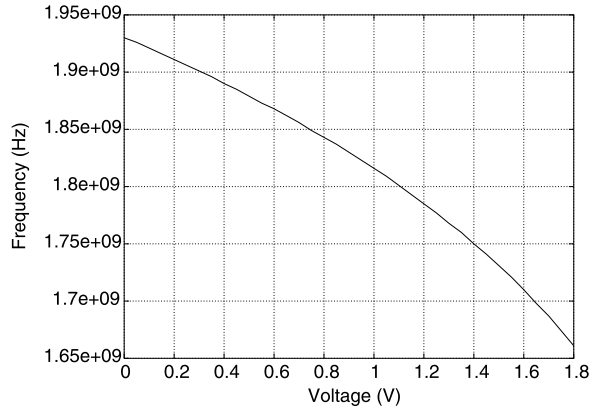
$$Q_b = \frac{1}{\omega C_{\text{switch}} R_{Qon}} \quad (4.19)$$

From this equation it is clear that R_{Qon} should be small. A way to reduce the ‘on’ resistance of the transistor (and thus increase Q_b), is by putting transistors in parallel. However, this not only leads to an increase in power consumption, but also to a reduction in the tuning range, because the parasitic capacitance increases.

The worst case occurs when the edge of the spectrum is just out of the band. When this happens, the frequency should be shifted by such an amount, that the spectrum remains in band, which means that it does not go over the other edge of the band. Therefore, the maximum shift should be $(1856 \text{ MHz} - 1804 \text{ MHz}) - 20 \text{ MHz} = 32 \text{ MHz}$. The incremental capacitance, when shifting from, for instance, 1.83 GHz to 1.862 GHz is approximately 60 fF. The required coarse tuning range has to be around 100 MHz at 1.83 GHz (see Table 4.2). Therefore, the coarse tuning range varies between 1.78 GHz and 1.88 GHz. Given a 4.19 nH inductance the two extreme capacitance values are 1.76 pF and 1.56 pF. To go from the minimum capacitance of 1.756 pF (the fixed capacitance) at 1.88 GHz to the maximum capacitance of 1.76 pF at 1.78 GHz, 4 switched capacitances of 60 fF are necessary, assuming the parasitic capacitance is very small. Table 4.3 shows simulation results for the switch, as displayed in Fig. 4.21, with a 60 fF switch capacitance.

Considering the high quality factor of the inductances, the Q factor of the switching branches should be a factor five or more higher in order to not degrade the VCO performance. From Table 4.3 it is clear that, to reach this result, several transistors need to be place in parallel. However, in this case, $\frac{C_{\max}}{C_{\min}}$ is very low and the base

Fig. 4.22 Varactor coarse tuning curve



current approaches several tens of μA , when all switches are on, causing a total base current of a several hundreds of μA .

Of course, it is also possible to achieve coarse tuning with varactors, in a similar fashion as the varactor fine tuning. When the voltage over the varactor is -0.9 V and the oscillator is in the center frequency, the coarse varactor capacitance should be 1.65 pF . For a single varactor cell, the capacitance at -0.9 V bias voltage is 2.927 fF . The number of parallel varactors $N_{\text{var,coarse}}$ can be calculated by

$$N_{\text{var,coarse}} = \frac{1.65\text{ pF}}{2.927\text{ fF}} \approx 565 \quad (4.20)$$

Figure 4.22 shows the coarse tuning curve after simulating the VCO over the complete voltage range. Between 0.2 V and 1.6 V tuning voltage, the tuning range is around 200 MHz . This is much larger than the required tuning range of around 100 MHz . Concluding, for both the coarse and fine tuning mechanism a varactor tuning element has been chosen as an optimal solution.

The biasing current for the LC VCO should be chosen so that three conditions are met:

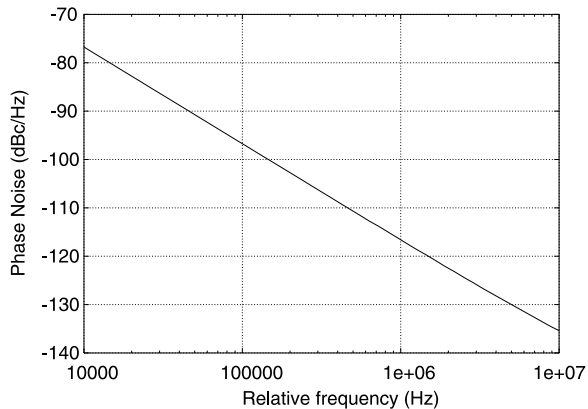
- Minimum loop gain for oscillation start-up
- Sufficient output signal swing
- Phase-noise requirement

The minimum loop gain can be easily evaluated setting the negative conductance of the cross-coupled pair at least twice the conductance given by the inductor series resistance. This leads to the following relation for the oscillation starting-up condition:

$$g_m \geq \frac{R_L C}{L} \quad (4.21)$$

where g_m is the transconductance of each transistor in the cross-coupled pair and R_L is the inductance equivalent series resistance. Given $R_L = 1.6\ \Omega$ at 1.83 GHz , each transistor in the cross-coupled pair must have a transconductance larger than 1.4 mS .

Fig. 4.23 Phase noise plot for the topology with AC coupled cross-coupled pair



This translates in a current larger than 35 μA for each transistor and therefore, an overall bias current larger than 70 μA .

The signal amplitude at the oscillator output is governed, to the first order, by the AC impedance of the lossy tank:

$$V_{\text{out}} = I_{\text{bias}} 2\pi Q_L f_0 L \quad (4.22)$$

where I_{bias} is the bias current, Q_L is the inductance quality factor, f_0 is the oscillation frequency and V_{out} is the peak-to-peak voltage at the oscillator output. The required current for a 300 mV peak to peak output is around 210 μA making this requirement dominant respect to the starting-up requirement.

For the topology from Fig. 4.18(a), using the AC cross-coupled pair, the tail current is approximately 240 μA . For the current mirror biasing topology from Fig. 4.18(b), the tail current is 290 μA . The current mirror, apart from requiring more components, requires more tail current than the AC coupled cross-coupled pair to acquire the same amplitude mainly because of higher losses in the biasing current source and because of the reduced available headroom.

Phase noise simulations have been performed on both the topologies shown in Fig. 4.18 at the previously derived bias currents. Figure 4.23 illustrates the simulation results of the topology with the AC coupled cross-coupled pair. Each of the feedback transistors has 10 devices in parallel. At 450 kHz, the phase noise is approximately -109.7 dBc/Hz. This value exceeds the required phase noise by more than 6 dB, giving enough margin for possible implementation losses.

Looking at the noise summary, presented in Table 4.4 for the topology with AC coupled cross-coupled pair, the noise contribution of the feedback transistors is rather high. A way to reduce this, is to place more transistors in parallel.

Figure 4.24 illustrates the simulation results when 20 transistors are placed in parallel. At 450 kHz, the phase noise is now approximately -111.2 dBc/Hz, which is an improvement of 1.5 dBc/Hz.

The noise summary after this modification is given in Table 4.5. Clearly, from this table, the noise contribution is now in the first place dominated by the losses

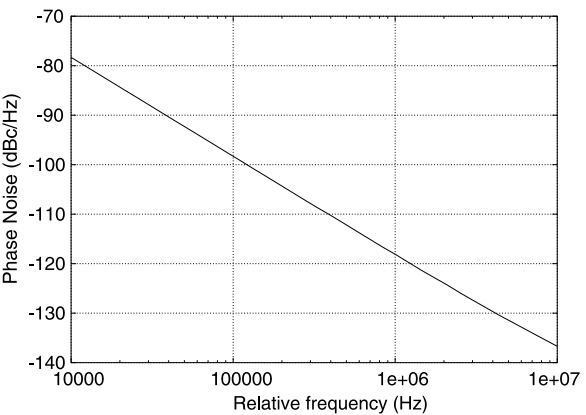
Table 4.4 Noise summary for the VCO topology with AC coupled cross-coupled pair, with noise contribution ranking (only the first 6 contributors are shown)

Device	Noise contribution
Variable base resistance feedback transistors	23.2%
Coil loss resistance	23.1%
Collector current through feedback transistors	22.9%
Constant base resistance feedback transistors	13.7%
Base current through feedback transistors	4.41%
Coarse varactor array loss resistance	4.13%

Table 4.5 Noise summary for the VCO topology with AC coupled cross-coupled pair and 20 feedback transistors in parallel, with noise contribution ranking (only the first 6 contributors are shown)

Device	Noise contribution
Coil loss resistance	29.2%
Collector current through feedback transistors	28.5%
Variable base resistance feedback transistors	14.7%
Constant base resistance feedback transistors	8.70%
Coarse varactor array loss resistance	5.13%
Base current through feedback transistors	4.25%

Fig. 4.24 Phase noise plot for the topology with AC coupled cross-coupled pair and 20 feedback transistors in parallel



in the coil, while the noise contribution of the base resistance becomes much less dominant.

Figure 4.25 illustrates the simulation results of the topology biased via a current mirror, with 10 devices in parallel for each feedback transistor. At 450 kHz, the phase noise is approximately -109 dBc/Hz, so slightly worse than the topology with the AC coupled cross-coupled pair. Contrary to Table 4.4, the noise summary, presented in Table 4.6 for the topology with current mirror, shows that the dominat-

Fig. 4.25 Phase noise plot of the topology biased via a current mirror

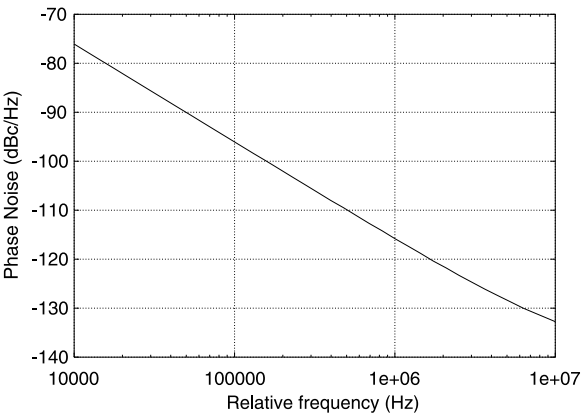


Table 4.6 Noise summary for the VCO topology with current mirror biasing, with noise contribution ranking (only the first 6 contributors are shown)

Device	Noise contribution
Coil loss resistance	20.5%
Collector current through feedback transistors	19.3%
Variable base resistance feedback transistors	16.1%
Collector current through current mirror transistors	10.1%
Constant base resistance feedback transistors	9.79%
Variable base resistance current mirror transistors	2.83%

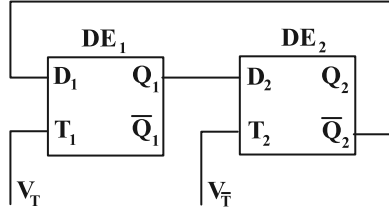
ing noise source are already the losses in the coil. Therefore, increasing the number of parallel transistors in the cross-coupled pair does not have any impact on the VCO phase-noise performance.

From all the previous considerations, we conclude that the AC-coupled topology with a bias current of 240 μ A and a varactor tuning mechanism for both the coarse and the fine tuning is the optimal solution. The area can be minimized by using a single differential inductor.

Frequency Divider Design

The use of a bipolar process, while enabling a low power oscillator, poses some challenges in the design of the frequency divider. The main target in the design of the frequency divider is to reach the highest operating frequency with the minimum power consumption. It is reasonable to assume that parasitic capacitances and resistances must be minimized in order to keep the internal circuit time constants as low as possible. This will improve the maximum achievable frequency for a given bias current. Therefore, the number of interconnections has to be kept as small as possi-

Fig. 4.26 Block diagram of a binary divide-by-two frequency divider using the master-slave principle



ble. This translates in a small number of components, especially of components with three or more terminals. In frequency dividers these components are the transistors.

The most straightforward way to have a divide-by-two frequency divider is to use two latches connected in a master-slave configuration as shown in Fig. 4.26. The required V_T signal in Fig. 4.26 is generally generated by a level shifter [83]. Following the approach in [83] it is possible to give a rough estimation of the minimum power consumption of a master-slave based bipolar frequency divider.

The max operating frequency $f_{\max}$ of this divider can be approximated by the following relation:

$$f_{\max} = \frac{1}{2(\tau_{LS} + \tau_{LT})} \quad (4.23)$$

where τ_{LS} is the delay of the level shifter and τ_{LT} is the delay of the latch stage. If we suppose to have $\tau_{LS} = \tau_{LT} = \tau$, the maximum allowed delay $\tau_{\max}$ is

$$\tau_{\max} = \frac{1}{4f_{\max}} \quad (4.24)$$

It is possible to express τ as a function of the latch current I_{LT} in the following way:

$$\tau = aI_{LT} + \frac{b}{I_{LT}} \quad (4.25)$$

where a and b are constants, which depend on the latch topology and on the technology used and they are defined more in detail in [83]. From the previous equation it is possible to derive the latch current as a function of the delay τ and the parameters a and b :

$$I_{LT} = \frac{\tau}{2a} \left(1 - \sqrt{1 - \frac{4ab}{\tau^2}} \right) \quad (4.26)$$

For $\frac{4ab}{\tau^2}$ much smaller than one, the following approximation can be used:

$$\sqrt{1 - \frac{4ab}{\tau^2}} \cong 1 - \frac{2ab}{\tau^2} \quad (4.27)$$

and therefore, from (4.24), (4.26) and (4.27), the minimum required current $I_{LT,\min}$ for the latch can be expressed as follows:

$$I_{LT,\min} = 4bf_{\max} \quad (4.28)$$

Considering $f_{\max} = 1.83$ GHz and supposing to use $b = 7.02 \times 10^{-14}$ [83] the required latch current is 206 μA . Therefore, the master slave divider, excluding the level shifter requires 412 μA .

The level shifter current consumption can be approximated as follows:

$$I_{LS} = \frac{C_{je,buf}}{(1.1^2 - 1)\tau_F} \quad (4.29)$$

where

$$\tau_F = \frac{1}{2\pi f_T} \quad (4.30)$$

where f_T can be assumed roughly equal to 20 GHz and $C_{je,buf}$ is the emitter junction capacitance of a source follower and it can be assumed equal to 41 fF. The estimated current required in the level shifter is roughly equal to 780 μA . The circuit based on the master-slave principle turns to be rather complicated. This makes it very difficult to have a circuit with few parasitic elements because two separate latch elements are required that control one other alternately.

To reduce the circuit complexity and therefore, the number of parasitics, in a practical implementation, other principles must be used. For example a great simplification would consist in using only one of such bistable elements. The work generally performed in a master-slave based divide-by-two frequency divider can be done by the short-term memory effect always present in the base-emitter capacitance of the transistor.

One of these kinds of dynamic frequency dividers is the so-called traveling-wave frequency divider [84]. Its schematic is shown in Fig. 4.27. The circuit consists of three differential amplifiers: the input amplifier T_5 , T_6 and the differential amplifiers T_1 , T_3 and T_2 , T_4 with complementary control. The collector and base termi-

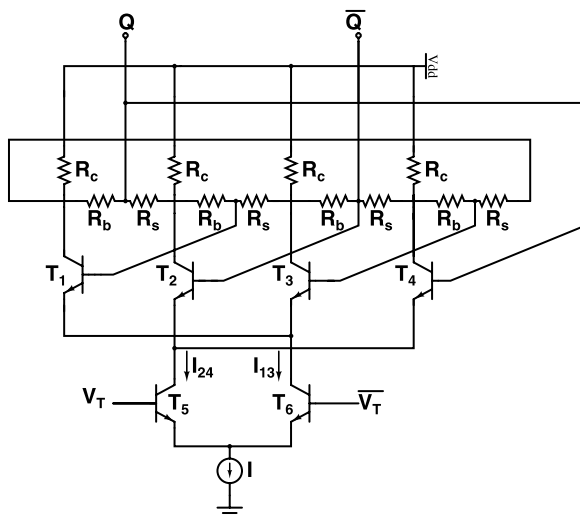


Fig. 4.27 Traveling-wave divider schematic

nals of the transistors T_1 to T_4 are coupled in a cyclic manner via resistances R_b and R_s .

The circuit has a cyclic configuration. Therefore, it is possible to describe its operation during only a quarter of the period of the output voltage $V_Q - V_{\bar{Q}}$. The voltages during the remaining three-quarters of the period can be easily derived from the first quarter by cycling interchange of the transistors T_1 to T_4 . To simplify the description, the variation of the currents I_{13} and I_{24} during the switching of differential amplifiers T_5 and T_6 will be approximated by a linear variation. The switching time t_s is determined by the rise time of V_T . In Fig. 4.28 are shown the collector and the base voltages of the transistors T_1 to T_4 at five successive instants during the switching time.

Figure 4.28(a) illustrates the situation at the starting of the switching event ($t = 0$). It is supposed that at the beginning of the switching event $I_{13} = I$ and $I_{24} = 0$. Furthermore, we can assume that I_{13} flows only through transistor T_1 . Most of the T_1 collector current flows through the load resistance R_c , while a small fraction flows via R_b and R_s , which are supposed to be equal in this particular example, through the collector resistances of T_2 and T_4 . A really tiny current component flows through the collector resistance of T_3 . The effects on the output voltages are the following: T_1 has a low collector voltage, T_3 a high collector voltage and T_2 and T_4 an intermediate collector voltages. Given the fact that R_b and R_s are supposed, in this example, to be equal, the base voltages (in dash in Fig. 4.28) are mid-way between the collector voltages.

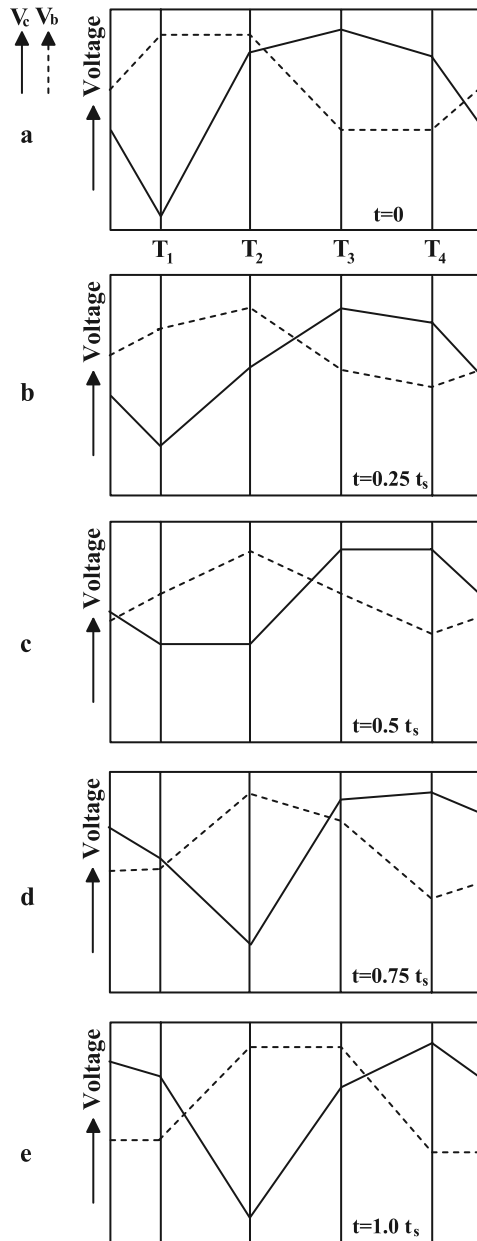
At $t = 0.25t_s$ (Fig. 4.28(b)) $I_{13} = 0.75I$ and $I_{24} = 0.25I$. Most of the I_{24} flows through T_2 , since this transistor has a higher base voltage than T_4 . Therefore, the collector voltage of T_2 has decreased while the collector voltage of T_1 has increased because of the current I_{13} has decreased.

In the middle of the switching time ($t = 0.5t_s$), $I_{13} = I_{24} = 0.5I$. I_{13} flows through T_1 , while I_{24} flows through T_2 resulting in a equal base voltages for the transistors T_1 and T_3 (Fig. 4.28(c)). The current continues to flow still only through T_1 because the switching process is faster than the transistor internal switching time, which has the same order of magnitude as t_s .

In Fig. 4.28(d) the situation at $t = 0.75t_s$ is depicted. At this moment $I_{13} = 0.25I$. It is assumed that this current still flows through T_1 only though the base voltage of T_1 is lower than the base voltage of T_3 . This situation occurs only for a very short switching time t_s . If t_s becomes comparable with the internal switching time of the transistors, a part of the current I_{13} will flow through T_3 decreasing the collector voltage of T_3 and increasing the collector voltage of T_1 . It follows that the base voltage of T_4 increases, while the base voltage of T_2 decreases. This condition is less favorable to the end result of conducting the current I_{24} through T_2 only. Concluding, if the switching time t_s is larger than a critical switching time t_{sc} , the base voltages of transistors T_2 and T_4 are too close to each other and the circuit does not work anymore as a divide-by-two frequency divider.

Figure 4.28(e) shows the situation at the end of the switching time t_s when all the current flows only through transistor T_2 and the base voltage of T_2 is higher than that of T_4 . This is a stable situation.

Fig. 4.28 The collector voltage V_c and the base voltage V_b of the transistors T_1 to T_4 during one step of the switching process. After four steps the starting position is reached again



The curves in Fig. 4.28 shift to the right at every switching time. Therefore, after four switching times they return in the initial position. Four switching times correspond to two periods of the trigger voltage V_T . Therefore, the circuit behaves as a divider-by-two in the frequency domain.

For an optimal circuit behavior the input voltage of the differential amplifier T_1 , T_3 should be as small as possible because it must fastly change polarity during the switching time. On the other hand, given the fact that we want I_{24} to flow through T_2 only, the input voltage of the differential amplifier T_2 , T_4 should be as large as possible. This behavior can be obtained by changing the ratio between R_b and R_s .

In the designed prototype, an optimal size of the transistor has been chosen in order to obtain this effect without using resistances and the circuit has been optimized to work around 1.83 GHz. Increasing the ratio between R_b and R_s while increasing the input sensitivity at a certain frequency of operation, reduces the critical switching time making the circuit less sensitive at low frequencies. An optimal transistor sizing, reduces the number of parasitics and allows to have an optimal sensitivity around the required oscillator tuning range. Because the divider has to work only in a limited frequency range (from 1.804 GHz to 1.856 GHz), low frequency behavior is not important. Therefore, R_b is set to zero, while R_s to infinite, maximizing in this way the input sensitivity.

This is an important feature of the TWD. Indeed, the system requires to divide-by-two only in a certain frequency range. The TWD has a higher sensitivity compared to a static master-slave based divider-by-two. But this happens only in a limited frequency range and at the expenses of a lower sensitivity at frequencies far from the optimal. A higher sensitivity means a lower required output swing for the VCO and therefore, a lower VCO power consumption according to (4.22). This reduction allows still to meet the phase noise spec of -103 dBc/Hz at 450 kHz (see Table 4.2). Indeed, the simulated phase noise at 210 μ A bias current for the AC-coupled LC VCO is around -111.2 dBc/Hz at 450 kHz far from the carrier (see Sect. 4.2.4), which is around 8 dB better than required. The start-up condition requires only 70 μ A and therefore, it is possible to reduce the oscillator amplitude (and therefore, its bias current) by improving the TWD sensitivity.

The simulated input and output transient waveforms are shown in Fig. 4.29. These waveforms are obtained in the condition of an unloaded divider and they

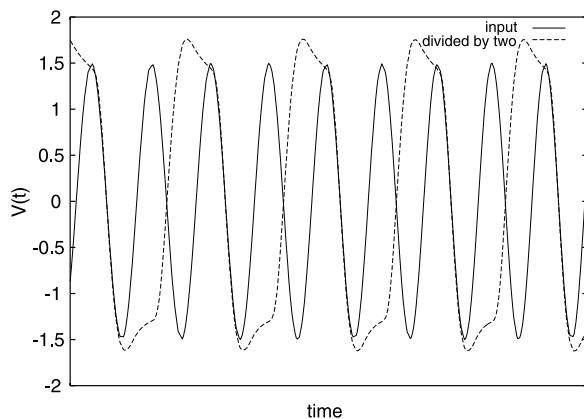


Fig. 4.29 Input and output transient waveforms for the TWD

are differentially probed. The input signal frequency is 1.83 GHz. The simulated divider current consumption is equal to 150 μA .

Output Buffer

The TWD has an output swing in the order of a few hundreds of mV_{pp} when it is not heavily loaded. Therefore, to minimize the load on the divider the output buffer shown in Fig. 4.30 has been used. The first stage of the buffer is a common-emitter gain stage (with a gain of about 2). It is used to minimize the load seen by the divider and to boost the amplitude of the signal driving the output stage. Indeed, the output stage needs to have a large size to conduct the required current to generate the wanted output power on the 50 Ω load. The current used in the pre-amplifier stage is around 160 μA .

Supposing that a -18 dBm output power is required,⁹ the required bias current on the output stage can be easily calculated as follows. The required voltage swing on the 50 Ω load resistance can be evaluated using the following equation:

$$V_{\text{out,rms}} = \sqrt{\frac{10^{\frac{-18}{10}}}{1000}} \times R_{50\ \Omega} \approx 28\ \text{mV} \quad (4.31)$$

Now, the peak voltage (V_p) is equal to $\sqrt{2}V_{\text{out,rms}} \approx 40\ \text{mV}$. Finally, the peak current is $\frac{V_p}{50} \approx 800\ \mu\text{A}$. Therefore, the overall output stage current consumption is around 1 mA from a 1.8 V supply.

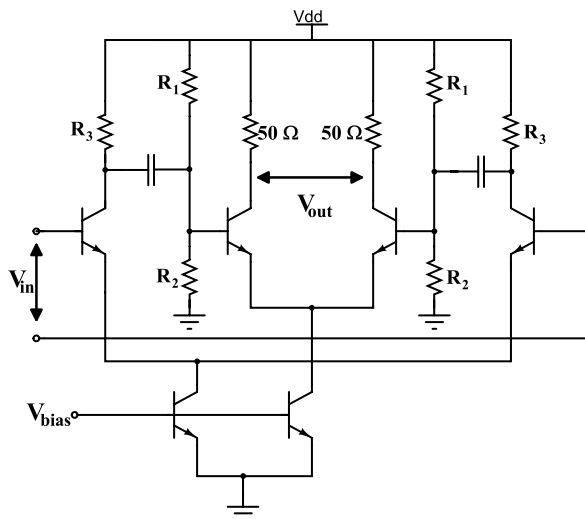
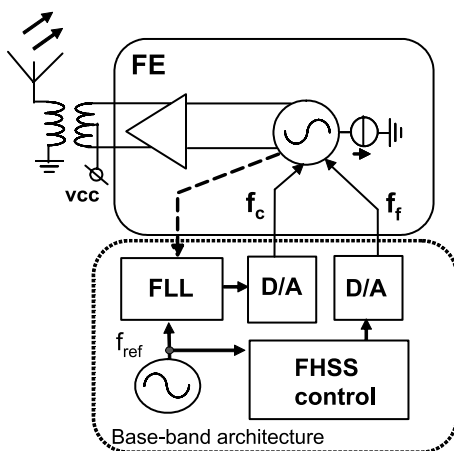


Fig. 4.30 Output buffer

⁹Some margin is required, compared to the earlier mentioned -25 dBm (see Table 4.2), in order to account for cable and board losses in a real prototype.

Fig. 4.31 Power-VCO based frequency hopping transmitter architecture



4.2.5 Power-VCO Based Architecture

The schematic block diagram of this architecture is shown in Fig. 4.31. This architecture, like the VCO-divider based architecture, utilizes the frequency pre-distortion concept in order to directly modulate the VCO. In the LC-divider based architecture some of the power is “wasted” in blocks other than the PA (or the output buffer) because of implementation issues. For example, in the LC-divider based architecture the VCO runs at twice the required frequency in order to avoid frequency pulling from the antenna. Therefore, a divider is required in order to obtain the required center frequency. These choices increase the overall required power consumption.¹⁰

The proposed architecture is able, by using the current reuse concept, to reduce the complete RF section to a single RF block coupled to the antenna via a balun. Therefore, the divider becomes obsolete, and the VCO runs at half the frequency of the LC-divider based architecture. As it is explained later in this section, isolation from the antenna is obtained by using a current buffer, which directly drives a 50 Ω matched antenna, via a balun.

The output of the VCO is used to drive an FLL for initial frequency calibration, while a low-accuracy reference signal is used to clock the hopping generator (FHSS control in Fig. 4.31). The digital output of the FHSS control block is converted into an analog voltage via a low power DAC and this voltage is then fed to the VCO varactors. In this way the capacitance and therefore, the transmitter output frequency, is varied in order to transmit a data bit over an assigned frequency bin.

¹⁰The optimal power consumption in terms of transmitter efficiency is obtained when all the power used in the RF section is effectively radiated from the antenna.

Power-VCO Design Procedure

As shown in Sect. 4.2.4, due to the use of the tail current source, cross-coupled based oscillators present a phase-noise performance worse than classical types of oscillator with one of the active device ports grounded. This, in the end, requires a larger DC current for a given phase noise specification. For this reason, in this case, a Colpitts based oscillator has been chosen and its schematic is shown in Fig. 4.32.

The Colpitts oscillator is generally arranged in a common-base or a common-collector configuration. Due to its inverting relation between input and output, the common emitter configuration requires an additional inverting stage and therefore, it is not easily applicable in a Colpitts configuration. The common-base configuration has a relatively low input impedance, reducing the in-circuit Q of the tank, which is a disadvantage at high frequencies. On the other hand it shows the highest output impedance for the same output power. To have, at the same time, a high input and output impedance, a common collector configuration with a cascoded buffer stage may be used avoiding any intermediate buffer stage through current reuse.

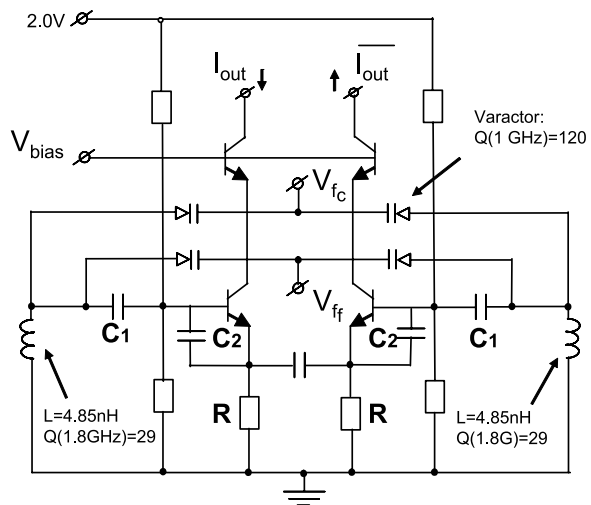


Fig. 4.32 Circuit schematic of the differential cascoded Colpitts power-VCO

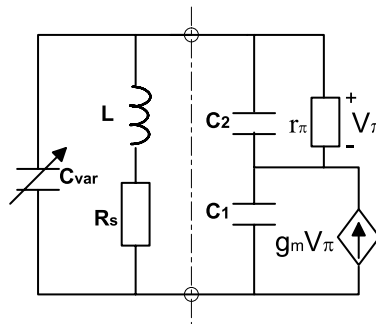


Fig. 4.33 Common collector Colpitts small signal AC equivalent

The AC small signal model for a single-ended common collector Colpitts oscillator is shown in Fig. 4.33. The positive feedback realized by the capacitances C_1 and C_2 produces a negative resistance in series with two capacitances when the β of the transistor is sufficiently large. It can be proven that the impedance looking in the base of the transistor can be approximated by the following equation

$$Z_B(j\omega) = \frac{1}{j\omega C_1} + \frac{1}{j\omega C_2} - \frac{g_m}{\omega^2 C_1 C_2} \quad (4.32)$$

The oscillation frequency is

$$\omega = \frac{1}{\sqrt{L(C_v + C_{12})}} \times \frac{1}{\sqrt{1 - R_s^2 \frac{C_v}{L}}} \quad (4.33)$$

where C_{12} is equal to $\frac{C_1 C_2}{C_1 + C_2}$, C_v is the varactor capacitance and all the tank losses have been lumped in the resistance R_s . In reality C_1 takes in account also the parasitic base-emitter capacitance of the transistor. The necessary condition to have a steady-state oscillation is that the negative resistance lumped element is bigger than the total loss of the tank, which means:

$$g_m > \omega^2 (C_v + C_{12})(C_1 + C_2)R_s \quad (4.34)$$

A relation between bias current and output power at the fundamental frequency can be derived considering that the AC current component at the resonant frequency is the same as in the common-base configuration [85].

$$I_{\omega_0} = 2I_B \frac{I_1\left(\frac{V_m}{V_T}\right)}{I_0\left(\frac{V_m}{V_T}\right)} \quad (4.35)$$

where I_1 is the modified Bessel function of the first kind and order one. The current component at the fundamental will divide itself between the impedance seen looking toward the tank and the impedance constituted by the capacitance C_1 in parallel with the emitter degeneration resistance R .

Calling these two impedances Z_1 and Z_2 respectively and calling the equivalent impedance between the base and the emitter of the BJT Z_π , the amplitude of the sinusoidal voltage across the base-emitter junction can be expressed by the following

$$V_m = 2I_B \frac{I_1\left(\frac{V_m}{V_T}\right)}{I_0\left(\frac{V_m}{V_T}\right)} \left| \frac{Z_1}{Z_2 + Z_1} Z_\pi \right| \quad (4.36)$$

Equation (4.36) should be solved numerically to obtain the value of V_m and therefore, the current component at the fundamental frequency. Simulated and predicted results for both the output current and base-emitter signal voltage V_m are represented in Fig. 4.34. As can be seen the theoretical analysis can accurately predict the required DC current for a certain signal current at the fundamental frequency and therefore, for a certain output power.

Fig. 4.34 Common collector Colpitts small signal AC equivalent

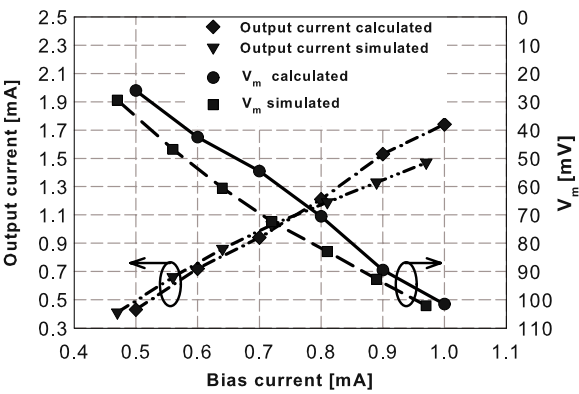
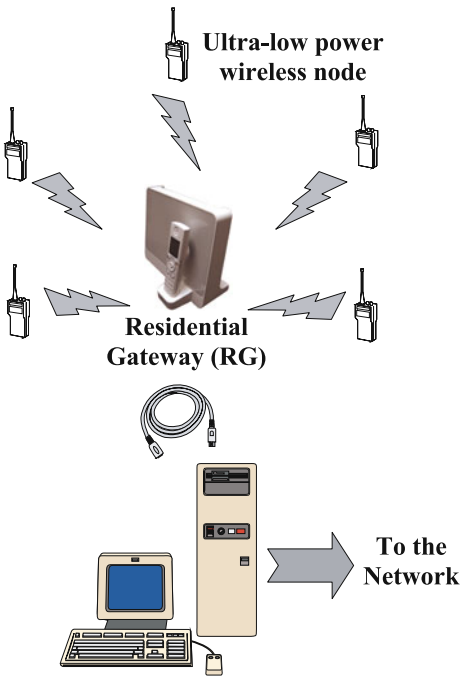


Fig. 4.35 Asymmetric wireless sensor network



4.3 Receiver Architecture

In the asymmetric scenario the transmitter node has a very limited power budget. The receiver node, commonly called RG receives the data from the transmitter nodes and, via a wireless or a wired connection, it moves the data through the network to a collection point. The collection point can, therefore, be very far away from the location where the wireless sensor network is deployed. An example of such a network is shown in Fig. 4.35.

While in Sect. 4.2.3 we focused on the analysis and design of the ultra low power transmitter node, this section will focus on the RG. Given the fact that the RG has a virtually unlimited power budget, this section will not focus on the design of the RG itself, because this is outside the scope of this book. The section will mainly put attention on the most important receiver requirements in order to have a robust wireless link.

Indeed, as mentioned in Sect. 4.2, two main non-idealities need to be solved at the receiver side in order to correctly demodulate the incoming data. These non-idealities are the following:

- Different center frequencies of the transmitter and the receiver
- Residual frequency error after frequency pre-distortion in the transmitter

The first non-ideality comes from the low-accuracy (1% accuracy) reference signal used on the transmitter side. This choice allowed to design a crystal-less transmitter. The second non-ideality arises from the finite precision of the frequency pre-distortion transmitter architecture described in Sect. 4.2 and from temperature or supply variation, which can slowly pull the VCO frequency.

4.3.1 RX-TX Center Frequency Alignment Algorithm

Before any data transmission the RX node (RG in this scenario) and the TX node need to align their center frequencies within an accuracy of few parts per million in order to avoid any degradation of BER performance.

The alignment requires the TX to send some kind of information on its center frequency so that the receiver center frequency can be modified in order to match the TX center frequency within a few parts per million accuracy. Given the fact that during this time the TX node is using energy to transmit such information, it is important that this energy overhead is kept to the minimum. Therefore, the frequency recovery algorithm implemented at the receiver side has to be accurate and fast. Fast means that the required time to achieve the required frequency synchronization within a certain accuracy must be small compared to the time required to transmit the data packet. In this way the energy overhead will be negligible.

Several methods exist to recover the frequency offset [86, 87], but they make use of the phase information embedded in the transmitted signal and therefore, they can only recover frequency offsets smaller than the data-rate. For low data-rate applications a new approach is required, which can recover frequency offsets much larger than the data-rate.¹¹

Another possibility is to use a PLL and, after the calibration process has been completed, to store the tuning information on the main capacitor of the loop filter [88]. This approach will anyhow require a PLL, and depending on the size of

¹¹Considering the channel allocation mentioned in Sect. 4.2, the maximum offset the algorithm needs to recover is around 9.3 MHz. The system data rate, on the contrary, is in the order of few kilobit per second.

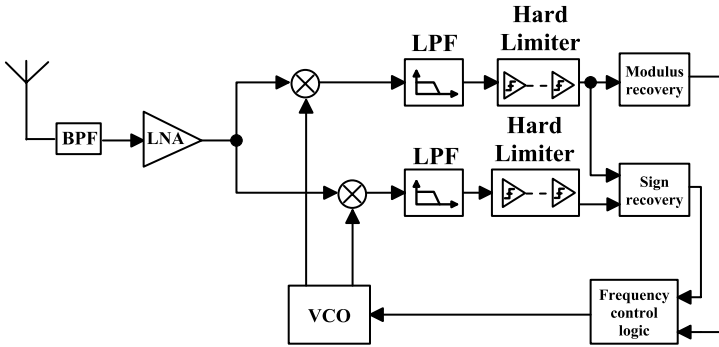
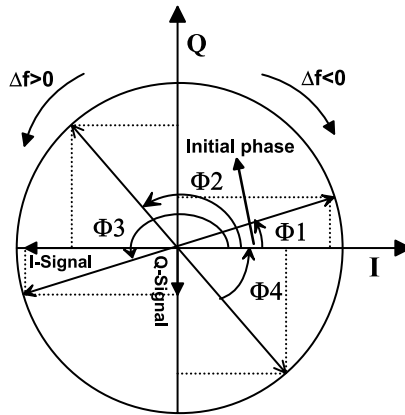


Fig. 4.36 Direct-conversion architecture with frequency offset recovery circuitry

Fig. 4.37 I - Q plane representation of the frequency offset component with different initial phases



the packet sent, it is possible that the required storage capacitance can be large. This will increase chip area and therefore, device cost. The algorithm developed to align the prototype TX and RX center frequencies has been described in [77, 89].

The frequency offset recovery algorithm can be used in both direct-conversion and super-heterodyne receiver architectures and its schematic block diagram is shown in Fig. 4.36 in the case of a direct-conversion receiver. At the beginning of the transmission, the TX node will send a pilot tone at the carrier frequency. This pilot tone is then mixed down in quadrature with the receiver LO frequency, which is supposed to be at the beginning of the process in the middle of the ISM band. Generally, the mixing products contain two components. One component has a frequency equal to the sum of the mixing tones while the second component is equal to the difference between the same tones. While the first component is filtered out by the LPF, the second one contains the information about the frequency offset between the TX and RX LOs.

The frequency offset acquisition is obtained evaluating the modulus and the sign of the offset. Furthermore, the phase of the incoming signal is also unknown and the recovery algorithm should be able to compensate also for it. The frequency com-

ponent at the output of the LPF will range within around ± 9 MHz. This frequency component can be depicted in the I - Q plane for different initial phase (Φ_1 to Φ_4) as shown in Fig. 4.37.

The speed at which the phasor rotates is the frequency offset modulus, while the direction of the rotation gives the sign of the frequency offset. The latter information allows the frequency control logic in Fig. 4.36 to increase or decrease the VCO control voltage by an amount proportional to the frequency offset modulus, increasing or decreasing accordingly the VCO frequency.

If the phasor rotates clockwise, then the difference between the TX carrier frequency and the RX LO frequency is negative, while if the phasor rotates counter-clockwise the same difference is positive and this does not depend on the initial phase of the incoming tone. A flow chart of the implemented algorithm is shown in Fig. 4.38.

The evaluation of the modulus requires at the receiver side a fast and accurate reference clock. Due to the fact that the RG will be mains supplied and its cost can be higher than the TX cost by some order of magnitude, its implementation does not pose any problem.

The reference clock at the RG should be fast enough to reduce the error in the evaluation of the frequency offset modulus to a negligible value. Indeed, if N is the number of reference clock periods contained in half of the frequency offset signal period, then this value in reality can be $N \pm 1$ due to the initial phase difference between the two signals. The error in the estimation of the frequency offset modulus, at each iteration, can be upper bounded by the following relation:

$$\Delta f \simeq \frac{f_{\text{ref}}}{2N^2} \quad (4.37)$$

while the estimated frequency offset modulus is:

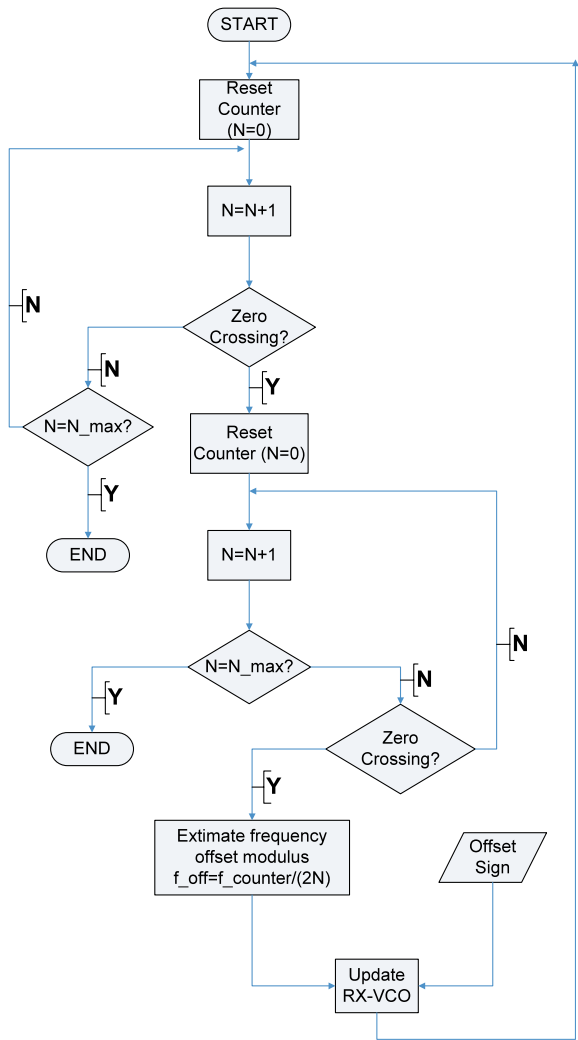
$$|f_{\text{off}}| = \frac{f_{\text{ref}}}{2N} \quad (4.38)$$

where f_{ref} is the frequency of the reference clock.

The choice of the optimal reference clock frequency is not trivial. It is necessary to minimize the acquisition time in order to minimize the transmitter power consumption. In Fig. 4.39(a) the acquisition time versus the initial frequency offset and the reference clock frequency is plotted. Four different regions are visible. In the region A, the required time is roughly constant and independent from the frequency of the reference clock. Indeed, for small initial frequency offsets, the acquisition time depends from the offset value itself and the algorithm is repeated only one time to reach the required accuracy in the frequency acquisition. This is visible in Fig. 4.39(b). Of course, the higher the reference clock frequency, the wider the initial frequency offset range that requires a single iteration to achieve the required accuracy.

In the region B, the initial frequency offset requires more iterations, generally two, sometimes one. Increasing the reference clock frequency leads to a more accurate estimation at the first iteration but never accurately enough. Therefore, a second

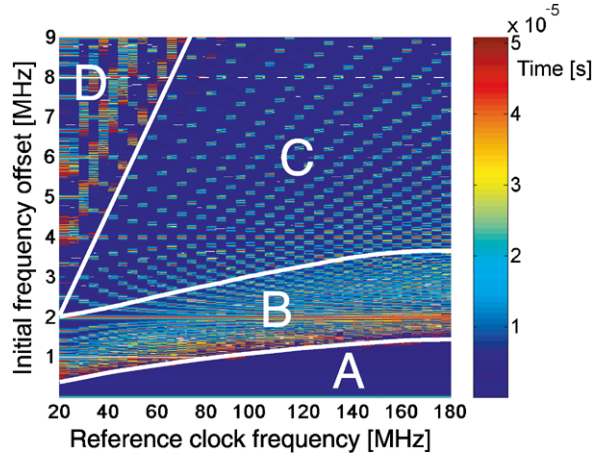
Fig. 4.38 Frequency offset recovery algorithm flowchart



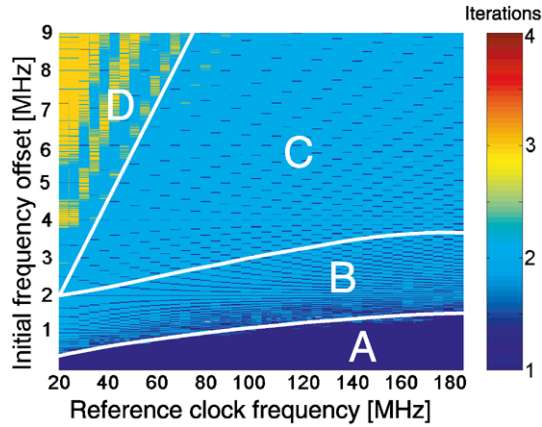
iteration is needed, but due to the more accurate first acquisition, the second iteration will require more time. Therefore, this area tends to enlarge increasing the reference clock frequency.

The area C requires between one and two iterations (mainly two). This is the area, which requires the smallest acquisition time (excluding the A area) and it tends to enlarge increasing the reference clock frequency. Obviously, the initial frequency offset is so large that the first acquisition is almost never sufficient. On the other hand, when a second iteration is required, it does not take a long time. Indeed, in the first acquisition the residual offset is still large enough requiring less time to reach the wanted accuracy.

Fig. 4.39 Iteration number and acquisition time vs. reference clock frequency (initial offset $\in [10 \text{ kHz}, 9 \text{ MHz}]$)



(a) Acquisition time



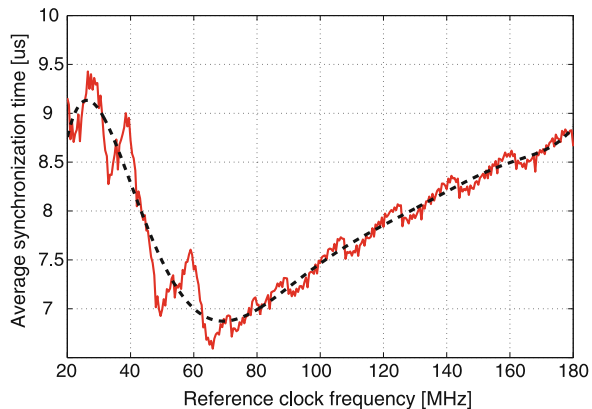
(b) Number of iterations

Finally, the region D is dominated by the large number of iterations required (mainly three iterations). In this case the reference clock frequency is too small. Therefore, the first iterations are very inaccurate requiring the algorithm to be repeated more than two times.

The criteria in the choice of the optimal reference clock frequency is based on the minimization of the average acquisition time for several transmitters. Considering the initial frequency offset uniformly distributed, and the whole frequency offset range subdivided in “ k ” slots, then the average acquisition time is

$$\text{Time}_{\text{avg}} = \frac{1}{\Delta f_{\text{off}}} \sum_{i=1}^k T_{\text{acq},k} \Delta f_k \quad (4.39)$$

Fig. 4.40 Average acquisition time vs. reference clock frequency



where Δf_{off} is the initial frequency offset range, $T_{\text{acq},k}$ is the acquisition time in the k th frequency slot and Δf_k the frequency width of the k slots. A picture of the average acquisition time versus the reference clock frequency is shown in Fig. 4.40.

The minimum average acquisition time is obtained for a reference clock frequency of around 70 MHz. The time required to reduced the initial frequency offset below the wanted maximum value is in any case less than 50 μs (see Fig. 4.39). Then, to complete the acquisition process, it is necessary to wait one more iteration as can be seen looking at Fig. 4.38. The last iteration will produce an overflow on the counter (the counter reaches N_{max}), which will end the acquisition process.

Considering a 70 MHz reference clock frequency and a 6 kHz maximum residual frequency error, then a 14-bit counter is required. This will require around 230 μs to produce an overflow. Therefore, the complete acquisition process requires less than 300 μs to be accomplished. At 10 kbps this is equivalent to 3 bits data overhead, which is a reasonable value in terms of energy consumption.

The next step in the proposed algorithm deals with the recovery of the frequency offset sign. For this purpose, a quadrature down-conversion scheme is needed. To understand the basic principle behind the algorithm it is useful to discuss Tables 4.7, 4.8, 4.9 and 4.10.

Supposing to look only at the transition type (which means in Fig. 4.37 when a phasor crosses the I or the Q axis), it is possible to have two situations, which lead to the same output even when the sign of the frequency offset is opposite. Indeed, looking at Table 4.7 to Table 4.10, independently from the initial phase difference between the incoming signal and the LO synthesized frequency, the frequency offset component experiences the same type of transitions either when the frequency offset is positive or it is negative. Therefore, to correctly evaluate the frequency offset sign independently of the initial phase difference between the incoming pilot tone and the RX LO synthesized frequency, additional information is needed.

Looking at Table 4.7 to Table 4.10, the required information can be found in the time. Indeed, when, for example, $\Delta f > 0$ and the initial phase is between 0 and $\frac{\pi}{2}$ (Table 4.7, 1st column), the Q signal crosses the zero (I axis in the I - Q plane) always after the I signal. In the first column of Table 4.7 an initial phase difference

Table 4.7 I and Q -signal transitions for initial phase difference between TX and RX oscillators $\in [0, \frac{\pi}{2}]$

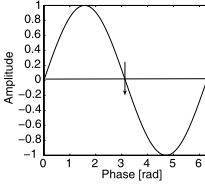
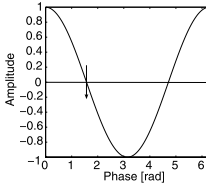
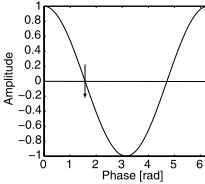
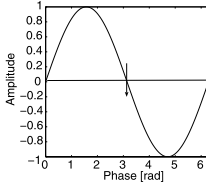
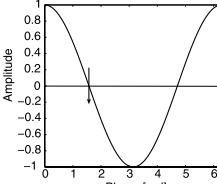
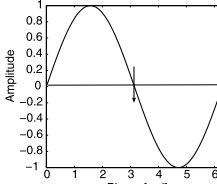
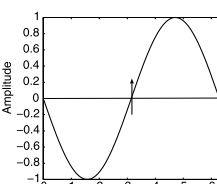
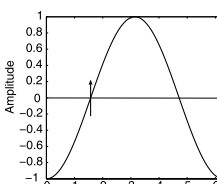
Signal	$[0, \frac{\pi}{2}], \Delta f > 0$	$[0, \frac{\pi}{2}], \Delta f < 0$
Q -Signal		
I -Signal		

Table 4.8 I and Q -signal transitions for initial phase difference between TX and RX oscillators $\in [\frac{\pi}{2}, \pi]$

Signal	$[\frac{\pi}{2}, \pi], \Delta f > 0$	$[\frac{\pi}{2}, \pi], \Delta f < 0$
Q -Signal		
I -Signal		

of zero radians has been supposed. When $\Delta f > 0$, the Q signal crosses the I axis always before the I signal. Therefore, when $\Delta f > 0$ and the initial phase difference is between 0 and $\frac{\pi}{2}$ the Q -signal lags and the I -signal leads, while when $\Delta f > 0$, even though the transition of the I and Q signal are still both negative, the I signal lags and the Q signal leads (the initial phase difference in Table 4.7, 2nd column is $\frac{\pi}{2}$).

Table 4.9 I and Q -signal transitions for initial phase difference between TX and RX oscillators $\in [\pi, \frac{3\pi}{2}]$

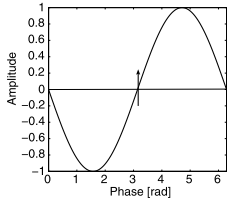
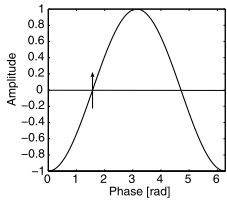
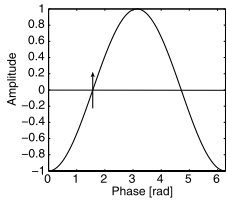
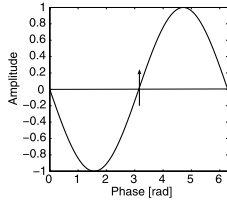
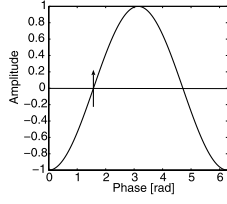
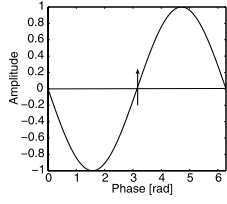
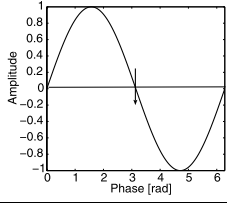
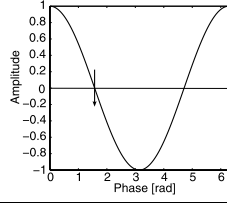
Signal	$[\pi, \frac{3\pi}{2}], \Delta f > 0$	$[\pi, \frac{3\pi}{2}], \Delta f < 0$
Q -Signal		
I -Signal		

Table 4.10 I and Q -signal transitions for initial phase difference between TX and RX oscillators $\in [\frac{3\pi}{2}, 2\pi]$

Signal	$[\frac{3\pi}{2}, 2\pi], \Delta f > 0$	$[\frac{3\pi}{2}, 2\pi], \Delta f < 0$
Q -Signal		
I -Signal		

The same considerations can be applied for the other possible phase relation ranges shown in Table 4.8 to Table 4.10, in which the initial phase difference has been chosen to highlight the leading and the lagging signal components. The evaluation process is summarized by the flowchart in Fig. 4.41.

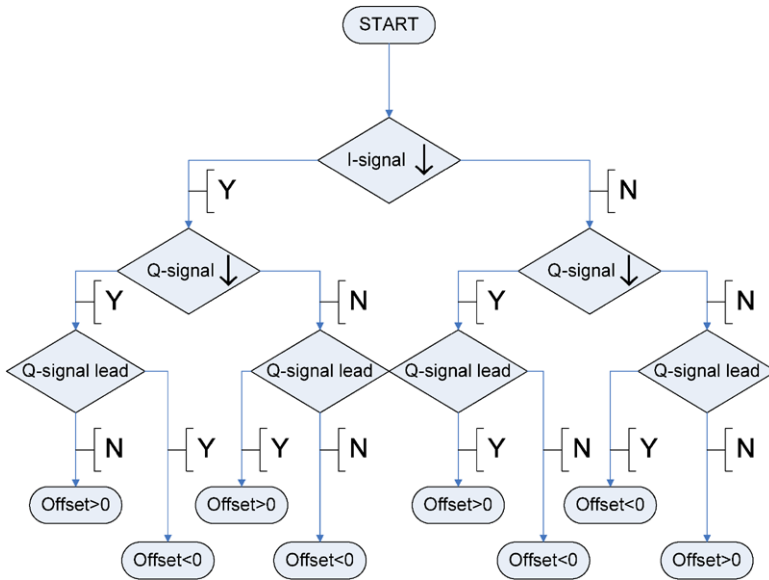


Fig. 4.41 Frequency offset sign recovery method flowchart

Simulation Results

To test the proposed algorithm, a full wireless link has been implemented using Simulink. The modulus recovery circuitry has been implemented using a 14-bit counter, two D-FF and a few logic ports. The sign recovery circuitry requires only 16 D-FFs and a few logic ports. Simulations show that starting from 9 MHz frequency error between TX and RX LOs, which corresponds to approximately 1% frequency accuracy, the frequency offset can be lowered in two steps to less than 6 kHz (6.5 ppm) in less than 300 μ s.

In Fig. 4.42 the frequency offset signal is plotted versus the recovery time. At the beginning of the acquisition process, the frequency offset is in the MegaHertz range (1.58 MHz). After the first correction the frequency is 916.556 MHz. The frequency difference is still too large and therefore, a second iteration is needed. At the end of the calibration process the synthesized frequency at the receiver side is 916.5804 MHz, which is only 400 Hz far from the TX carrier frequency (916.58 MHz). The inset of Fig. 4.42 shows the downconverted residual frequency error during the calibration process (at the output of the receiver mixers). The required acquisition time is approximately 42 μ s.¹²

¹²This does not include the last iteration required to confirm the frequency acquisition, which takes around 230 μ s.

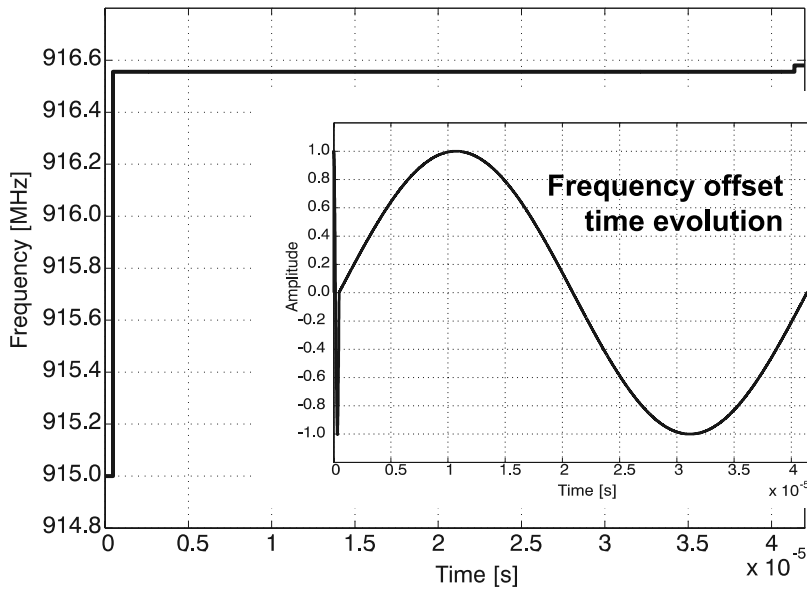


Fig. 4.42 Receiver LO synthesized frequency vs. time

4.3.2 Residual Frequency Error after Pre-distortion

This frequency error has two main sources:

- Temperature or supply variations
- Finite precision of the pre-distortion algorithm

These two sources require two different strategies, which will be further discussed in this section.

Temperature or Supply Variations

A straightforward way to cope with this problem is either to use an AFC loop or to adopt at the receiver side a demodulation algorithm, which is insensitive, in first approximation, to the frequency offset induced by those variations. While the first approach requires to increase the circuit complexity and therefore, the power consumption of the transmitter, the second approach can use the large processing power of the receiver to correctly demodulate the incoming data. In the following part of this section the second approach is described in detail [62].

Among all the demodulator architectures, four of them, suited for ultra-low power implementation, have been chosen and their performances have been simulated using Simulink models under different frequency offsets in an AWGN channel. These four types of demodulators are:

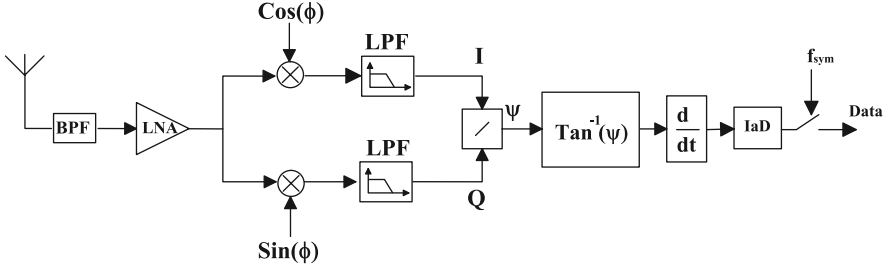


Fig. 4.43 Conventional Arctangent receiver

- Arctan-Differentiated deModulator (ADM) [90]
- Correlation demodulator [91]
- Digital Cross-Differentiate Multiply (DCDM) demodulator [92]
- Short-Time Discrete Fourier Transform (ST-DFT) demodulator [93]

Another widely used digital demodulator is the zero-crossing demodulator. Unfortunately in [94] it has been shown that when the frequency offset is equal to the 6% of the data rate, the SNR degradation reaches 4.8 dB. In the field of low data-rate applications, this will limit the maximum uncorrected frequency offset, without significant degradation in the BER performance, to a few hundreds of Hertz, which is quite difficult to achieve in practice without an AFC system.

In the *ADM demodulator*, the Frequency Modulation (FM) discriminator is followed by an INtegration and Dump (IaD) circuit and by a decision circuit, which retrieves the transmitted data. The block diagram of this demodulator is shown in Fig. 4.43.

The received signal is first downconverted to a complex baseband signal to obtain the In-phase (*I*) and the Quadrature-phase (*Q*) signal components and then these components are used to retrieve the transmitted information. The *I* and *Q* components are:

$$\begin{cases} s_I(t) = -\frac{A(t)}{2}[\cos(2\pi f_d \int_{-\infty}^t m(\tau) d\tau)] \\ s_Q(t) = -\frac{A(t)}{2}[\sin(2\pi f_d \int_{-\infty}^t m(\tau) d\tau)] \end{cases} \quad (4.40)$$

where $A(t)$ is the carrier frequency amplitude, f_d is the frequency deviation, $m(\tau) = \sum_{n=-\infty}^{n=+\infty} a_n d(\tau - nT_s)$ is the transmitted signal with $a_n = \pm 1$, and $d(t)$ the rectangular pulse over a symbol time T_s .

Now, taking the ratio between the *Q* signal and the *I* signal and extracting the arctangent function it is possible to recover the signal information. The output of the inverse tangent block is the following:

$$\phi(t) = 2\pi f_d \int_{-\infty}^t m(\tau) d\tau \quad (4.41)$$

The transmitted signal is then retrieved by differentiation of $\phi(t)$, which can be performed in the digital domain subtracting one sample from the previous one, or

by passing $\phi(t)$ through a filter that approximates the response $H(\omega) = j\omega$. The IaD block and the decision block will retrieve the digital transmitted data.

In the case in which the signal passes through an AWGN channel, remembering that the output of the discriminator is the derivative of the phase in (4.41) ($\dot{\psi} = d\psi(t)/dt$), and considering that the IaD filter integrates over one symbol period, the output of the IaD is [95]

$$\Delta\psi = \Delta\phi + \Delta\eta + 2\pi N \quad (4.42)$$

where $\Delta\phi$ is the signal component, $\Delta\eta$ is the continuous phase component and $2\pi N$ is the click noise component due to spikes at the discriminator output when the SNR is below the threshold point, in a symbol period.

When a frequency offset is present, an additional phase term appears in (4.42):

$$\Delta\psi = \Delta\phi + \Delta\eta + 2\pi N + \Delta\theta \quad (4.43)$$

where $\Delta\theta$ is the offset induced phase difference. If we consider a noiseless environment ($2\pi N = 0$) and no continuous phase component present ($\Delta\eta = 0$), it is easy to notice from (4.43) that, supposing a digital one has been transmitted, a negative frequency offset equal to f_d and the threshold of the decision block placed at zero, $\Delta\psi = 0$ ($\Delta\phi = f_d\Delta T$) and the BER will approach 50%.

It is known that the optimum FSK detector is the *correlation detector*. Actually, this correlation detector is not widely used due to its high complexity. In [92] a new correlation type of detector has been proposed, which reduces the complexity of the system, avoiding to use the multipliers, which are replaced by simple XNOR logic ports. Nevertheless, no information is provided in the paper when the incoming data and the local oscillator at the receiver side are not perfectly synchronized in frequency. A block diagram of a correlation demodulator is shown Fig. 4.44.

It can be proven [96] that, in the presence of frequency offset, the detector matched to the incoming signal suffers from a signal attenuation equal to

$$\alpha = \frac{\sin^2 \pi \rho}{\pi^2 \rho^2} \quad (4.44)$$

where $\rho = f_{\text{off}}T_s$, with f_{off} the frequency offset and T_s the symbol period. When $\rho = 1$ the term in (4.44) goes to zero. Therefore, no more energy is collected by the detector matched to the incoming signal and the BER is close to 50%.

Therefore, the maximum residual offset should be less than a fraction of the data rate, which requires, for low data rate applications, an AFC system at the receiver side.

The *cross-differentiate multiply demodulator* has been largely used for FSK detectors in pager applications and a block diagram of a classical analog cross-differentiate and multiply detector is depicted in Fig. 4.45(b). Realizing the derivative of a signal and two multipliers in the analog domain can be quite hard bringing to large power consumption and increased hardware complexity.

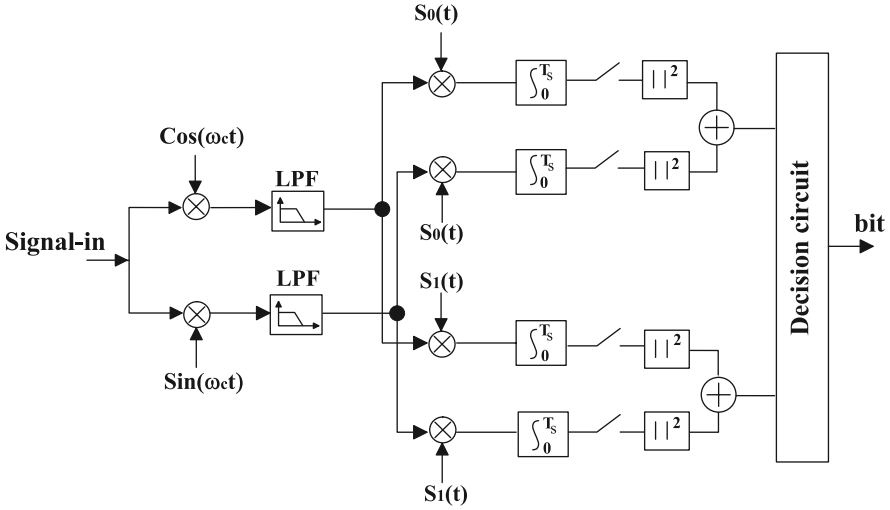
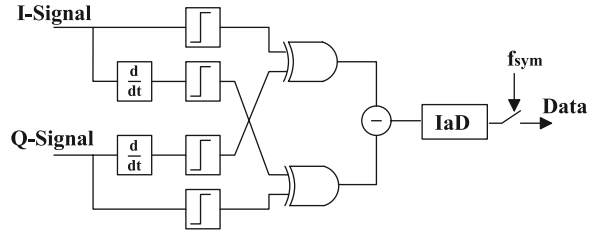
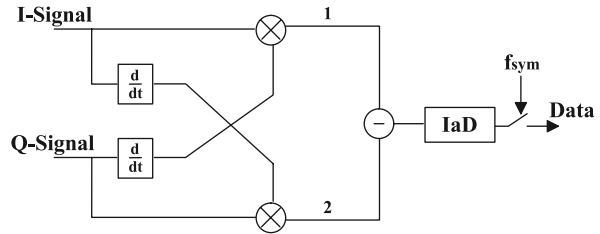


Fig. 4.44 Non-coherent correlator receiver

Fig. 4.45
Cross-differentiate-multiplier
demodulator

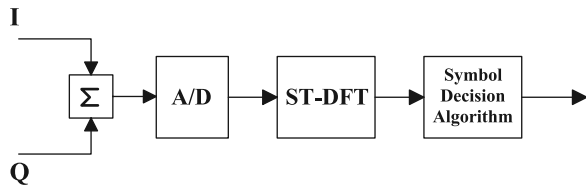


(a) DCDM demodulator



(b) Analog Cross-Differentiate Multiply (CDM) demodulator

Therefore, new topologies, still based on the cross-differentiate-multiply principle, have been developed in order to reduce the hardware complexity and therefore, the power consumption [93]. One of this topology is depicted in Fig. 4.45(a). Indeed, if a zero IF architecture is used, the *I* and *Q* signals can be easily digitized and the derivative can be obtained by simple subtraction between two consecutive samples. Furthermore, the complex analog multiplier can be replaced by a XOR logic port, which behaves as one bit multiplier.

Fig. 4.46 ST-DFT demodulator

Considering the simple analog representation given in Fig. 4.45(b), and considering for the I and Q signals the expressions given in (4.40), the signals at point 1 and 2 in Fig. 4.45(b), with a constant signal amplitude, are

$$\begin{cases} s_1(t) = \frac{A^2}{4} \cos^2[2\pi(f_d \int_{-\infty}^t m(\tau) d\tau + f_{\text{off}}t)] \\ \quad \times 2\pi(m(t)f_d + f_{\text{off}}) \\ s_2(t) = -\frac{A^2}{4} \sin^2[2\pi(f_d \int_{-\infty}^t m(\tau) d\tau + f_{\text{off}}t)] \\ \quad \times 2\pi(m(t)f_d + f_{\text{off}}) \end{cases} \quad (4.45)$$

From (4.45) it is clear that when $f_{\text{off}} = \pm f_d$ and $m(t) = \mp 1$, the BER should approach 50%. Even though (4.45) suggests the possibility of cancelling the frequency offset via proper encoding, it is easy to see that the presence of integrated noise (which has not been considered in (4.45)), will make the task not easy due to the fact that the output of the IaD stage will be a stochastic variable.

The *ST-DFT demodulator* is based on the Short-time DFT algorithm and has been proposed for Low-Earth Orbit (LEO) satellite communication systems [94].

A block diagram of a ST-DFT receiver is depicted in Fig. 4.46. The algorithm assumes differential encoding. Therefore, when the information bit “1” is transmitted, the modulation frequency is shifted with respect to the previous symbol. If the information bit “0” is transmitted the modulation frequency remains unchanged.

At the receiver side, supposing perfect synchronization between transmitter and receiver, ST-DFT is applied on N samples of the incoming signal. Due to the noise, different peaks will be present in the spectrum including the tone of the transmitted signal. Let us call $f_{\text{dec}}^{(m-1)}$ the decided frequency for the symbol $m-1$. Let us suppose that t peaks are present in the ST-DFT of the symbol m . Then, if the absolute value of the difference between the i -th peak and $f_{\text{dec}}^{(m-1)}$ is less than $2f_{\text{res}}$, where f_{res} is the fast Fourier Transform (FFT) resolution, the decision algorithm will output a logic zero. If the difference lies in the region $[2f_d - 2f_{\text{res}}, 2f_d + 2f_{\text{res}}]$ the decision circuit will output a logic one. If none of the previous conditions are met, the peak is disregarded as a noise peak and the algorithm is applied to the following peak.

It can be seen that this simple algorithm can easily track the offset. Indeed, supposing that the decided frequency at the $(m-1) \times T_s$ time is $f_{\text{dec}} = f_d + f_{\text{off}}$, where f_{off} is the offset frequency and f_d is the frequency deviation, then we have at the receiver side when the information bit “1” is transmitted:

$$|\Delta f| = |-f_d + f_{\text{off}} - (f_d + f_{\text{off}})| = 2f_d \quad (4.46)$$

where Δf is the frequency difference between the m and the $m - 1$ symbol. As can be seen from (4.46) the offset cancels out. When the information bit “0” is transmitted, the detected frequency difference is:

$$|\Delta f| = |f_d + f_{\text{off}} - (f_d + f_{\text{off}})| = 0 \quad (4.47)$$

Again the frequency offset cancels out.

For all the demodulator topologies a Simulink model has been realized (see Fig. 4.47) and the robustness against a static frequency offset has been simulated. For all the demodulators a data rate of 1 kbps has been chosen together with a modulation index $m = f(d)/D = 8$, where D is the data rate. Furthermore, because of the wideband FSK modulation employed the effect of the HPF is negligible and will be neglected in the following analysis.

Both the ADM and DCDM demodulators use the phase information embedded in the signal to recover the transmitted data. Indeed, the derivative of the phase of the incoming signal can be approximated by using the I and Q signals from the following equation:

$$\frac{d}{dt}(\phi(t)) \propto i(t) \frac{d}{dt}q(t) - q(t) \frac{d}{dt}i(t) \quad (4.48)$$

which is exactly the same operation performed in the DCDM demodulator (see Fig. 4.45(b)). It can be proven that these demodulators exhibit a threshold effect and therefore, the input SNR should be high enough to allow the demodulator to work above the threshold requiring data pre-filtering.

From Fig. 4.47(b) it can be seen that the correlation demodulator is quite weak when a frequency offset exists between the incoming signal and the local oscillator. Due to their similarities, the DCDM and the ADM demodulators exhibit similar performances. As can be seen from Fig. 4.47(b), the BER exceeds 5% when the offset exceeds 3 kHz. For these demodulator topologies an AFC system is required.

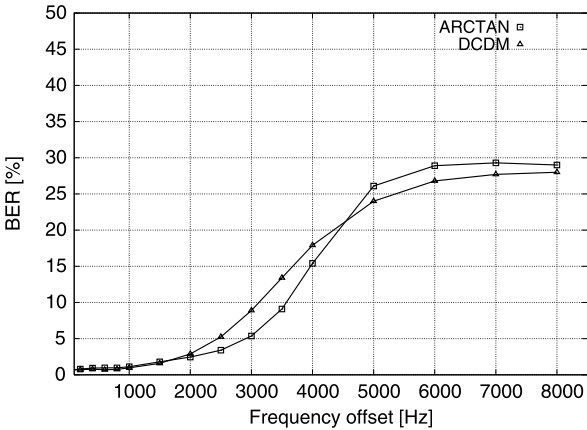
On the other hand looking at the performance of the ST-DFT demodulator it is possible to see its robustness against static frequency offset in a range that can be extended well above the 8 kHz shown in Fig. 4.47(c). In this way, frequency offsets induced by temperature variations and/or power supply variations can be easily tracked at the receiver side, keeping the transmitter as simple as possible. This directly translates in a higher transmitter power efficiency.

In the case of the ST-DFT demodulator, the FFT length has been chosen equal to 1024, the filters at the transmitter side and at the receiver side have been chosen to be a raised-cosine filter with roll-off factor equal to one. The window function is a Hanning window with a duration of two symbols.

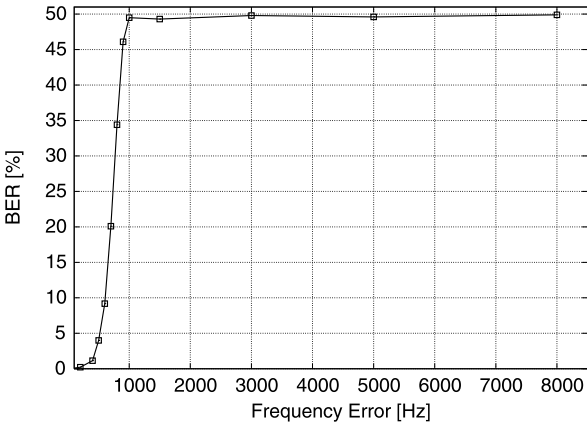
Finite Precision of the Pre-distortion Algorithm

As shown in [93], the capability of the ST-DFT algorithm to reject the frequency offset depends upon the condition that the offset is slowly varying. In other words,

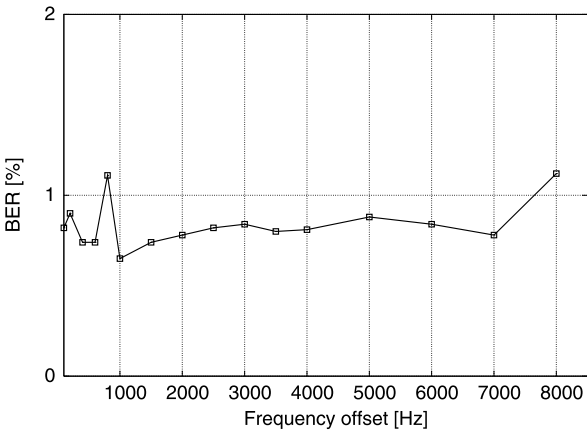
Fig. 4.47 BER vs. frequency offset with $E_b/N_0 = 12$ dB



(a) Arctan and DCDM.



(b) Correlation



(c) ST-DFT

the frequency offset should vary at a rate smaller than the data-rate. Therefore, the offset between two consecutive bits can be considered the same and it will be canceled out when differential encoding is applied. In this way also the frequency error due to temperature and power supply variations can be tracked and actively canceled out without requiring any additional circuitry.

While differential encoding can cancel out, when applied to two consecutive bits, frequency errors induced by temperature or power supply variations, it cannot cancel the frequency error induced by the statistical properties of the DAC and of the C - V varactor characteristic. Indeed, in the particular case of the proposed FHSS system, each bit is sent on a different channel, which is affected by a different offset due to the INL distribution of the DAC as well as the varactor C - V characteristic spread.

This means that two different bits will have two different offsets and, therefore, the simple differential encoding cannot cancel it out. Therefore, a straightforward way to cope with such a problem is to use a dedicated look-up table for each IC. Though this approach guarantees an easy solution to the previous problem, it will increase the cost of the final wireless node due to an increase in testing and calibration cost.

Another solution is to send information about the offset present at that particular hopping frequency when hopping begins. The principle is depicted in Fig. 4.48. In the figure the offset is sent as a high logic level but it can be chosen to be a low logic level as well. Due to the fact that the offset and the data are now sent on the same frequency bin, they are both affected by the same frequency error (due to the statistical properties of the DAC INL and the C - V characteristic of the varactor). The differential encoding, therefore, can provide in this case the final cancellation. The drawback of such a solution is a decrease in the effective data rate.¹³ Given the large modulation index employed ($m > 5$) and the low data-rate there will be no severe drawback at the receiver side. Therefore, this technique was chosen for the proposed implementation.

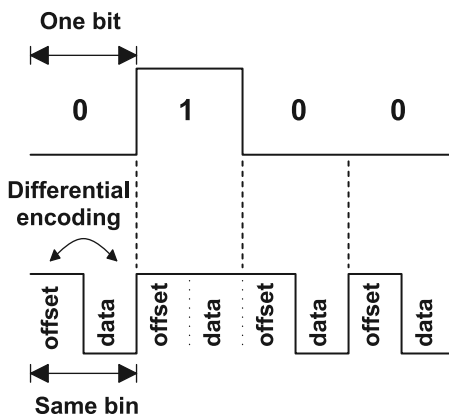


Fig. 4.48 Residual offset cancellation technique

¹³If we suppose that for 50% of the symbol period we transmit the offset information and for the remaining 50% we transmit the effective data, the effective data-rate is halved.

4.4 Implementation and Experimental Results

In this section the implementation and the measurement results for the TX-only prototype are discussed. The section is divided in two parts. The first part deals with the implementation of the ultra-low power TX node, while the second part deals with the realization of the RG. The section ends with a measurement summary of the ultra-low power TX-only node proposed in this chapter.

4.4.1 TX Node Implementation

The high-level block diagram of an ultra-low power TX prototype implementing the features described in the previous sections of this chapter is depicted in Fig. 4.49. The RF front-end is designed as an integrated circuit using the SOA technology. The antenna has been designed by TNO Industries and as a DSP processor the Microchip PIC16F627A [97] has been used.

The prototype should be able to harvest the required energy to transmit data from the ambient. In this specific case solar cells given by Philips Research are used but other sources of energy can be used as well. The energy harvested from the ambient in any form, will be stored in a ultra-flat battery developed by Front-edge Technology and discussed in Chap. 1.

The Application block in Fig. 4.49 senses the environment and produces a digital signal that can be further processed by the DSP. All the interfaces between the different prototype sub-blocks are also shown in Fig. 4.49. The application specific sub-block, the harvesting element and the storage sub-block and their interfaces are not discussed further in this section because they are outside the scope of this book. Therefore, the section focuses on the three remaining sub-blocks and their interfaces:

- DSP
- RF front-end
- Interface between the DSP and the RF front-end
- Antenna

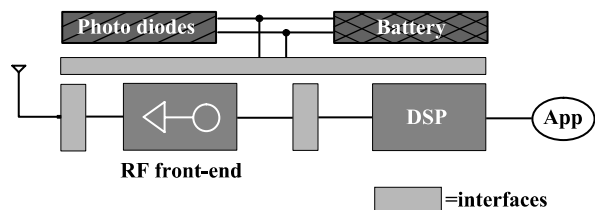


Fig. 4.49 High level block diagram of the TX-only demonstrator

DSP

The PIC16F627A microprocessor has a reduced set of instructions (only 35), which makes it easy to program and it features a 4 MHz internal factory-trimmed oscillator with 1% accuracy.

The stand-by current is lower than 100 nA at a supply voltage of 2.0 V while the supply current is around 120 μ A at 1 MHz operating frequency. The DSP wake-up time is lower than 4 μ s. These low power features made this microprocessor suitable for the TX-only prototype.

This is one of the many on-the-shelf microprocessors, which use low-power technologies in order to reduce the power consumption of the DSP. Moving from on-the-shelf products to a more customized DSP it is possible to consider the Coolflux DSP [98] from NXP semiconductor. The 65 nm CMOS implementation reduces the overall power consumption to about 25 μ W per MHz from a 0.9 V power supply. This product is available only as an Intellectual Property (IP) block and therefore, the Microchip product has been used in the prototype.

RF Front-End

As mentioned before in this chapter, two different RF front-ends have been implemented. The first prototype uses a more classical solution consisting of an oscillator working at twice the required frequency and a divide-by-two frequency divider in order to avoid frequency pulling. A block diagram of the first implemented SOA RF front-end, including the pin location is depicted in Fig. 4.50. The oscillator is an LC based oscillator using a differential coil to save space on chip.¹⁴ The oscillator is directly modulated by the signal coming from the DSP and processed by the interface between the DSP and the RF front-end. This interface implements the digital frequency pre-distortion concept and it is disclosed further in this section. Two frequency dividers are present. The output of the first frequency divider drives the antenna using the buffer disclosed in Sect. 4.2.4. The signal from the second divide-by-two frequency divider is output using the same buffer and it is used for the initial calibration using an FLL loop as described in Sect. 4.2. Both the buffers are 50 Ω matched for testability.

The silicon implementation of the RF front-end is shown in Fig. 4.51. As can be seen from the picture, most of the silicon area is used to implement the differential inductor used in the LC VCO. The area occupied by this prototype is around 6 mm².

The block diagram of the second implemented SOA RF front-end, including the pin location is depicted in Fig. 4.52. As it can be seen from the picture, the RF front-end is reduced to a single block and it has a reduced number of pins. The output does not require any additional buffer. The VCO provides itself the driving capability for the antenna. The coupling is obtained via an external balun, which acts also as a

¹⁴Compared to two coils.

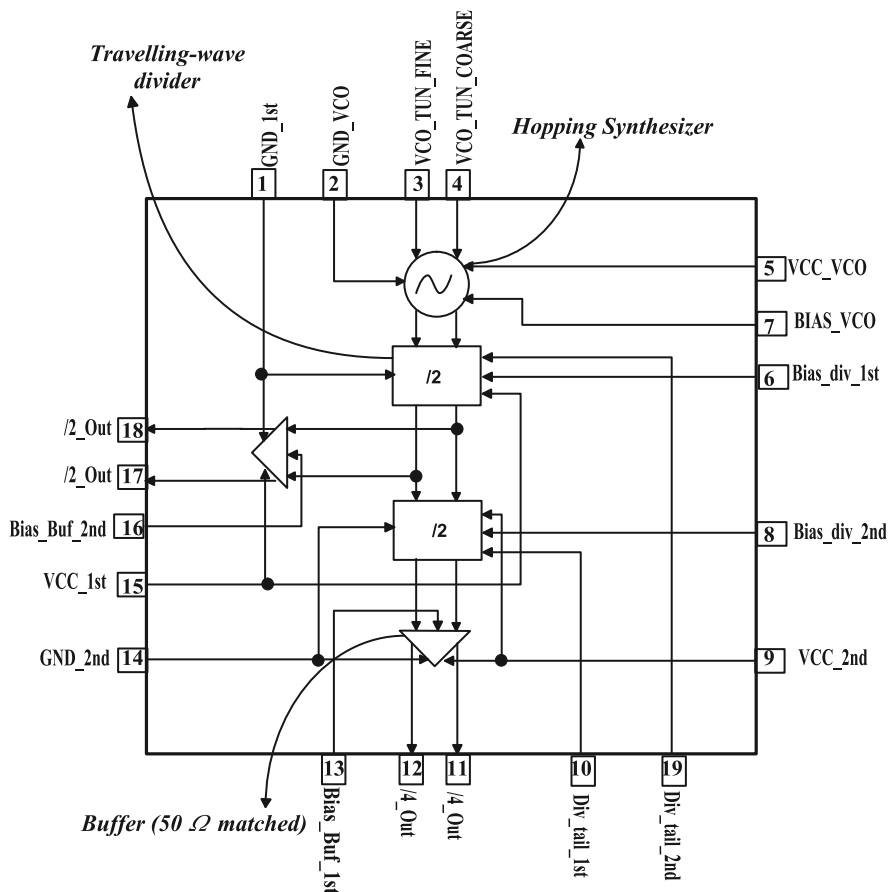


Fig. 4.50 Block diagram of the 1st IC prototype

differential to single ended converter. The single ended output of the balun is then connected to the antenna.

The realized RF front-end is depicted in Fig. 4.53. Differently from the first prototype it does not use a differential inductor. Indeed, the power VCO works at half the frequency used by the LC VCO in the 1st prototype and no differential inductor was available at that frequency. This translates in a much larger area, which can be avoided by designing a differential inductor optimized for the 915 MHz center frequency. The area occupied by the prototype is 3.6 mm².

Interface Between the DSP and the RF Front-End

The interface between the integrated RF front-end and the DSP core has been implemented using on the shelf components. The interface is the core of the digital

Fig. 4.51 1st IC prototype silicon implementation

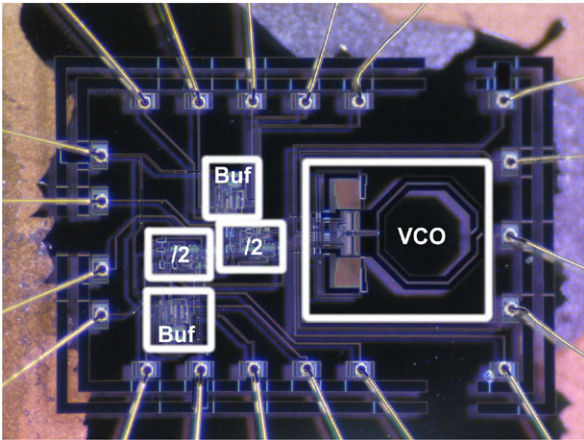


Fig. 4.52 Block diagram of the 2nd IC prototype

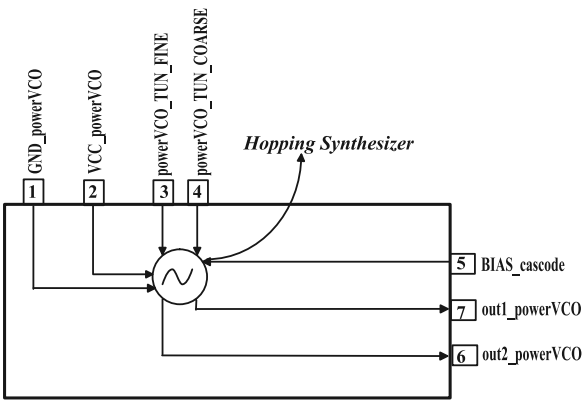
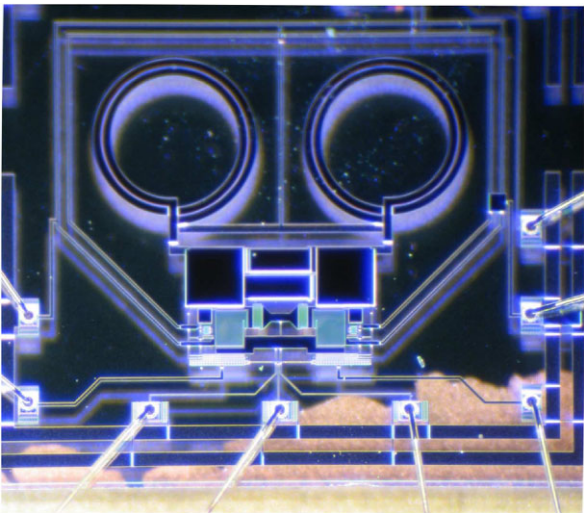


Fig. 4.53 2nd IC prototype silicon implementation



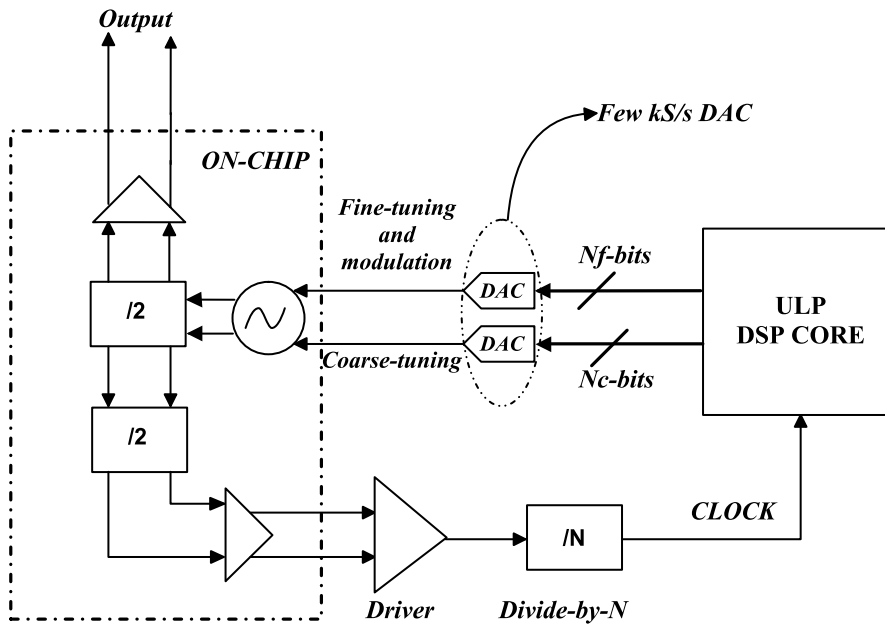


Fig. 4.54 Schematic block diagram of the transmitter prototype

frequency pre-distortion and it is mainly composed by a shift register and a DAC. The shift register is required in order to load the 12-bit word used in the DAC.

The microprocessor stores all the 12-bit words into a ROM and outputs those words as a serial bit stream at a frequency of 1000 words per second. The speed at which the words are output is the hopping speed. In this prototype a 1 khop/s has been chosen but this rate can be varied within a reasonable range without affecting the overall TX power consumption. Before the word can be loaded into the DAC, 12 parallel bits are required and therefore, they are temporary stored into the shift register.

A schematic block diagram of the RF front-end and DSP core including the interface between these two sub-blocks is shown in Fig. 4.54. The on-chip block depicted in the figure refers to the first prototype, but it can be easily substituted by the more simple (a single RF block) front end used in the second prototype. Two DACs are present. One DAC is used for the initial coarse calibration and it is switched off after the coarse frequency tuning has been achieved. The second DAC effectively implements the interface between the digital world (the DSP) and the analog world (the RF front-end). An external divide-by- N frequency divider is used to close the FLL loop together with the DSP and the internal divide-by-two frequency divider. The FLL loop assures the initial coarse frequency calibration and it is switched off after the coarse frequency tuning is achieved.

The fine tuning DAC needs to work at a very low frequency (few ksamples/s) and therefore, it does not pose any problem in terms of power consumption. In this prototype the AD7392 [99] has been chosen. It is a micropower DAC. This DAC has

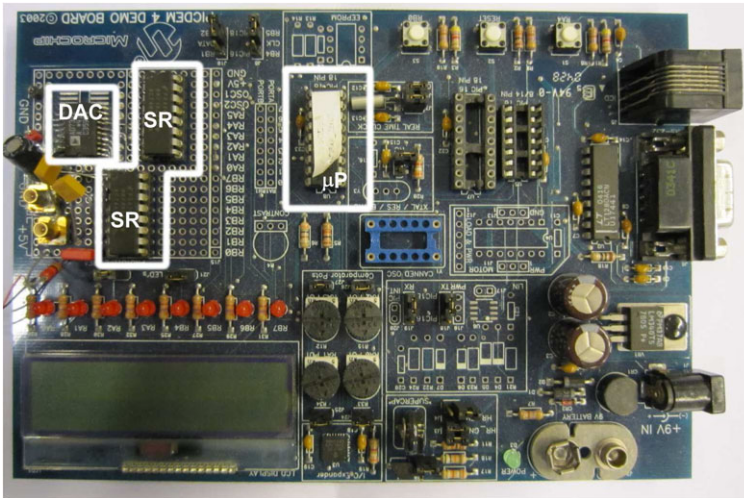


Fig. 4.55 PCB implementation of the DSP core plus the interface between the DSP core and the RF front-end

a maximum current consumption of only 100 μA and a 100 nA current consumption when it is shutdown. Moreover, it has a 2.7 V power supply and it has a maximum operating frequency of 17 ksample/s. For higher data rates (up to 125 ksample/s) the AD5341 [100] can be used for a slightly higher current consumption. The AD5341 consumes 115 μA from a 3 V power supply while consuming only 80 nA when it is shutdown.

The DSP core and the required interface between the DSP core and the RF front-end has been implemented on a dedicated board shown in Fig. 4.55. The μP is the microchip microprocessor, the SRs are the two 8-bit shift registers used¹⁵ and the DAC is the AD7392.

Antenna

The main characteristic an antenna should have to be used in an indoor environment is the omni-directionality. In the indoor environment a line of sight path is rarely present and therefore, the communication link should generally rely on reflected paths. For this prototype a very simple microstrip patch antenna has been used. In its most basic form, a microstrip patch antenna consists of a radiating patch on one side of a dielectric substrate, which has a ground plane on the other side as shown in Fig. 4.56.

For a rectangular patch, the length L of the patch is usually $0.3\lambda_0 < L < 0.5\lambda_0$, where λ_0 is the free-space wavelength. The patch is selected to be very thin such

¹⁵Only 12 out of 16 bits are used and after the 12 bits are loaded, a reset signal is used to reset the two registers for the next hopping bin synthesis.

Fig. 4.56 Structure of a microstrip patch antenna

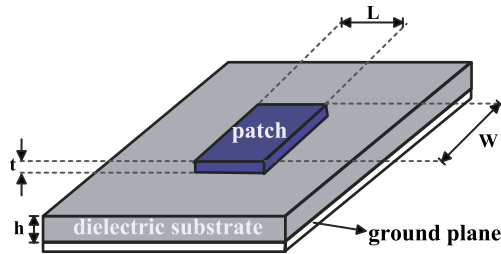
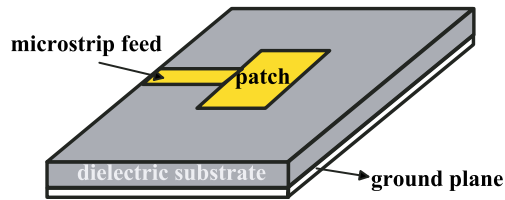


Fig. 4.57 Microstrip line feed



that $t \ll \lambda_0$ (where t is the patch thickness). The height h of the dielectric substrate is usually $0.003\lambda_0 < L < 0.05\lambda_0$. The dielectric constant of the substrate (ϵ_r) is typically in the range $2.2 < \epsilon_r < 12$.

We chose to focus on an optimal design in order to obtain a good efficiency, a large bandwidth and an efficient radiation. This choice requires the following characteristics:

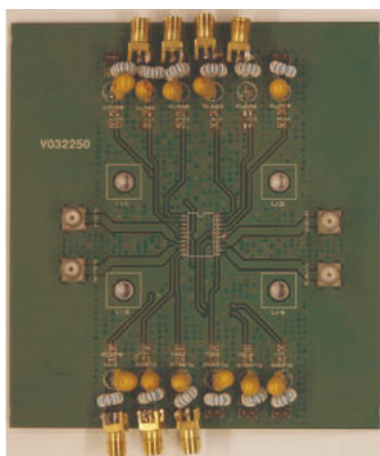
- Thick dielectric substrate
- Low dielectric constant

Unfortunately, this generally leads to a larger antenna. Though antenna integration is very important for an autonomous wireless node, antenna design is outside the scope of this book and therefore, we optimized the antenna in terms of bandwidth, efficiency and radiation regardless of the final dimension.

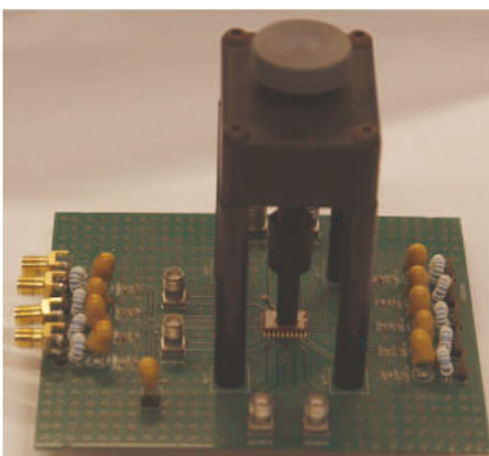
Microstrip patch antennas can be fed by a variety of methods. These methods can be classified into two categories: contacting and non-contacting. In the contacting method, the RF power is fed directly to the radiating patch using a connecting element such as a microstrip line. In the non-contacting scheme, electromagnetic field coupling is done to transfer power between the microstrip line and the radiating patch. The first method has been chosen for the antenna used in this prototype.

In this type of feed technique, a conducting strip is connected directly to the edge of the microstrip patch as shown in Fig. 4.57. The conducting strip is smaller in width as compared to the patch and this kind of feed arrangement has the advantage that the feed can be etched on the same substrate to provide a planar structure. The purpose of the inset cut (microstrip feed in Fig. 4.57) in the patch is to match the impedance of the feed line to the patch without the need for any additional matching element (50 Ω impedance matching is required in this particular prototype). This is achieved by properly controlling the inset position. Hence this is an easy feeding scheme, since it provides ease of fabrication and simplicity in modeling as well as impedance matching. The realized microstrip patch antenna is shown in Fig. 4.58.

Fig. 4.58 Microstrip patch antenna with microstrip line feed matched to $50\ \Omega$



(a) PCB layout for the power-VCO based IC prototype



(b) PCB including the IC and the tower for fast PVT measurements

Fig. 4.59 PCB used for PVT test for the power-VCO based IC prototype

IC Board for PVT Measurements

The PCB used to test the power-VCO based IC prototype¹⁶ is depicted in Fig. 4.59(a). As stated in Sect. 4.2.2, an important point in the measurement consisted in evaluating the results under process spread conditions. For this reason, it was mandatory to use only one PCB for all the ICs measured. In this way, the spread due to the PCB itself is cancelled and the measurement results show directly the performance variation of the prototype due to the process spread. To obtain this, a tower has been used, which can keep the IC prototype correctly aligned. Following

¹⁶The PCB for the LC-divider based prototype is very similar and therefore, it has not been shown in this book.

this procedure, the IC pins are correctly contacted to the PCB without requiring to be soldered and the PCB remains unchanged over the all set of measurements.

The PCB with the described tower is shown in Fig. 4.59(b). This measurement setup has been used to obtain all the measurement results in this section as well as for example the results shown in Fig. 4.8.

4.4.2 RG Implementation

The RG has been implemented using measurement instruments and a PC in order to apply all the algorithms required for the correct reception of data and described throughout this chapter. The schematic block diagram of the RG receiver is shown in Fig. 4.60.

The receiver antenna is the same antenna used at the transmitter side and described in Sect. 4.4.1. The downconversion part is realized via the National Instrument PXI-5600. It is composed of a local frequency reference and a downconversion mixer. After the incoming signal has been downconverted to baseband a high sampling rate digitizer is used to convert the incoming data into the digital domain. The sensitivity of the receiver chain is around -88 dBm.

The measurement results throughout this chapter has been obtained using the PXI-5620, which is a 12 bits, 64 Ms/s digitizer. After the data has been converted in the digital domain, a Labview based interface is used to collect the data. The final conversion to a demodulated bit stream using the ST-DFT algorithm described in Sect. 4.3.2 is not performed in real time because of data limit with the computer Peripheral Component Interconnect (PCI) interface. Therefore, the data is saved into an internal high speed memory buffer and when the available 16 Mbytes are filled, the measurement is stopped and the data is transferred for further analysis. The ST-DFT algorithm is then performed using Matlab and the received data stream is compared with the transmitted data stream to obtain the BER.

4.4.3 Measurement Results

In this section the measurement results are given. In general, both the RF front-ends can be used in the architecture described in this chapter and based on frequency

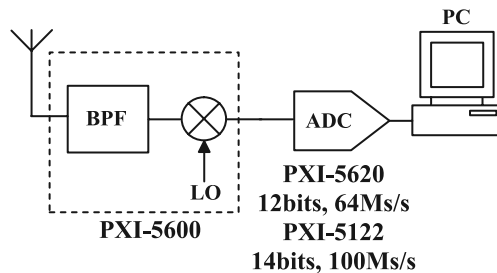


Fig. 4.60 RG receiver implemented using a PC and measurement instruments

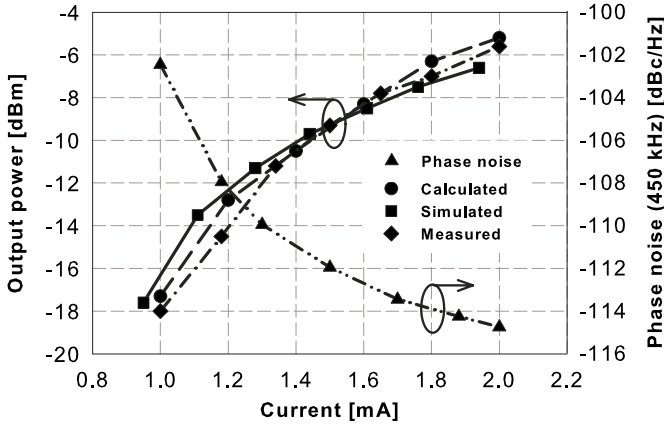


Fig. 4.61 Output power and phase noise at 450 kHz offset versus bias current for the power-VCO based RF front-end

pre-distortion. To avoid repetitions, the measurement results are referring to the power-VCO based prototype. The few differences in the measurement results are summarized in Table 4.11. These differences are mainly a wider tuning range for the VCO-divider based topology with respect to the power-VCO topology and the requirement for an output buffer, which is not present in the power-VCO based RF front-end.

For short-range wireless communication, power ranges between -25 dBm and -5 dBm can assure a robust communication over 10 meter distance in the indoor environment. From Fig. 4.61 it can be seen that the required bias current for the power-VCO based front-end ranges approximately between 1 mA and 2 mA to obtain an output power between -18 dBm and -5 dBm. Predicted, simulated and measured results show a good agreement allowing minimization of the current consumption for a given set of output power and phase noise specifications.

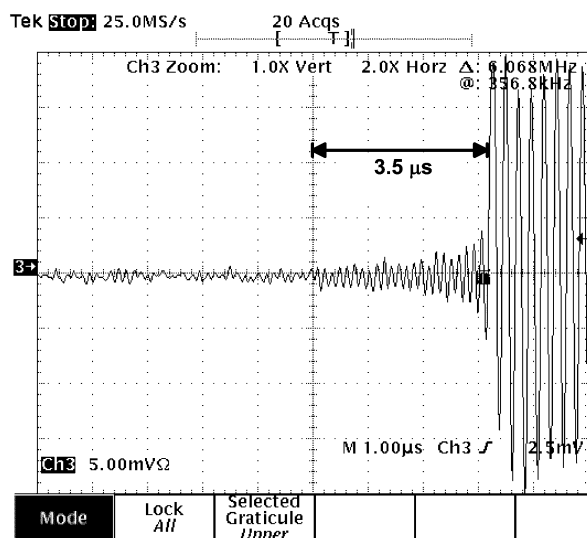
In Fig. 4.61 the phase noise varies between -102 dBc/Hz and -115 dBc/Hz when the bias current changes between 1 and 2 mA, which is in line with the specification given in Table 4.2.¹⁷ At the lower current consumption (1 mA), the output power on a common $50\ \Omega$ antenna is around -18 dBm (excluding cable losses).

Given the fact that the system has to be duty-cycled to reduce the average power consumption, it is important that VCO start-up time is small compared to the data transmission time. In Fig. 4.62 the measured oscillator start-up time is plotted.¹⁸ The start-up time is around $3.5\ \mu\text{s}$. Even considering a 10 kbps data rate, this is a

¹⁷The phase noise specification in Table 4.2 already includes a 3 dB margin to account for implementation losses. Therefore, it is possible to state that the -102 dBc/Hz exceeds by 2 dB the minimum phase noise requirement (see Sect. 4.2.3).

¹⁸The oscillator frequency in Fig. 4.62 is an alias of the synthesized frequency due to the sampling rate used in the measurement setup. Indeed, the 25 Ms/s rate is lower than required by the Nyquist theorem.

Fig. 4.62 Oscillator start-up time for the power VCO based architecture



small fraction of the transmission time and it contributes to minimize the TX node start-up time.

To demonstrate all the concepts, algorithms and techniques described in this chapter, a link between a TX node and an RG was created. The TX node used the power VCO based architecture as RF front-end, while the RG has been described in Sect. 4.60. The environment was a common office environment and the relative distance between the TX node and the RG was around 8 meters.

In Fig. 4.63 the signals throughout the whole chain, from the look-up table in the microprocessor up to the evenly spaced FH spectrum are shown. In the measurement results shown in Fig. 4.63, the 64 channels are addressed sequentially rather than in a pseudo-random fashion. This choice allows to see the predistorted DAC output.¹⁹ The DAC output directly modulates the VCO and as it can be seen from Fig. 4.63, an almost linear frequency grid has been obtained with a maximum inter-channel error smaller than 5 kHz.

The transmitted power was set to -18 dBm, and when board and cable losses are included, a -25 dBm radiated power is obtained. Under these conditions a BER measurement has been performed over 10000 transmitted bits at 1 kbps data rate and 1 khop/s hopping rate. The RG was not able to hop at 1 khop/s due to settling time limitations. A possible solution was to digitize the whole bandwidth. Unfortunately, digitizing the whole bandwidth caused a data rate above the handling capability of the PC PCI interface. For these reasons, the TX hopped among 6 channels over the

¹⁹Indeed, because all the channels are addressed in a sequential way, the DAC output should look like a voltage ramp. Due to the frequency pre-distortion, it is clearly visible in Fig. 4.63 that the DAC output has a curved shape instead of a ramp like shape.

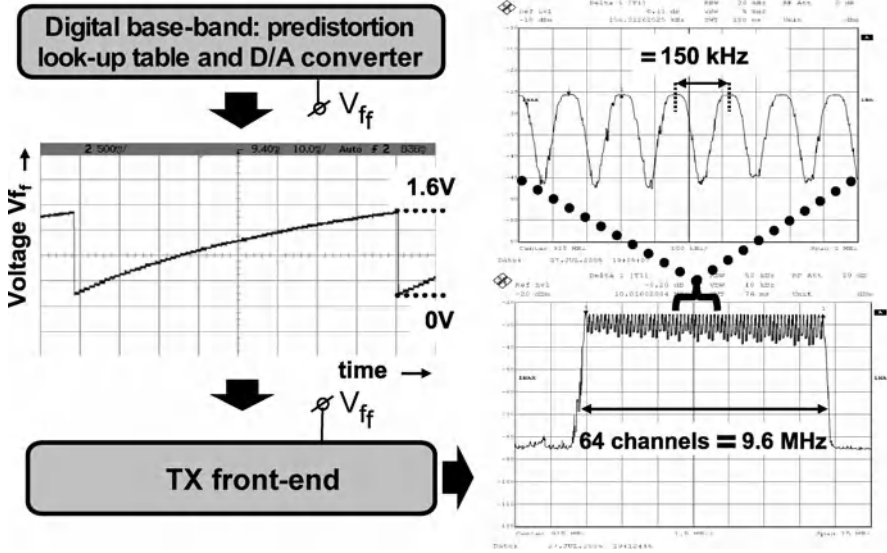


Fig. 4.63 Pre-distortion chain with measured DAC pre-distorted output voltage and output spectrum of TX front-end

64 it is possible to synthesize.²⁰ The receiver was not hopping but it digitized all the 6 channels at the same time. The resulting BER when the TX node and the RG are in a NLOS condition was measured to be smaller than 1.1%. As mentioned in Sect. 4.2.3, this BER is sufficient to obtain a reasonable QoS when FEC (Reed-Solomon codes for example) is applied.

The spectrum as received from the RG in real time before any demodulation is applied is shown in Fig. 4.64. The frequency axis is centered around 915 MHz ISM band center frequency. In reality, the spectrum shown has been already downconverted to an IF frequency by the RG. The frequency shown in the picture has been obtained by adding the LO frequency used for downconversion in the RG mixer. Also this measurement has been obtained by narrowing the possible hopping bins to 6 because of measurement instrument limitations. It should be noticed that when calibrated to the same output power the two described front-ends give the same output spectrum. Only the absolute position of the frequency peaks will vary due to the process spread on the varactor $C-V$ curves (DAC and board are kept the same).

The overall measured transmitter power consumption is 2.4 mW at -18 dBm output power and 4.4 mW at -5 dBm output power from a 2 V power supply for the power VCO based prototype. The LC VCO based prototype consumes 5.4 mW at -25 dBm transmitted power mainly due to the not power optimized output buffer as clearly visible in the power breakdown of Fig. 4.65. When the output buffers are

²⁰In practice this is a worst case condition because the system could benefit from a much smaller than designed frequency diversity.

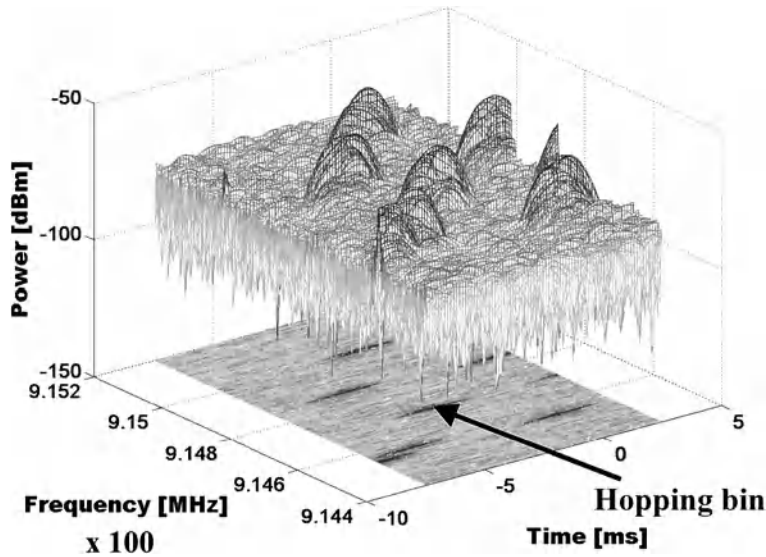
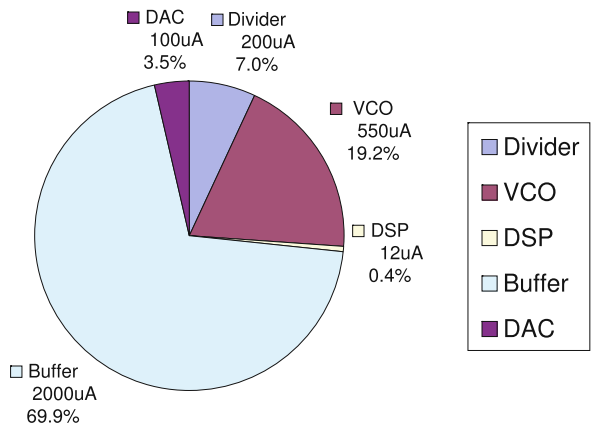


Fig. 4.64 Received power spectrum (8 meters distance, NLOS condition, and -25 dBm transmitted power)

Fig. 4.65 Power consumption breakdown by block for the LC-VCO-divider front-end (buffer not optimized for low-power)



not considered the hopping synthesizer ($\mu P + \text{DAC} + \text{VCO} + \text{divider}$) consumes $870 \mu\text{A}$ from a 1.8 V power supply (around 1.6 mW).

The baseband and the mixed signal circuitry including the DSP dissipate 0.4 mW , which is mostly consumed in the DACs. If the coarse calibration DAC (see Fig. 4.5) is switched off after the first initial calibration, the power consumption reduces to about 0.2 mW for both prototypes.

Table 4.11 RF front-ends performance summary

Parameter	FE1	FE2(1 mA)	FE2(2 mA)	Unit
Technology	SOA		SOA	
Active chip area	2.8		3.6	mm ²
V_{cc}	1.8		2	V
Max. Front-end current (20 samples)	2.8	1	2	mA
Max. Front-end dissipation (20 samples)	5	2	4	mW
Baseband dissipation			0.4	mW
Total active FHSS TX dissipation	5.4	2.4	4.4	mW
Min. coarse tuning range (20 samples)	94.4		50	MHz
Min. fine tuning range (20 samples)	8.9		10	MHz
Worst case phase noise (20 samples)	−109	−102	−115	dBc/Hz @ 450 kHz
Output power	−25	−18	−5	dBm
Raw BER (8 meter distance, $P_{out} = -25$ dBm)			<1.1E-2	

4.4.4 Benchmarking

A summary of the performance for both the prototypes, including the baseband circuitry is shown in Table 4.11. The measurement results confirmed the capability of the proposed architecture to minimize the complexity of the transmitter and therefore, to reduce the power consumption compared to the state-of-the-art FH transmitters, as shown in Table 4.12, by a factor 4.

It can be seen that a further reduction in the overall power consumption can be obtained by employing simpler modulation schemes together with dedicated technologies [14, 101] or by matching to high ohmic antennas [102] and using one or two channels communication link [14, 101–103]. In this work, the primary goal was the robustness of the wireless link together with the low power consumption, while employing standard technological solutions. In this sense the method compares favorably with all the state-of-the-art spread-spectrum based solutions, even in terms of efficiency.²¹

4.5 Conclusions

This chapter focused on an architecture suitable for one-way link communication. In this asymmetric scenario, the TX node is severely power constrained, while the RX node has, virtually, an unlimited power budget. Starting from this assumption, the

²¹In this book the transmitter efficiency is defined as the ratio between the radiated power and the overall transmitter power consumption.

Table 4.12 Comparison with state-of-the-art WSN transmitters

Ref.	Freq. [MHz]	Process	Architecture	Modulation	Data-rate [kbps]	$P_{\text{out}}^{\text{a}}$ [dBm]	P_{tx}^{b} [mW]	V_{supply} [V]	$\eta_{\text{tx}}^{\text{c}}$ [%]
[103]	2400	0.13 μm CMOS	One channel	FSK	300–500	−5	1.8	0.4	17.8
[104]	2400	0.18 μm CMOS	DSSS	BPSK	200	>0	18	1.8	5.6
[14]	1900	0.13 μm CMOS	Two channels	OOK	40	1.6	8.4	1.2	17.2
[102] ^d	900	0.25 μm CMOS	One channel	FSK	20	−6	1.3	3	19.2
[105]	900	0.18 μm CMOS	DSSS	BPSK	40	0	28.8	1.8	3.5
[101] ^e	1900	0.13 μm CMOS	One channel	OOK	50	0	1.6	0.3	31.5
[106]	2400	0.18 μm CMOS	DSSS	BPSK	200	1	42	1.8	3
This work	900	SOA Bipolar	FHSS	FSK	1–10	−5	4.4	2	7.4

^aRadiated power
^bTransmitter power
^cTransmitter efficiency
^dA 400 Ω antenna is used
^eFBAR resonators are used

novel architecture proposed in this chapter exploits the benefit of direct modulation of the VCO to synthesize the required hopping bins.

Non-idealities coming from the non-linear relation between frequency and varactor capacitance, the DAC, and other sources of non-linearities, are compensated via digital frequency pre-distortion. The remaining frequency error is corrected, together with any other frequency error caused by process spread, temperature and supply variations, at the receiver side via a set of dedicated algorithms. In this way, the TX node is kept simple, thus minimizing its power consumption. The hopping speed can be varied without increasing, in a first approximation, the overall power consumption. This allows to trade the transmitted power for hopping speed without any power consumption penalty.

Two RF front-ends are proposed. The first front-end employs a classical LC oscillator followed by a frequency divider. The second front-end combines the up-conversion with the power amplification using a power-VCO based on a Colpitts oscillator. Both in theory and in measurements the second front-end outperforms the first implementation, which also requires an output buffer to drive the antenna. The RF front-end exploiting the power-VCO concept uses only 1 mA from a 1.8 V supply. The output power is only -25 dBm including cable and antenna losses but this is sufficient to get a BER smaller than 1.1% in a real indoor environment with antennas 8 meters apart and in a NLOS condition. It has also been proven that this BER is sufficient for a reasonable QoS when FEC is used and in particular Reed-Solomon codes are implemented.

Chapter 5

A Two-Way Link Transceiver Design

In Chap. 1 the space of applications for an autonomous wireless sensor network has been divided in applications that require only to transmit data and applications that require to transmit and receive data. The first ones only require a transmitter, while an RG is used to receive the data and eventually to distribute the information using a wireless or a wired communication link. The second ones require a two-way link. In this category also fall all the networks that require a multi-hop based protocol to function properly.¹ This second case presents a much higher degree of complexity because both the receiver and the transmitter must be designed to be ultra-low power.

The first consequence of this constraint is that it is not possible to implement complex digital algorithms to compensate for system non-idealities and therefore, accurate references are mandatory. The first system choice we adopted is, therefore, to use a crystal based reference. Though this increases the wireless node cost, it allows to avoid complex frequency recovery algorithms minimizing, in this way, the node power consumption. Moreover, though the cost per node increases, no RG is required. This means that the cost of the RG used in the one-way link can be spread among several nodes. Therefore, though a two-way link network will probably cost more than a one-way link network, the cost difference can be reduced because no infrastructure is required for this kind of network.

This chapter focuses on the design of a two-way link transceiver. First some general design guidelines are given. Second the transmitter architecture is described and an ultra-low power fast hopping frequency synthesizer is proposed. Then, the receiver system analysis is described and simulation and measurement results are given. Finally, conclusions are drawn and an outlook to a possible future work is given.

¹For example, in a remote area it can be very difficult to deploy an RG or any other kind of infrastructure to support a one-way link network. In this case, if the network is deployed over a wide space, a multi-hop technique must be used in order to avoid the wireless node to transmit over a very long distance, which will translate in an exponential increase of the transmitted power and therefore, of the overall power consumption of the wireless node.

5.1 Transmitter Design General Guidelines

A two-way link requires any node in the network to be able to transmit as well as to receive data. A two-way link network, though it does not require any infrastructure to operate, does present a big challenge when its power consumption is reduced to a level that can allow autonomous operation. The reason resides on the fact that more hardware is required (a transmitter and a receiver) and this hardware is used more often than in a one-way link network especially if a multi-hop protocol is adopted.

In a one-way link network we have shifted most of the complexity from the transmitter to the receiver. The reason is that the receiver is mains supplied and therefore, it could be designed to handle complex digital functions able to correct for the non-idealities present in the transmitter. This allows to reduce the transmitter complexity and, therefore, to reduce the transmitter power consumption. Unfortunately, in a two-way link both the transmitter and the receiver functions are power constrained. The first outcome of this different scenario is that, while in the one-way link network it was possible to remove the crystal, in a two way link this is not anymore possible. Indeed, in a one-way link network the required frequency synchronization between the transmitter and the receiver is mainly performed at the receiver side by the RG. While this is surely possible in a system with no power constraints like the RG, it is very difficult to achieve in the two-way link network.

In Chap. 3 we have shown that at low transmitted power (below 0 dBm) most of the power is used in the pre-PA part of the frequency hopping transmitter. We have also pointed out that most of this power is used in the hopping synthesizer but no evidence for this statement has been given up to this point. For this reason, in Table 5.1 4 examples of frequency hopping transmitters are listed. It can be seen from Table 5.1 that when the transmitter only part is considered, 40% or more of the power consumption is used in the frequency hopping synthesizer. This percentage drops down to more than 20% where both functions (RX and TX) are used. Generally, it is possible to assume that every wireless node in the network works in half duplex mode and therefore, almost half of the total power consumption during active mode is used to power the frequency hopping synthesizer. Therefore, the design of a truly autonomous frequency hopping transceiver needs to reduce the synthesizer power consumption an order of magnitude below the current state-of-the-art.

Table 5.1 Power consumption of some commercially available FHSS products

Product	Freq. band [MHz]	Power consumption TX [mW]	Synthesizer power consumption [mW]
CC1100 [107]	915	24.3	14.8
CC2500 [108]	2400	27	13.3
ADF7020 [109]	915	36.6	23 ^a
SX1223 [110]	915	28	10

^aDeduced from AD synthesizer products

5.2 Transmitter Architecture

In order to achieve this target, it is important to minimize the range of frequencies that the synthesizer must synthesize as well as its maximum operating frequency. For this reason, the choice of the transmitter architecture becomes fundamental. In this book, we adopted a direct SSB upconversion architecture. This architecture is shown in Fig. 5.1. This transmitter architecture has been proposed for frequency hopping systems in [35]. The benefits of this architecture are:

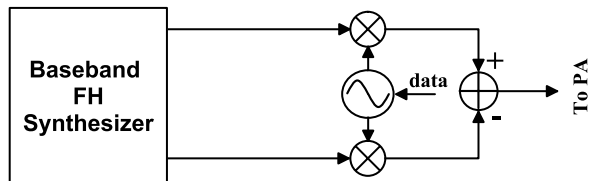
- Up-conversion is performed via a quadrature modulator so no closed loop system is used
- SSB up-conversion scheme halves the number of channels to synthesize at base-band

The absence of a closed loop allows to achieve an almost instantaneous upconversion of the baseband signal to the wanted band center frequency.² Halving the number of required bins to synthesize reduces the maximum operating frequency of the baseband synthesizer and therefore, its power consumption. The data is applied on the local oscillator used to upconvert the signal to the band center frequency. By shifting the LO frequency by an amount equal to the frequency deviation³ it is possible to upconvert the modulated data to the wanted frequency bin.

To avoid degradation of the BER at the receiver side especially when the SNR is poor, the local oscillator must have a good frequency stability. Good frequency stability means that the frequency deviation from the ideal center frequency under temperature and supply variations must be smaller than the instantaneous data rate. This can be obtained using a crystal reference and a PLL loop. Though feasible, this architecture can be costly in terms of power consumption especially if a frequency deviation in the order of few hundreds of kilohertz is considered.

Given the low data rate generated by any wireless node of the network, a simpler and more optimal solution in terms of power is to use a SAW based oscillator. If the oscillation frequency is changed between two values close to the resonant frequency of the SAW based resonator, an FSK modulator is obtained. Moreover, if the frequency deviation is much smaller than the SAW resonant frequency all the

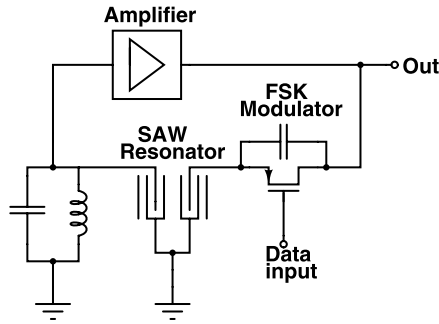
Fig. 5.1 Transmitter architecture employing SSB direct-up-conversion scheme



²In a closed loop system like a PLL the required frequency is synthesized after a certain amount of time called lock-in time. Generally, the lock-in time depends on the accuracy required in the synthesis of the frequency. If a part per million accuracy is required, the lock-in time can be large limiting the maximum achievable hopping speed.

³We suppose to use an FSK modulation.

Fig. 5.2 Block diagram of the Surface Acoustic Wave (SAW) stabilized LO with embedded FSK modulation functionality



stability properties of the SAW resonator do not change and therefore, a stable and accurate FSK modulator can be obtained.

A simple block diagram of a SAW based oscillator able to generate two tones at a slight different frequency around the band center frequency is depicted in Fig. 5.2 [111]. The FSK modulation is achieved via the parallel combination of a capacitor and a MOS switch. When the switch is turned on and off depending on the data input, different phase shifts are introduced in the feedback loop. In this way, the oscillator outputs two different frequency around the self resonant frequency of the SAW resonator. The most important advantages of this topology (MOS switch and parallel capacitance) are the following:

- No current consumption
- Low area
- Low complexity

The switch requires a special attention in the design. Indeed if not properly designed, it can degrade the phase noise performances considerably. Other possible alternatives to a SAW stabilized oscillator consist in replacing the SAW resonator with an FBAR [112, 113] resonator, which has, generally, a better form factor, or with a MEMS based oscillator [14].

In Chap. 3 and in the previous part of this chapter, we have shown that the bottleneck in the design of an autonomous frequency hopping transceiver is the frequency hopping synthesizer. The next section describes a sub-mW baseband fast frequency hopping synthesizer suitable for integration in an autonomous frequency hopping transceiver for a two-way link scenario.

5.3 Synthesizer Design

The frequency hopping synthesizer is the core of any frequency hopping system. In Chap. 3 we have shown that the design of an agile sub-mW frequency hopping synthesizer is not a trivial task and that common architectures like PLL and DDS fail to achieve both targets at the same time. In Chap. 4 we have shown that in a one-way link scenario the two targets can be obtained by direct modulation of a VCO and by

using calibration algorithms at the receiver side in order to recover from the non-idealities of the silicon implementation. The use of those algorithms was possible because the receiver was mains supplied and therefore a virtually unlimited power budget could be used. This hypothesis, unfortunately, is not true for a two-way link scenario. Therefore, a new strategy has been used in order to design a frequency hopping synthesizer that is, at the same time, agile and low power.

5.3.1 Baseband Frequency Hopping Synthesizer Specifications

The first step in the synthesizer design consists in defining the required specs in terms of hopping speed, accuracy, SFDR and power consumption. These specifications, explained in the sequel, are the following:

- Power consumption: smaller than 0.5 mW
- Fast hopping ($\gg$ data-rate)
- Good SFDR: > 30 dB
- Accuracy: within a fraction of the data rate
- Small channel spacing: < 1 MHz

As already mentioned in the previous chapters, the power consumption of the frequency synthesizer is part of the pre-PA power and in general, can be considered as a “wasted” power. Minimizing the synthesizer power consumption optimizes, in terms of efficiency, both the transmitter and the receiver. In the transmitter part, the PA power becomes dominant while at the receiver side the LNA-Mixer-LPF-ADC path dominates. Given the fact that a complete transceiver power consumption should not exceed a few milliwatts, we set the baseband frequency hopping synthesizer power consumption to be below 0.5 mW.

As shown in Chap. 2, fast hopping is required in order to increase the system immunity to channel non idealities like fading. Indeed, if the same bit is transmitted on several channels, the probability of a bit error, when a simple majority criteria is used to reconstruct the transmit data, will be greatly reduced. Therefore, though the instantaneous SNR will be decreased due to fading, this fast hopping technique allows to recover from this degradation without increasing the transmitted power. It is important that while increasing the hopping speed the power consumption of the synthesizer remains roughly independent of the hopping speed.

SFDR is an important requirement in order to avoid jamming of many communication channels caused by spurs of the transmitted signal. If the spur levels are too high, a channel used by another transmitter communicating at the same time⁴ can be unusable and this decreases the effectiveness of the frequency hopping scheme against interferers. It has been shown in [43] that an SFDR of 30 dB or larger is

⁴This transmitter can be further away from the receiver with respect to the jamming node. Therefore spurs of the jamming node can effectively overwhelm the signal received from the wanted transmitter.

enough to avoid such a situation in wireless sensor networks. We designed the system for a 40 dB SFDR in order to have some margin over degradations due to silicon implementation.

Frequency accuracy is an important parameter in order to avoid that complexity increases considerably at the receiver side. Given the fact that the system employs a simple BFSK modulation scheme, any frequency inaccuracy at the transmitter or the receiver side can cause an increase in the BER. In Chap. 4 some demodulator topologies have been analyzed focusing on their sensitivity to residual frequency offset between transmitter and receiver. From that analysis, in order to avoid any limitation in the choice of the frequency demodulator, an accuracy better than a fraction⁵ of the data rate is required.

Finally, a small channel spacing is required in order to pack as many channels as possible within the wanted band. Indeed, the advantages given by the frequency spread spectrum technique against channel non idealities and interferers is proportional to the number of channels used. The larger the number of channels, the more robust is the communication link even in the presence of strong non-idealities. A channel spacing below 1 MHz allows to have more than 83 channels in the 2.4 GHz ISM band, which allows to have a fairly large processing gain.

5.3.2 Baseband Frequency-Hopping Synthesizer Architecture

In a DDS architecture accuracy and agility are guaranteed by the use of an accurate crystal based reference signal and by a feedforward architecture. The proposed architecture keeps these two characteristics of a DDS while trading flexibility for power. Indeed, a DDS is able to generate a wide range of frequencies with a very fine and accurate frequency step that in some architectures can reach a few Hertz. The proposed two-way link wireless sensor node requires to generate only a relatively small set of frequencies with a resolution of 0.5 MHz. The proposed architecture, therefore, has a lower flexibility compared to a DDS architecture but its power consumption is an order of magnitude lower, and therefore it can be integrated in an autonomous wireless node for a two-way link network.

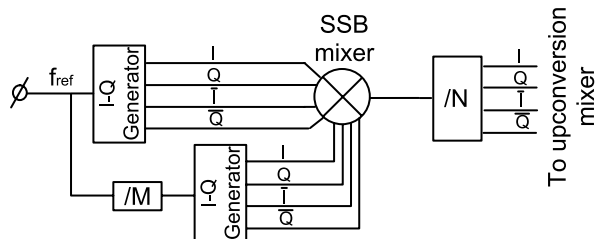
The proposed solution to have an agile, accurate but still ultra-low power frequency hopping synthesizer is to synthesize the required frequency bins according to (5.1) [114].

$$f_{\text{bin}} = f_{\text{ref}} \times \left(1 \pm \frac{1}{M} \right) \times \frac{1}{N} \quad (5.1)$$

where f_{ref} is the stable and accurate reference signal, f_{bin} is the required frequency bin and M and N are integer numbers. It is important to notice that the reference

⁵In Sect. 4.3.2 we shown that the correlation based demodulator is the most sensitive frequency demodulator architecture to frequency offset. With a fraction of the data-rate we intend that the frequency offset has to be below roughly one fifth of the data-rate for an acceptable BER.

Fig. 5.3 Simplified block diagram of the proposed baseband frequency hopping synthesizer



signal is constant for all the required values of f_{bin} . Therefore, all the required bins must be synthesized by changing the integer numbers N and M .

Looking at (5.1), it can be seen that the frequency synthesizer can be implemented by using a SSB mixer with the capability to choose between the two sidebands and two programmable frequency dividers. Given the fact that all the required frequencies are synthesized from a crystal locked reference signal, all the synthesized bins will hold the stability and accuracy properties of the reference signal (e.g. in the order of parts per million).

Equation (5.1) has several possible solutions. The implemented solution uses an $f_{\text{ref}} = 288$ MHz for an inter-channel spacing of 0.5 MHz. The 288 MHz reference signal can be generated in different ways which are discussed more in detail in Sect. 5.4.

Table 5.2 shows the required M and N factors to obtain all the frequency bins in the range between 0.5 MHz and 14 MHz with a step of 0.5 MHz. Some baseband frequency bins (4, 8 and 12 MHz in Table 5.2) can be obtained by direct division of the reference input by N . The total number of baseband channels is 28 and given the SSB up-conversion scheme employed it is possible to synthesize 56 channels around the desired band center frequency.

A simplified block diagram of the proposed architecture is depicted in Fig. 5.3. Given the fact that the architecture is feed-forward and no feedback loop is present, its settling time is set by the internal nodes time constants. This means that the proposed architecture is suitable for a fast hopping frequency synthesizer. Moreover, if the complete architecture is implemented using digital blocks its power consumption can be reduced below 0.5 mW.

As it can be seen from Table 5.2, all the required N and M division ratios are multiples of two and three. Therefore, only divide-by-two and divide-by-three frequency dividers are required. If implemented as analog dividers, they can consume a relatively large power. Therefore, they need to be implemented digitally. Moreover, an SSB mixer requires quadrature signals. This means that low-power quadrature generators are required (see Fig. 5.3). The baseband synthesizer must also be able to generate quadrature signals in order to support the SSB up-conversion scheme as shown in Fig. 5.1.

The easiest way to generate quadrature signals is to use a divide-by-two frequency divider. A low-power implementation for a generic divide-by- k (where k can be any integer number) frequency divider uses digital Flip-Flop (FF). While a fully digital implementation is optimal for the power consumption, it has several

Table 5.2 One possible solution for (5.1)

Freq. bin [MHz]	M	N	Sideband ^a
0.5	4	144	L
1	4	72	L
1.5	4	48	L
2	4	36	L
2.5	12	48	L
3	4	24	L
3.5	16	36	L
4		$\frac{f_{\text{ref}}}{72}$ ^b	
4.5	4	16	L
5	12	24	L
5.5	24	24	L
6	4	12	L
6.5	24	24	U
7	12	24	U
7.5	12	16	L
8		$\frac{f_{\text{ref}}}{36}$ ^b	
8.5	36	16	L
9	8	12	L
9.5	36	16	U
10	12	12	L
10.5	16	12	L
11	24	12	L
11.5	48	12	L
12		$\frac{f_{\text{ref}}}{24}$ ^b	
12.5	48	12	U
13	24	12	U
13.5	16	12	U
14	12	12	U

^aL = Lower sideband and
U = Upper sideband

^bSSB mixer not used

drawbacks in terms of SFDR because of the high harmonic content of the signals involved in the SSB mixing process (see Fig. 5.3).

The analog signals used are square-waves that contain, besides the fundamental component, all the odd harmonics of the fundamental frequency. The third harmonic, for example, is only 9.5 dB below the fundamental and very close to it. This means that, after SSB upconversion to the center of the ISM band, it is very difficult to filter out this component. A pre-filter is required, but given the fact that the third harmonic is very close to the fundamental, this can be quite demanding in terms of complexity and power. A 30 dB extra rejection on the third harmonic would require a 3rd order filter which can be power hungry.

Also the level of the harmonics before the divide-by- N frequency divider must be well controlled. Indeed, also these harmonics are transferred to the output of the frequency synthesizer and can potentially reduce the synthesizer SFDR. In order to alleviate the drawbacks coming from a mostly digital implementation of this architecture, a selective rejection of close-in harmonics is adopted [115].

Harmonic Rejection Based Frequency Synthesizer

A harmonic rejection circuit is, in general, a circuit which passes some frequencies unattenuated, while it attenuates some other frequencies. Therefore, a filter can be also considered a harmonic rejection circuit. The major drawback of an analog filter is that its complexity and therefore, its power consumption increases by increasing the filter order.

If we suppose for a moment that it is possible to implement a harmonic rejection of the 3rd and the 5th harmonics of a wanted tone, it is instructive to see if removing these harmonics is sufficient to meet the required 40 dB SFDR.

Figure 5.4 shows the proposed architecture without any harmonic rejection applied. The level of harmonics at various points of the architecture are also shown. Because no harmonic rejection is applied in point A and B we have the fundamental components ($f_{\text{ref}}/2$ and f_{ref}/M) but also their odd harmonics. The largest harmonic is the third harmonic, which is only 9.5 dB below the fundamental. The impact of these harmonics on the architecture is different depending on if it is the third harmonic of the $f_{\text{ref}}/2$ signal or the third harmonic of the f_{ref}/M signal. The latter is much more difficult to attenuate after SSB mixing is performed because M can be large (see Table 5.2). A large M can produce a tone, which after SSB mixing, is very closely spaced around the fundamental of the $f_{\text{ref}}/2$ signal. For example, if we suppose that the fundamental of the signal in B is 1 MHz, its third harmonic is at 3 MHz. This harmonic mixes with the fundamental tone in A (144 MHz) creating a tone at 141 MHz only 9.5 dB below the wanted tone ($144 + 1 \text{ MHz} = 145 \text{ MHz}$ when the upper sideband is selected). If we want to attenuate this tone by 30 dB, the required analog filter should have several poles.

Of course, the third harmonic of the tone in A is also only 9.5 dB below the fundamental. But this harmonic is generally placed 3 times further away from the wanted tone. Nevertheless, a certain degree of attenuation is required for the third harmonic of the tone in A if no harmonic rejection is used for it.

To estimate the level of the harmonics at the synthesizer output, it is necessary to consider the effect of a digital divider on a spectrum having a large fundamental tone (at $f_{\text{ref}}/2 + f_{\text{ref}}/M$) and several other tones 9.5 dB or more below the fundamental. The effect of a frequency divider on an input signal which contains a certain number of spurs is analyzed in [116]. If a spur is present in the input signal at a frequency separation equal to f_1 from the wanted signal, after an integer division by a factor N , the output spectrum contains the same spur around each odd harmonic of the wanted signal at the same frequency separation $\pm f_1$. The amplitude of those spurs will be reduced, with respect to their input amplitude by a factor equal to $2N$. This can

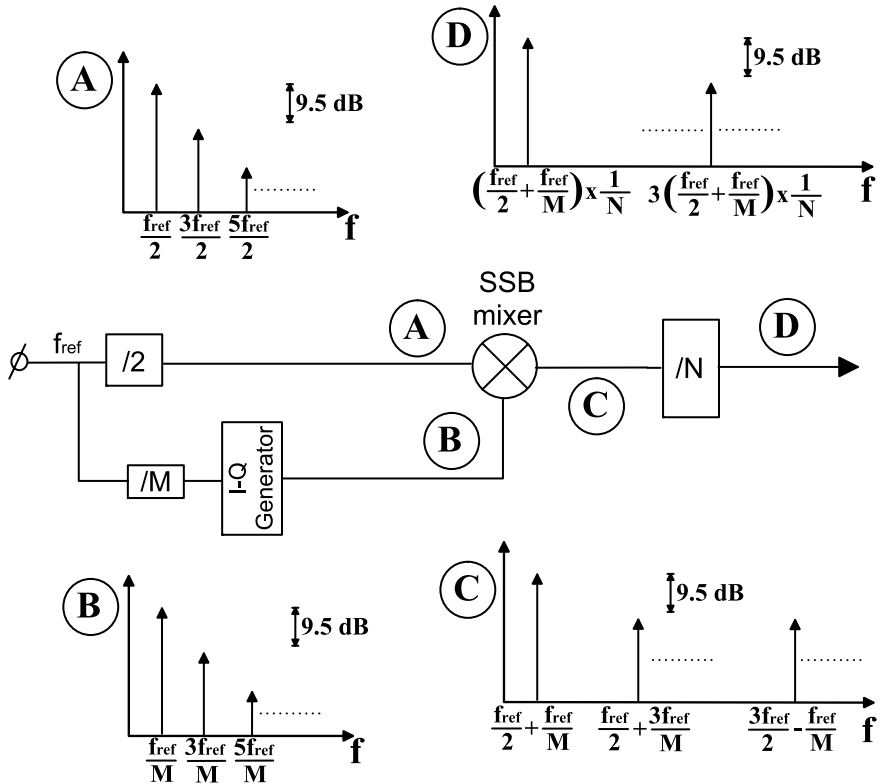


Fig. 5.4 Harmonics level for the proposed synthesizer when no harmonic rejection is applied

be intuitively explained by considering the spurs around the fundamental tone like phase noise. The factor 2 accounts for the fact that the spurs come from a SSB mixing and therefore, going through the divider the amplitude content is lost.

If we look at Table 5.2 we can see that the minimum value for N is 12 giving an extra 27.6 dB ($24\times$) attenuation for the unwanted spurs. This is sufficient to attenuate the third harmonic of the $f_{\text{ref}}/2$ signal (see the signal at point A in Fig. 5.4) by more than 40 dB (and therefore, more than the required specification) using a simple LPF placed after the SSB mixer (see Fig. 5.5).

The 3rd and 5th harmonics of the fundamental tone in B are removed using Walsh function based harmonic rejection as proposed in [115]. The 7th and 9th harmonics of the fundamental tone in B are 17 dB and 19 dB below the fundamental and will therefore be more than 40 dB below the fundamental tone after the divide-by- N block (see Fig. 5.5).

After the divide-by- N block, at point D in Fig. 5.4 the third harmonic of the wanted tone ($(f_{\text{ref}}/2 + f_{\text{ref}}/M) \times 1/N$) is just 9.5 dB below the fundamental if no harmonic rejection is applied. This is a situation similar to the case of the signal in point B. Indeed, after the SSB up-conversion to the center frequency of the ISM band, the third harmonic is very close to the wanted tone requiring a high order filter

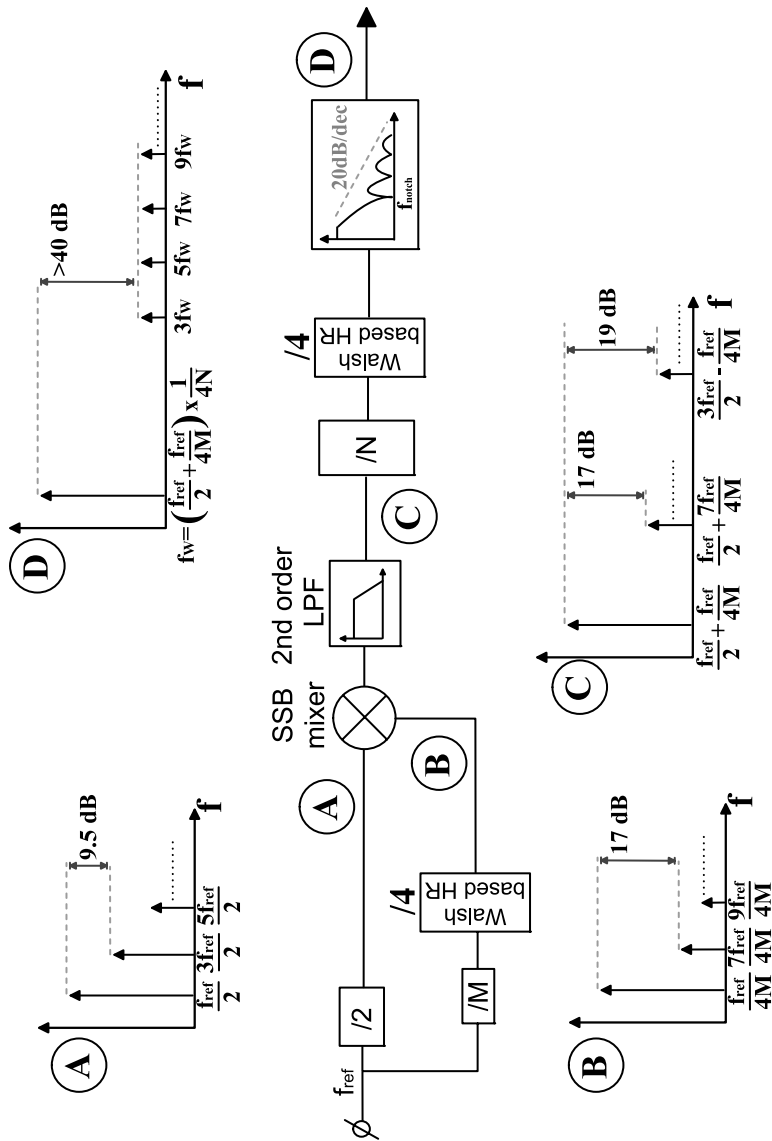


Fig. 5.5 Harmonics level for the proposed synthesizer including harmonic rejection and filters

to obtain at least a 30 dB extra rejection. Therefore, after the divide-by- N block, another harmonic rejection block based on the Walsh functions is used to remove the 3rd and the 5th harmonics of the wanted tone.

After removing the third and the fifth harmonics of the wanted tone in D, the two largest remaining tones are the 7th and 9th harmonic, 17 dB and 19 dB below the fundamental respectively. These two tones are still too close to the wanted tone after the final SSB up-conversion and must, therefore, be removed. The task is now easier, compared to the case in which the 3rd and the 5th harmonics are present, because, before up-conversion, their distance from the wanted tone is relatively large. For this reason an analog filter can be used.

This filter is a combination of a 1st order LPF and a notch filter (see Fig. 5.5). The LPF attenuates all the harmonics which are around one decade or more away from the wanted tone (giving in this way around 20 dB extra attenuation). The notch filter is designed in order to remove the 7th and the 9th harmonics of the wanted tone. The extra required attenuation of this filter is, for a 40 dB SFDR, around 23 dB. This is feasible considering that the LPF also attenuates the 7th and 9th harmonics.

As mentioned in this section the SSB mixing requires quadrature signals. Quadrature signals at point A in Fig. 5.5 are generated by simply dividing the reference frequency signal by a factor of 2. The generation of a quadrature signal at point B is done inside the Walsh function based harmonic rejection block. Moreover, to ease the implementation of this block (see Sect. 5.3.3), and of the second Walsh shaper at the end of the chain, a division by 4 in frequency is implemented. For this reason (5.1) changes into the following equation:

$$f_{\text{bin}} = f_{\text{ref}} \times \left(\frac{1}{2} \pm \frac{1}{4M} \right) \times \frac{1}{4N} \quad (5.2)$$

This modification does not change the values shown in Table 5.2 except for a scaling factor. Indeed, the values of M and N are chosen to be a multiple of 4. The programmable N and M frequency dividers, therefore, divide by the values given in Table 5.2 scaled down by a factor 4. The reference signal is chosen to be twice the required signal at point A in Figs. 5.5 and 5.4 to enable the usage of a simple and low power digital frequency divider for quadrature generation. These implementation constraints set the relatively high required reference signal used in this architecture.

Concluding Remarks

Concluding, the architecture shown in Fig. 5.5 allows to meet all the required specifications for a fast frequency hopping synthesizer for WSNs. The frequency accuracy is guaranteed from the fact that the generated frequencies are locked to a crystal based frequency.

It has been proven that it is possible to synthesize at least 28 channels with a fine inter-channel spacing of 0.5 MHz. The number of upconverted channels is then doubled by using an SSB up-conversion scheme.

The system does not contain any closed loop and, therefore, its settling time is set by the internal nodes time constants. This allows to achieve a high hopping speed with potentially no variation in the synthesizer power consumption.

Finally, by selectively removing harmonics of the signals involved in the synthesis two benefits are obtained. First, it is possible to use digital (and in principle low power) frequency dividers. Secondly, the required analog filter specifications are greatly relaxed.

The following section describes the circuit level implementation of all the blocks composing the architecture shown in Fig. 5.5.

5.3.3 Baseband Frequency Hopping Synthesizer Implementation

This section focuses on the design at transistor level of the various blocks composing the proposed synthesizer architecture. The main blocks analyzed in this section are the following:

- SSB mixer
- Programmable dividers
- Walsh function based harmonic rejection
- LP-notch filter

Mixer Design

The schematic block diagram of the SSB mixer is depicted in Fig. 5.6. The first thing that can be noticed from Fig. 5.6 is the use of a double balanced passive mixer. A balanced structure is used in order to remove the $f_{\text{ref}}/2$ components that otherwise can be of the same order of magnitude as the wanted sideband. The use of a passive mixer allows to have high linearity and low power consumption because no bias current is used. The major drawback is that it does not have any conversion gain but rather a conversion loss of at least 3.9 dB [117]. Nevertheless, this is not a limiting factor because the input signals $f_{\text{ref}}/4M$ and $f_{\text{ref}}/2$ are relatively large signals. This allows to have large conversion losses without major impact on the SNR.

Given the required SFDR and given that the frequencies involved are below 200 MHz a single-quadrature mixer topology has been used.⁶

The mixer linearity has to be sufficiently high in order to avoid to reintroduce, due to distortion, any odd harmonic (mainly third and fifth) in the wanted signal.⁷ For

⁶In [118] it is shown that a double quadrature mixer achieves a larger sideband suppression with respect to the single quadrature counter part. In general with 1% accuracy in amplitude and phase 40 dB sideband rejection can be achieved. At very high frequencies this can turn out to be a problem but at the low frequencies involved in this architecture, 40 dB image attenuation is achievable.

⁷The system is designed for 40 dB SFDR. The divider-by- N in Fig. 5.3 gives 27.6 dB spur attenuation (see Sect. 5.3.2), which means that the third harmonic generated by the mixer has to be at least 13 dB below the fundamental. This is not a very difficult requirement to achieve.

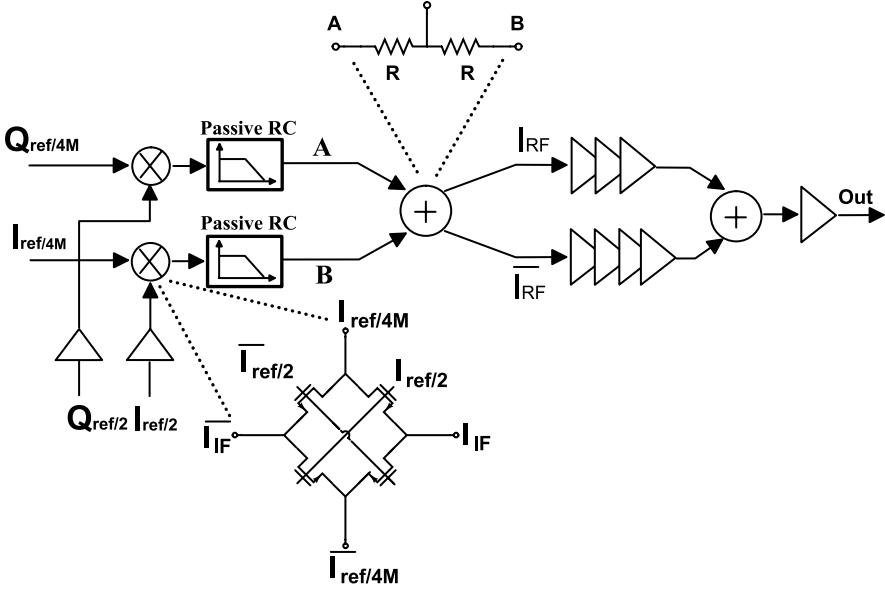


Fig. 5.6 SSB mixer block diagram and partial circuit level implementation

this reason, the divide-by- $4M$ signal needs to be attenuated before being applied to the source/drain of the MOS switches. By carefully selecting the values of the resistances in the Walsh interpolator (see Sect. 5.3.3) with respect to the low-pass filter and switch on-resistance it is possible to attenuate the divide-by- $4M$ signal at the output of the Walsh interpolator. This can be understood with the help of Fig. 5.7. We consider, for simplicity, only one mixer, no differential signals and a single pole RC filter following the mixer. When the switch is on, the divide-by- $4M$ signal is connected via the switch on-resistance R_{on} to the IF node and its 180 degrees shifted version to the $\overline{IF}$ node (see the inset in Fig. 5.6). The resistances R_1 , R_2 and R_3 are the Walsh interpolator resistances while R_{filt} is the RC filter resistance. The filter is followed by a feedback amplifier (the first stage of the following limiting amplifier, see Fig. 5.8). This amplifier sets the node A in Fig. 5.7 to a virtual ground. Supposing now that $R_1 = R_2 = \sqrt{2}R$ and $R_3 = R$ it is possible to prove that the voltage level at the interpolator output for $R \gg R_{on} + R_{filt}$ is approximately the following:

$$V_x \approx \left(\frac{1}{\sqrt{2}}a + \frac{1}{\sqrt{2}}b + c \right) \frac{R_{on} + R_{filt}}{R} \quad (5.3)$$

Therefore, the voltage at the source/drain of the mixer switches is the interpolated output of the Walsh shaper attenuated by the ratio between the R in the interpolator and the sum of the switch on-resistance and the RC filter resistance.

The resistance of the switches in the mixer should be sufficiently lower than the input resistance of the following stage. At the same time, these switches should be small enough to allow a relatively high operating frequency of 144 MHz set by the $f_{ref}/2$. The following stages are the passive RC filters, which implement the passive

Fig. 5.7 Mixer working principle, single ended signals, no SSB operation and single pole RC LPF

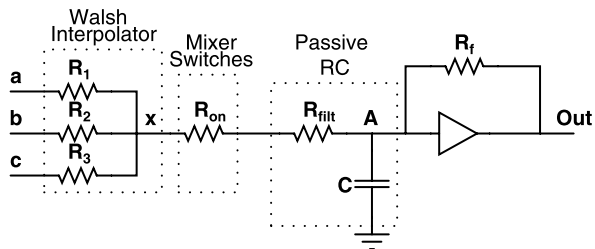
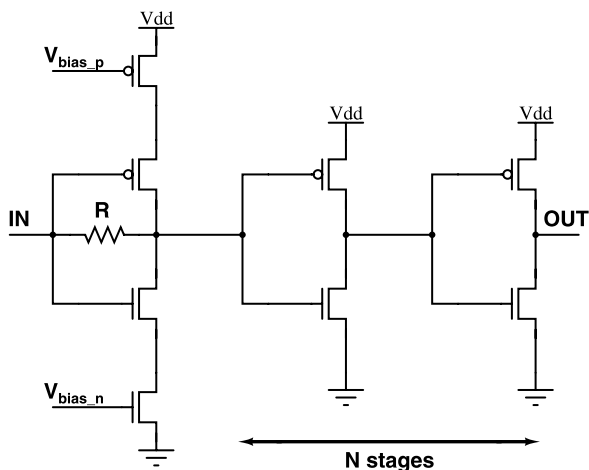


Fig. 5.8 Schematic of the hard limiters in Fig. 5.6: N is equal to 3 for the path A and to 4 for the path B



filter in Fig. 5.5 after the SSB mixer. The passive filters are simple 2nd order RC LPF which attenuate harmonics of the $f_{\text{ref}}/2$ as mentioned in Sect. 5.3.2.

The SSB selection is performed according to the following equation

$$I_{\text{SSB}} = I_{\text{ref}/2} I_{\text{ref}/4M} + Q_{\text{ref}/2} Q_{\text{ref}/4M} \quad (5.4)$$

where $I_{\text{ref}/2}$ and $I_{\text{ref}/4M}$ are shown in Fig. 5.6, while $Q_{\text{ref}/2}$ and $Q_{\text{ref}/4M}$ are the quadrature counterparts of $I_{\text{ref}/2}$ and $I_{\text{ref}/4M}$. The addition in (5.4) is performed by using a simple passive interpolation. The sideband selection is obtained by swapping the sign⁸ of the $f_{\text{ref}}/2$ signal [118]. This is obtained by a simple set or reset operation of the frequency divider at point A in Fig. 5.5.

Though not shown in Fig. 5.5 all the signals are differential signals. Therefore, the last step consists of subtracting these differential signals to obtain a single ended signal that drives the divide-by- N block in Fig. 5.5. After sideband rejection, the signal has a peak-to-peak value in the order of 100 mV. For this reason limiters are used after sideband rejection is performed. The limiter eliminates any amplitude variation and allows the resulting waveform to be input to the next stage, which is a digital divider (divide-by- N block in Fig. 5.5). The hard limiter schematic is

⁸This means switching the differential signals.

shown in Fig. 5.8. The first stage of the hard limiter is a current starved inverter with resistive feedback. The current is set to $15\ \mu\text{A}$. The resistance R is a $30\ \text{k}\Omega$ resistance. The local feedback on the inverter allows to set the DC voltage at the input of the current starved inverter to roughly its trip point. If the current is not limited, the power consumption of this first stage can become very big. Given the fact that the circuit needs to handle relatively low frequencies, the current is limited using a PMOS and an NMOS transistors acting as current sources.

At the output of the current starved inverter, the signal is sufficiently strong to drive inverter logic gates. Because differential signals are used in this circuit, the difference between the positive and the 180 degrees shifted signal is obtained using 3 inverter stages for I_{RF} in Fig. 5.6 and 4 inverter stages for $\overline{I_{\text{RF}}}$ in Fig. 5.6. The final summation is performed by using a passive interpolation. After the passive interpolation a digital buffer is used to reconstruct a digital signal suitable for the next divide-by- N stage (see Fig. 5.5).

Programmable Dividers

The frequency dividers needed to implement the proposed architecture are digital frequency dividers in order to minimize their power consumption. In Fig. 5.9 the architecture of the divide-by- M frequency divider is shown. The divide-by- N frequency divider uses the same architecture but with a larger number of bits.

The frequency dividers must be programmable and in order to optimize their power consumption and area, they are custom designed. The input reference signal is directed, using the multiplexers, through one or more frequency dividers depending on the division ratio required to synthesize a given frequency bin. The multiplexers are set using the bits B_i . When a frequency divider is not used, its input is grounded reducing its power consumption to virtually zero. In this way only working dividers and multiplexers contribute to the overall power consumption.

Only divide-by-two and divide-by-three frequency dividers are required in this architecture. It is known that a divide-by-two frequency divider can give I and Q quadrature signals with a precision that depends on the input signal duty-cycle. If the

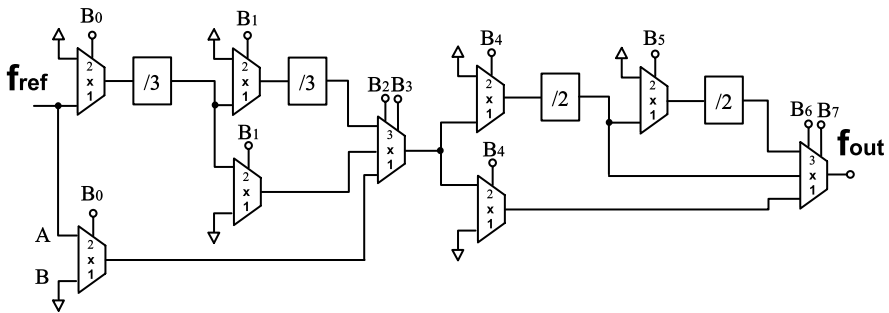


Fig. 5.9 Schematic block diagram of the programmable divide-by- M frequency divider; the divide-by- N frequency divider is designed using the same architecture

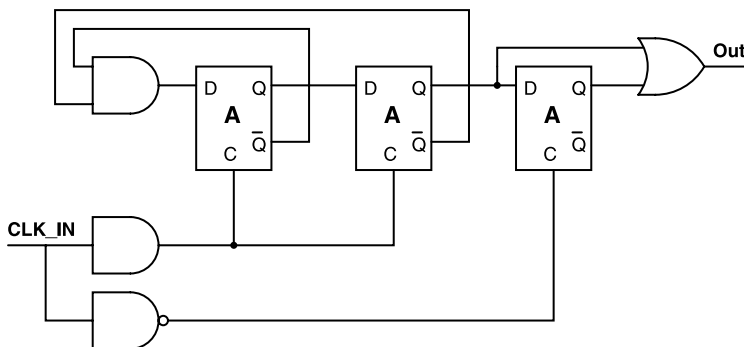


Fig. 5.10 Divide-by-three frequency divider

input signal duty-cycle is 50%, the quadrature relation between the I and Q signals at the output of the frequency divider is perfect. If the input signal does not have a perfect 50% duty-cycle, the I and Q signals will not present a phase difference of 90 degrees, but it will differ from that proportionally to the duty-cycle error.

Quadrature signals are required in the Walsh function based harmonic rejection blocks (see Sect. 5.3.3). Any error in the phases synthesized by this block causes a lower harmonic rejection and therefore, a higher system SFDR. Therefore, all the required frequency dividers must output a 50% duty-cycle waveform with less than $\pm 1\%$ error.

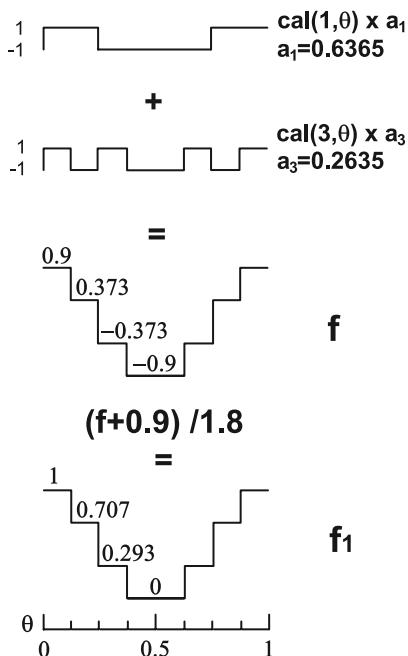
While generally a master-slave based divide-by-two frequency divider outputs a 50% duty cycle regardless of its input signal duty cycle, this is not generally true for a divide-by-three frequency divider. Therefore, the divide-by-three frequency divider has been synthesized as a state machine as shown in Fig. 5.10. Actually, it is a counter in which only the states 0, 1 and 6 are used, but globally it has 8 states (from 0 to 7). Therefore, it is necessary to make sure that, if at power up or for any other reason, the counter comes into one of the unused states, the divider will move out of this state and does not lockup. The design shown in Fig. 5.10 makes sure that there is no unused state in which the counter can lockup.

Walsh Function Based Harmonic Rejection Block

The Walsh functions [119] constitute a set of orthogonal functions. Given the fact that they can have only two values (+1 and -1) and that they are linear functions, they are very suited for digital computation.

The Walsh function can be seen as the equivalent of the Fourier series in the digital domain. Therefore, as any signal can be decomposed in Fourier series, any signal can be decomposed in Walsh series as well. While the Fourier series decompose the wanted signal in an infinite summation of sinewaves and cosinewaves, the same signal can be also decomposed in an infinite summation of bipolar linear functions called “cal” and “sal”, where “cal” stands for Cosine wALsh and “sal” stands for Sine wALsh.

Fig. 5.11 Synthesized function for $n = 3$



It can be proven that to synthesize a sinewave $g(\theta) = \cos(2\pi\theta)$ an infinite number of cal functions are required [120]. It can also be proven that if the Walsh series of $g(\theta)$ is truncated to the 2^{n-1} cal term the resulting waveform has a spectrum given by the following equation [120]:

$$\begin{aligned} \widehat{g(\theta)} = & \left(\text{sinc}^{-2} \frac{1}{2^n} \right) \left\{ \cos(2\pi\theta) \right. \\ & + \sum_{q=1}^{\infty} \frac{(-1)^q}{q2^n + 1} \cos[(q2^n + 1)2\pi\theta] \\ & \left. - \sum_{q=1}^{\infty} \frac{(-1)^q}{q2^n - 1} \cos[(q2^n - 1)2\pi\theta] \right\} \end{aligned} \quad (5.5)$$

where $q = 1, 2, 3, \dots$, and $\text{sinc}(x) = \frac{\sin(\pi x)}{\pi x}$.

Considering for example $n = 3$, see Fig. 5.11, it can be seen from (5.5) that neither the third nor the fifth harmonic is present. The wanted signal is slightly attenuated and the first and largest two harmonics in the spectrum are the seventh and the ninth with amplitudes 17 dB (1/7) and 19 dB (1/9) below the fundamental.

Given $n = 3$, two cal functions are required to synthesize the wanted signal (called $\text{cal}(1, \theta)$ and $\text{cal}(3, \theta)$ ⁹ in Fig. 5.11). When they are superimposed after scal-

⁹ θ is defined between 0 and 1 and it represents the phase of the signal modulus 2π .

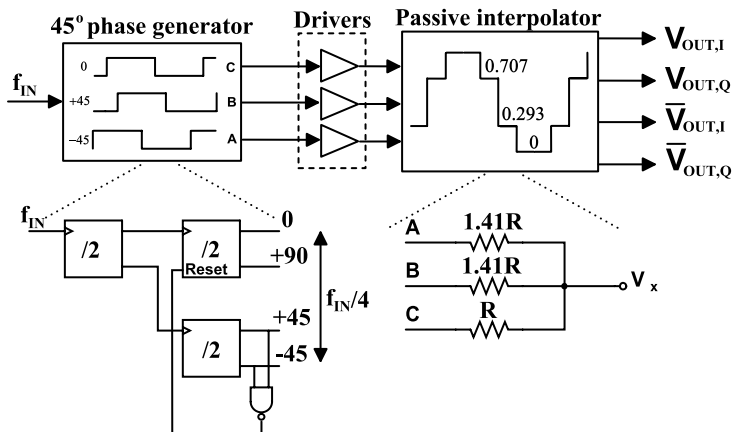


Fig. 5.12 Walsh function harmonic rejection block

ing by a_1 and a_3 respectively, the resulting waveform looks like the one depicted in Fig. 5.11 (called function f in Fig. 5.11). Given the fact that the two cal functions can also be negative, the resulting waveform is also a bipolar waveform, which is not very well suited for single-supply voltage IC processes. As stated before, the Walsh functions are linear and therefore, they can be multiplied by a constant or shifted by a constant value without changing their properties. Adding 0.9 and dividing by 1.8 results in the waveform f_1 in Fig. 5.11, which can be implemented using a single supply voltage.

The implemented Walsh function harmonic rejection block is depicted in Fig. 5.12. It is composed by two main blocks. The first block is a ± 45 degrees generator. This block generates the required phases (0, 90, $+45$ and -45 degrees) for the interpolator block. The interpolator is a passive interpolator and therefore, drivers are used in between the phase generator block and the interpolator.

To obtain a waveform that does not contain the third and the fifth harmonic of the fundamental tone, the resistances of the interpolators should be correctly weighted. The resistors driven by the signal in A and B in Fig. 5.12 ($+45$ and -45 degrees) are $\sqrt{2}$ larger than the resistor driven by the signal in C (0 degrees). The phase generator also outputs the 90 degrees shifted waveform. If correctly combined, these phases can produce not only the required waveform but also a 90 degrees shifted version and all the required differential (180 degrees shifted) waveforms as well. Therefore, the phase generator output is fed to a set of drivers and interpolators, each of them generating a third and fifth harmonic free waveform but with different phases (0, 90, 180 and 270 degrees). These phases are required to drive the SSB mixer or the upconversion mixer (after the LPF-notch filter) as shown in Fig. 5.5. The drivers are CMOS logic buffers or inverters depending on the required signal phase.

The phase generator block is composed by three divide-by-two frequency dividers. The first divider outputs the I and the Q signal. These two 90 degrees phase shifted signals drive two other frequency dividers generating the required phases. To

be able to have the correct phase alignment between the four phases required, a reset sequence is generally used. In the proposed architecture the simple addition of a logic gate avoids the use of any start-up sequence to correctly set the outputs of the frequency dividers. The NAND gate senses the $+45^\circ$ and -45° degrees signals and resets, eventually, the frequency divider that outputs the 0° and 90° degrees phases. If the phases are aligned correctly, the NAND gate does not have any influence on the circuit, but if the initial phase relation among the output signals are not correct, the NAND gate aligns those phases within 4 clock cycles. If the initial phases do not have the correct relation the system functionality is compromised.

The Walsh based harmonic rejection block requires its input frequency to be four times its output frequency in the current implementation. Furthermore, it is important to understand how resistor matching as well as phase mismatch (between the multiphase signals), do affect the harmonic rejection. As an example we discuss the effect of the aforementioned non idealities on the rejection of the third harmonic.

It can be proven (see Appendix), that if a mismatch between two resistances in the passive interpolator but no phase error is present, the residual third harmonic amplitude level is:

$$A_3 = \frac{\varepsilon}{3} \quad (5.6)$$

where ε is the mismatch between the R resistance and the $\sqrt{2}R$ resistances in the interpolator. If a phase error between the multiphase signals but no amplitude error is present, and supposing the phase error small, the residual third harmonic amplitude is:

$$\phi_3 = \sqrt{2}\varphi \quad (5.7)$$

where φ is the phase unbalance between the $\pm 45^\circ$ signals and the 0° signal. Considering the levels of the fundamental tones in the same conditions, it can be proven that, in the two cases, the rejection of the third harmonic is the following:

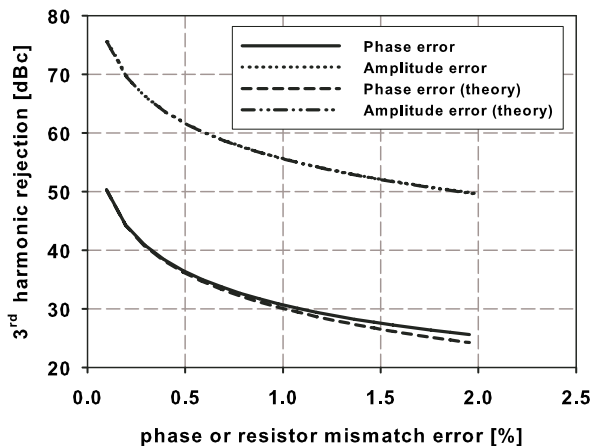
$$HR3_A = \frac{6\sqrt{2}}{\varepsilon} \quad (5.8)$$

$$HR3_\phi = \frac{2}{\varphi} \quad (5.9)$$

To properly compare the effects of mismatches in the resistance values and the error in the phase relation between the multiphase signals, we normalized the errors to the R value and to one period (360°) of the multiphase signal. The result can be expressed as a percentage of the two quantities and the third harmonic rejection can be plotted, as a function of this normalized error. The plot is shown in Fig. 5.13.

The effect of a phase mismatch between the signals at the passive interpolator input can be severe and it does have an impact larger than the effect on the mismatch between the passive interpolator resistance values. For this reason, the design of the digital signal distribution to the phase interpolators is very important.

Fig. 5.13 Effect of the resistor mismatch or phase mismatch on the 3rd harmonic rejection



LP-notch Filter

The LP-notch filter is used to remove higher order harmonics at the output of the baseband synthesizer before they enter the final high frequency upconversion stage. The notch is used to remove the 7th and the 9th harmonic of the wanted tone, while the LPF attenuates higher order harmonics. Both filters are first order filters and they are implemented as passive structures to minimize the power consumption.

The schematic block diagram of the filter is shown in Fig. 5.14. The system requires four of these filters given the quadrature differential signals needed at the output.

The tunable LPF is realized using a simple R - C structure in which the C is tuned in a thermometer fashion while the R is kept constant. The output of the tunable LPF is coupled to the notch filter using a source follower as a buffer.¹⁰ The source follower must be optimized in terms of distortion. Indeed, its odd harmonic distortion cannot be removed anymore and can reduce the system SFDR. The even distortion is greatly reduced by using differential signals. In the source follower, the main source of distortion comes from the modulation of the threshold voltage via the bulk effect. For this reason a triple well transistor has been used and the bulk has been connected to the source. The use of a triple well transistor reduces the operating frequency of the source follower via the n -well parasitic capacitance. Given the low frequencies involved, this was not a major obstacle.

The last block of the filter is the notch filter. The notch has to be low power, but also it should not require a very precise tuning mechanism. The higher the Q of the filter the more accurate the tuning mechanism has to be. On the other hand, the lower the Q of the filter the more the losses the wanted signal experiences while notching out the unwanted harmonics. The input of the notch filter, though is not a rail to rail signal anymore, has a still sufficiently high amplitude to sustain some

¹⁰This avoids pole shifting due to reciprocal loading effect between the LPF and the notch filter.

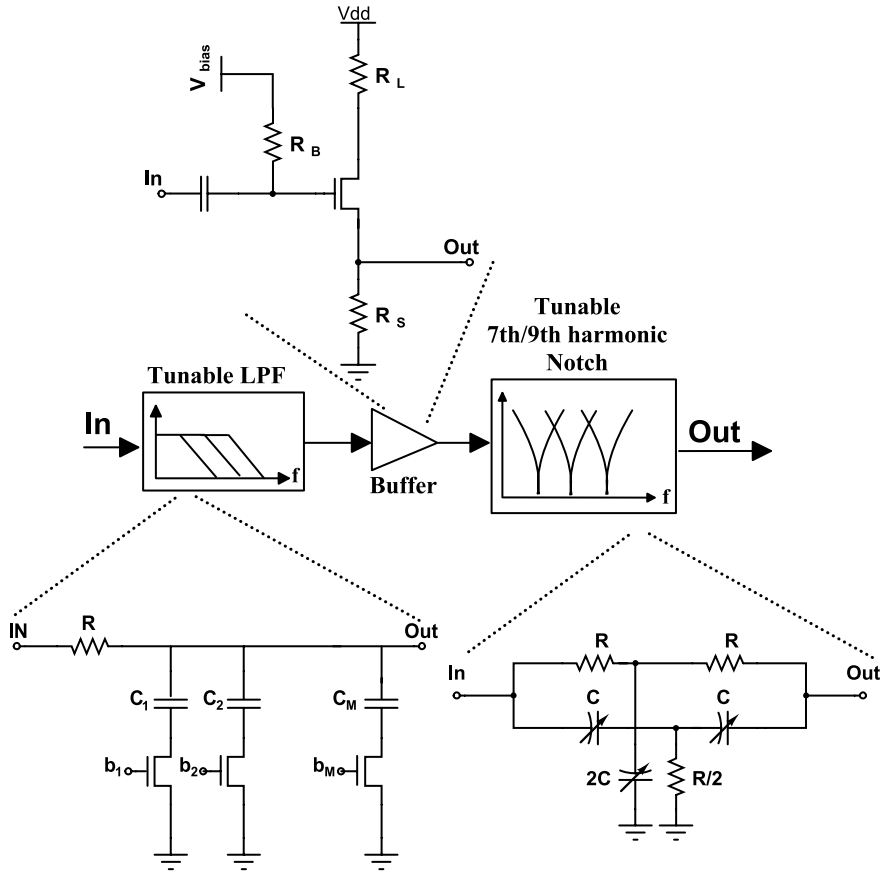


Fig. 5.14 Schematic block diagram of the LP-Notch filter and its circuit level implementation

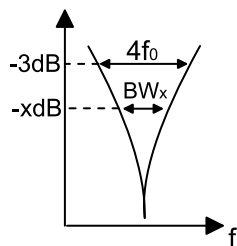
losses. Avoiding fine tuning of the filter avoids extra circuitry saving power. For this reason, a twin-T passive notch filter has been used. This filter has a Q equal to 0.25. This means that it is possible to achieve the required extra harmonic suppression of the 7th and the 9th harmonics with a single notch filter. The position of the notch is equal to:

$$f_{\text{notch}} = \frac{1}{2\pi RC} \quad (5.10)$$

The position of the notch is changed by simply tuning the capacitance as shown in Fig. 5.14. The attenuation of the wanted signal due to the low Q of the filter can be easily calculated using the following equation:

$$Att = 10 \log_{10} \left(1 + \left(\frac{4f_{\text{notch}}}{BW_x} \right)^2 \right) \quad (5.11)$$

Fig. 5.15 Bode plot of the twin-T notch filter



where the meaning of BW is shown in Fig. 5.15. From (5.11) and Fig. 5.15, it is clear that the attenuation of the wanted signal is below 3 dB and therefore, does not constitute a problem.

5.4 Generation of a 288-MHz Reference Clock

As mentioned in Sect. 5.3 a 288 MHz reference clock is required in order to generate all the required frequency bins. Besides some very common reference signal characteristics, like frequency stability and good phase noise, two more requirements are necessary in order to integrate the reference signal generation block into an ultra-low power wireless sensor network node:

- Low power consumption
- Low cost

There are several architectures that can be used to generate a reference clock signal. We can distinguish them in two categories [121]:

- Tuned oscillator based
- Frequency synthesizer based

Frequency generation based on tuned oscillators comprises an amplification element and a feedback network consisting of a tuned element. The tuned element sets all the requirement of the reference oscillator and therefore, it is the core of this kind of frequency generation architecture. The amplifier (generally transistor based) is used to recover the inevitable losses of the tuned element.

A simple LC stabilized oscillator is generally not stable enough and therefore, it is not suitable to be used as a reference frequency. Using a transmission line in place of the LC circuit will preserve a higher accuracy and usually exhibits less drift with temperature but it can be cumbersome to implement especially at such low frequencies.¹¹ A ceramic resonator, widely used as a clock source for micro-processor applications, has a frequency tolerance of only ≈ 3000 ppm and hence is

¹¹The transmission line is used to implement the reactive element in a resonator. It is generally implemented as a $\frac{\lambda}{2}$ microstrip line. At low frequencies λ can be quite big increasing, therefore, the wireless node form factor.

not adequate for most radio applications. Coaxial resonators (often ceramic based) can provide very high Q elements and hence have an excellent phase noise. Their temperature stability is however poor and therefore, they require additional circuitry to compensate for the large temperature drift. SAW devices can provide the accuracy needed, giving typically 300 ppm stability over a wide temperature range for frequencies up to 1 GHz or beyond. They usually have an initial tolerance which can be tuned on manufacture (at a increased manufacturing price). Because they provide a low cost, on-frequency source, SAW based oscillators are very popular for low power radio solutions. Another very popular solution for the generation of a reference signal is based on quartz crystals as feedback tuned element. For a very cheap crystal, the temperature tolerance is about 15 ppm, whereas for a TCXO, about 3 ppm tolerance, or better, is possible. The major drawback of crystal based oscillators compared to SAW devices is that the oscillation frequency is typically below 30 MHz for a fundamental mode design. It can be pushed up to 100 MHz using a fifth overtone design. Unfortunately, oscillator design based on an overtone of the fundamental resonance frequency of a crystal is more difficult and translates in higher power and cost. Concluding, among all the tuned oscillator based topologies, the one based on the SAW resonator is the most viable solution for the generation of a low power, low cost reference signal.

Other viable ways to generate the wanted 288 MHz reference signal are based on architectures which do not directly use a tuned oscillator. As an example, if we make a signal passing through a non linear function, several harmonics are generated. Therefore, it is possible to generate the reference signal using one of those harmonics. A tuned circuit can be used to extract the wanted harmonic while rejecting all the other unwanted harmonics. Though feasible, this is cumbersome for two reasons. First the level of the harmonics decreases with the harmonic order and second a tuned circuit is required to extract the wanted harmonic. This means that generally non-linear blocks are cascaded and every time the second or the third harmonic is extracted. This translates in higher power consumption, and higher cost.

An alternative to using a non-linearity to generate harmonics is the use of linear mixing to create multiples of the fundamental frequency. A fully balanced multiplier acts as a frequency doubler and can, therefore, generate the 288 MHz reference signal from a lower frequency signal. Unfortunately, this technique is not very flexible and can be power hungry if low phase noise is required.

Injection locking is a common technique for providing stable microwave oscillators. The injection locking is based on the principle that injecting a tone close in frequency to the free running frequency of the oscillator forces the oscillator to oscillate at the frequency of the injected tone. This technique, however, finds little use in low frequency applications due to its extreme design and layout sensitivity.

The most common architecture to multiply a stable reference frequency is to use a PLL. The PLL generates at the output a frequency equal to N times the reference signal where N is the division ratio of the feedback divider. In this way, the reference signal can be a quartz based oscillator using the fundamental resonance frequency of the quartz, while the frequency multiplication is achieved via the PLL loop. Good stability and accuracy is achieved in this way. Moreover, the required 288 MHz

reference signal is a fairly low frequency signal for the latest CMOS technologies. This means that an integer PLL can be designed and because of the low frequencies involved, a fairly low power consumption can be expected. The penalty compared to a SAW based oscillator is the extra circuitry required for the PLL, which increases the die area and therefore, the cost. On the other hand, crystals are generally cheaper and have a better aspect ratio than SAW resonator giving a cost advantage for a PLL based frequency generator.

Lastly, another way to generate a reference signal is to use a DDS. Unfortunately, generating a 288 MHz signal with a DDS require a reference clock running at least at twice the frequency. This generally translates in a very high power consumption.

Concluding the most suitable architecture for the generation of the 288 MHz reference signal used in the baseband frequency hopping synthesizer is either a SAW based oscillator or an integer PLL using a quartz reference signal.

5.5 Receiver Design at System Level

In Sect. 2.4 the different receiver architectures have been described. It has been shown that, though the zero-IF receiver topology requires a quadrature downconversion, it remains simple to implement and offers a higher degree of integration. Reducing the external components translates in a reduction in the power consumption because the outputs of the various blocks can be matched to a much higher than 50 Ω impedance, reducing, in this way, the biasing current.

Unfortunately, homodyne receivers suffer from I/Q mismatch, even-order distortion, DC offset and flicker noise. Flicker noise and DC offset can be overcome for low data-rate application by employing a wideband FSK modulation followed by a high-pass filter which can be sometimes a simple AC coupling capacitance. This technique solves also partially the even-distortion problem in case the received signal contains some amplitude modulation. Indeed, if the information is carried by the frequency, the received signal can be hard-limited cancelling any unwanted amplitude modulation. I/Q mismatch can be reduced by proper layout and matching techniques. Given the aforementioned considerations, the proposed RX architecture is a zero-IF topology and is depicted in Fig. 5.16.

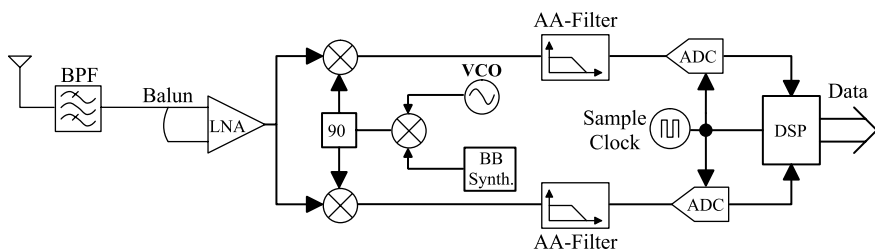


Fig. 5.16 Proposed zero-IF receiver architecture

The BPF after the antenna is used to remove out of band interferers that can saturate the LNA causing clipping and therefore, heavy distortion and system gain desensitization. After the filter the LNA needs to amplify the signal to a level that can be further processed by the following blocks. Quadrature downconversion is performed in order to translate the information signal to baseband. After the quadrature downconversion an in-phase (I) and a quadrature-phase (Q) paths are present. Therefore, all the following blocks are doubled in terms of hardware implementation. The AA-filter attenuates all the portions of the spectrum, which do not contain useful information. In this way, any destructive folding coming from the subsequent analog-to-digital conversion is avoided. Finally, the DSP can demodulate and reconstruct the transmitted data.

5.5.1 Receiver Link Budget Analysis

Any channel and demodulator imposes some boundary conditions on the RF front-end. These conditions can be specified in terms of SNR and can be translated to circuit concepts like NF, gain and distortion. Next, it is necessary to translate these boundary conditions into boundary conditions of the front-end sub-blocks. The following analysis refers to the 2.4 GHz ISM band but can be easily translated to any band of interest (like the 902–928 MHz band).

Propagation Link Budget Analysis

The first step consists in calculating the minimum received signal known as receiver sensitivity. In Sect. 1.4.3 it has been proven that for a 10 meters communication distance the required receiver sensitivity in the 2.4 GHz ISM band is -76.5 dBm. The required SNR at the demodulator input largely depends on fading conditions. In Sect. 4.2.3 the relation between SNR and BER when fading is present has been derived (see (4.5)). From (4.5) it follows that for a 1% BER the required SNR is 20 dB if fading is considered. Furthermore, this already large SNR has to be met also when interferences are present. Because these interferences are random in nature, the demodulator cannot differentiate them from the channel noise and will process them during the demodulation operation. In such a situation, the SNR is reduced even more.

Link Budget Analysis of Discrete Parts

The first step is to filter out the noise and the interferences which are out of the band of interest. This is accomplished by the BPF in Fig. 5.16. The quality factor of this filter is generally quite high. For the band of interest the required Q has to be $2400/83.5 \approx 29$. It is generally realized as an external LC network, but this

approach increases the form factor and the cost of the wireless node. Therefore, the integration of the BPF is highly desirable.

Several approaches can be used in the design of an integrated BPF. One possibility consists in using integrated passive components to realize the common LC ladder of a BPF. Unfortunately, on-chip inductors have very poor Q limiting de facto the Q of the entire filter. The Q is related to the losses in the inductive element of the filter and therefore, can be made very large if the losses are minimized. One way to accomplish this result is to use on-chip transformers [122]. The unloaded filter Q is given by

$$Q_{\text{unloaded}} \approx Q_{\text{inductor}} \frac{1 + K}{1 + \beta K} \quad (5.12)$$

where K is the coupling factor and β takes into account the losses in the transformer. Equation (5.12) shows that a small increase of K leads to a significant increase of the filter unloaded quality factor. Starting from a Q_{inductor} of about 2 a filter Q of about 10 has been achieved.

To increase the unloaded Q of the filter an active topology is needed [123]. Unfortunately, though a Q of about 3000 is achieved, the required active components (CMOS transistors) consume around 10 mW while achieving 3.1 dB gain. Consequently, this topology cannot be used for ultra-low power applications.

The last possibility relies on using FBAR resonators. Their size is between 5 and 8 times smaller than SAW based filters. A 1.9 GHz FBAR BPF has been realized in [124]. The total required volume is $1.0 \text{ mm} \times 1.0 \text{ mm} \times 0.7 \text{ mm}$. An insertion loss of 3.3 dB has been measured with 35 dB rejection in the upper stop-band and 25 dB in the lower stop-band.

Concluding, the two most promising ways to realize an on-chip BPF are the one based on on-chip transformers and the one based on FBAR resonators. Nevertheless, the Q of the first topology is still below the receiver specifications while the second one cannot be still considered fully integrated. Therefore, there is a large margin of improvement in this field for researchers. In this book, passive structures will be considered for the BPF, the balun and the TX-RX switch. The only required specification is the attenuation, which has been derived from data-sheet of off the shelf components.

Link Budget Analysis for Integrated Parts

The receiver chain has to assure a certain SNR at the demodulator input. This SNR has been previously evaluated around 20 dB when worst case fading is considered for a 1% BER. For noise calculation the noise bandwidth needs to be calculated. The noise bandwidth is the smallest bandwidth in the receiver chain and it strongly depends on the channel bandwidth.

Considering a 2 kbps data-rate and a modulation index equal to 5, from the Carson's rule a 24 kHz signal bandwidth is required. Increasing the data rate will increase the bandwidth and, therefore will reduce the available receiver NF. If the

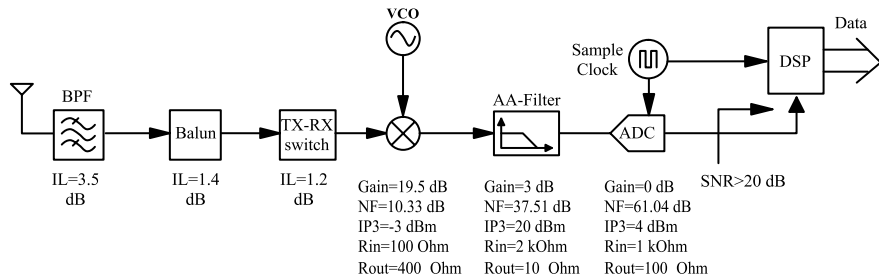


Fig. 5.17 Receiver front-end blocks specifications

bandwidth has to remain the same in order to relax the receiver requirements, then the modulation index has to be reduced. This will make the effect of flicker noise and DC offset more severe.

It is important to notice that any signal with a frequency higher than the Nyquist frequency, will fold to the Nyquist band after sampling. Therefore, an AA-filter is required before the signal is converted in the digital domain. If the signal is sampled at the minimum required sampling frequency (two times the maximum signal bandwidth), the specification on the AA-filter will become prohibitive. This means that generally a certain over-sampling ratio is needed at the ADC side in order to relax the AA-filter specifications. In this case an over-sampling ratio of four has been considered. Concluding, for a 24 kHz signal bandwidth a 192 kHz sampling frequency can be used. In Fig. 5.17 the specifications for each block of the receiver have been derived (only one path between the *I* path and the *Q* path has been considered).

While it appears that no LNA is present, in reality it is merged with the mixer. To meet the low power target, the simple cascaded LNA-mixer topology is not efficient. Therefore, a good strategy to minimize the power consumption is to reuse the bias current of the LNA to bias also the mixer. More details about this LNA-mixer topology is given in Sect. 5.5.2. The insertion losses of the passive components (BPF, Balun and TX-RX switch) are derived from commercially available product datasheets.

Any radio system has to be designed to achieve two targets:

- Discerning the signal from the noise
- Discerning the signal from other signals

To discern the wanted signal from the noise, a certain SNR is required. The worst case condition is when the signal is at its sensitivity level. In this case we have to assure that the SNR at the demodulator input (in practice at the input of the DSP core) has to be at least 20 dB. Now the system design starts with an initial NF for each block refining it till the SNR at the demodulator input hits the wanted target and the block NF is considered feasible. The signal and noise levels at the input of each block in the receiver chain are depicted in Fig. 5.18. From Fig. 5.19 it is possible to see that, given the NF specifications shown in Fig. 5.17, the SNR at the demodulator input hits the 20 dB specification. The decimation is performed inside the DSP block and it is in reality not a stand-alone block. The effect of the

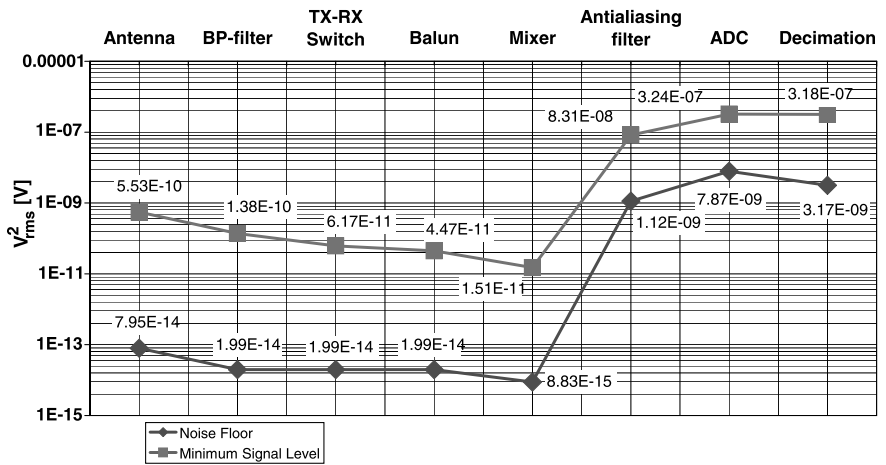


Fig. 5.18 Signal and noise levels along the receiver chain referred to the input of each block

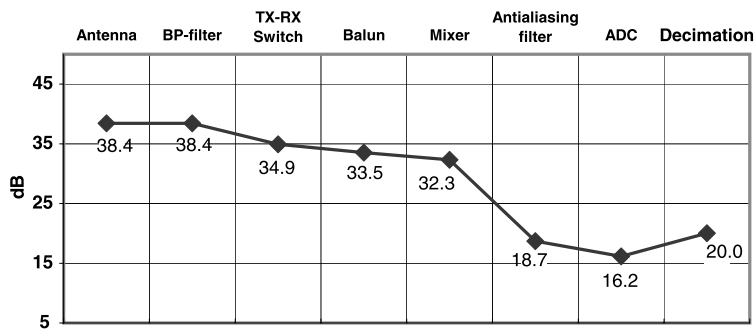


Fig. 5.19 SNR along the receiver chain referred to the input of each block

decimation is to reduce the noise bandwidth from 96 kHz (used in the ADC to relax the AA-filter specifications) to 24 kHz, which is the signal bandwidth. This improves the SNR from 16.2 dB to 20 dB.¹²

The most stringent requirement in terms of noise is set by the first gain block in the chain, which is the LNA-mixer. Indeed, all the subsequent blocks have a larger input signal and, therefore, their noise contribution has a lower impact on the SNR. All the NF are calculated with respect to the driving impedance.¹³

¹²The factor between the oversampled bandwidth and the signal bandwidth should cause a 6 dB improvement in the SNR if the ADC were noiseless. In reality the ADC contributes itself to the overall SNR and, therefore, the SNR improvement after the decimation is reduced to around 4 dB.

¹³In this book we adopted a different approach in the calculation of the NF of the RX blocks. A very common and widely used approach is to refer the NF to a common impedance (in general 50 Ω or 75 Ω). In this way the NF is conceptually disconnected from the real implementation. Referring

An ADC is in general specified with its SNR or its ENOB. Therefore, it is useful to derive, starting from the calculated NF, the required ADC SNR and ENOB. The ADC NF is close to 61 dB with respect to the output impedance of the AA-filter (10 Ω). Now remembering that after decimation the noise bandwidth is reduced by a factor 4, the ADC noise power can be expressed by the following equation:

$$V_{\text{rms,ADC}}^2 = 10^{\frac{NF}{10}} \times k T R_s \frac{f_{\text{sample}}}{8} \quad (5.13)$$

where NF is the ADC noise factor expressed in dB, k is the Boltzmann constant, T is the temperature expressed in degrees Kelvin, R_s is the source resistance of the driving stage and f_{sample} is the ADC sampling frequency bandwidth (192 kHz in this case). With these values, the noise contribution of the ADC is about 35.5 μV . Now supposing that the ADC FSR is about 0.4 V, the required ADC SNR in dB can be derived from the following equation:

$$\text{SNR}_{\text{ADC}} = 20 \log_{10} \left(\frac{\text{FSR}}{2 \sqrt{2 V_{\text{rms,ADC}}^2}} \right) \quad (5.14)$$

where $V_{\text{rms,ADC}}^2$ is the ADC noise contribution. From (5.14) the required ADC SNR is 72 dB. This corresponds to an ENOB, which can be calculated using the following equation:

$$6.02 \times \text{ENOB} = \text{SNR}_{\text{ADC}} - 1.76 - 20 \log_{10} \left(\frac{f_{\text{sample}}}{B_{\text{signal}}} \right) \quad (5.15)$$

where B_{signal} is the signal bandwidth (24 kHz in this example). Concluding, for the system under analysis, this corresponds to an ENOB equal to 10.7.

From (5.14) and (5.15) it can be seen that the SNR due to quantization can be improved by over-sampling the incoming signal. Taking over-sampling to the extreme, it is possible to achieve any desired SNR with one bit sampling at a sufficient high rate. Unfortunately, this technique is limited by two constraints:

- Higher ADC sampling rate and higher DSP working frequency
- Trade-off between AA-filter selectivity and ADC dynamic range

In the first case, the over-sampling places a bigger burden on the DSP because it increases the bit rate. It can be easily proven that if the SNR is kept constant, and bits of resolution are traded for a higher sampling rate, the overall bit rate increases. This can pose a severe overhead in the DSP stage in terms of power consumption.

In the second case, it should be noticed that any radio system has to work in an environment full of other radios using the same propagation medium (the air) and potentially interfering with each other. Therefore, every radio must have a certain dynamic selectivity to cope with this scenario. Considering a simple 2nd order Butterworth filter as an AA-filter, the signal and blockers level throughout the chain are

the noise figure to the driving impedance creates a connection between the real implementation and the noise contribution of the block to the overall system noise.

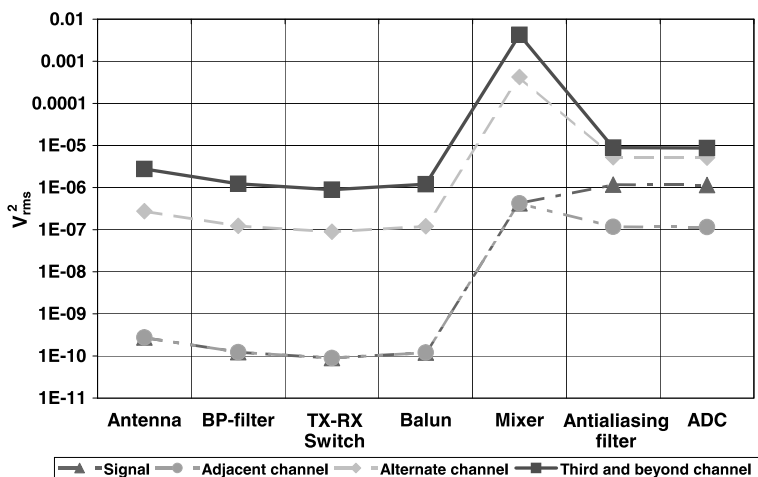


Fig. 5.20 Signal and blockers levels along the receiver chain

depicted in Fig. 5.20. As can be seen, only around 10 dB Carrier-to-Interferer (C/I) ratio is achieved in the adjacent channel while for the alternate and third and beyond channel the situation is even more dramatic because the interferer can have a higher power compared to the wanted signal.

Therefore, either the ADC has to have enough dynamic range¹⁴ (which means enough bits) and linearity to convert in the digital domain the interferer or the selectivity of the analog filter has to improve by increasing the order of the AA-filter. In this way using a digital post-filtering it is possible to achieve the required selectivity.

Due to the fact that in the newer technologies the cost in terms of both power and area of a digital function is foreseen to decrease with time, it looks more promising from a power and cost point of view to use a weak (for example a second order) analog filtering followed by a strong digital post-filtering compared to a pure analog approach (higher order AA-filter). A simple and straightforward way to implement such an architecture, especially when a high SNDR is required, is a $\Sigma - \Delta$ loop.

Another drawback of the environment in which a wireless radio needs to work is the fact that it needs to cope with tiny signals (as small as the sensitivity level) as well as large signals. As previously stated, the capability of a radio to handle small as well as large signal is measured in terms of its dynamic range. In order to be able to handle signals, which can be several dBs different in amplitude, radio systems generally use an AGC system. The AGC lowers the gain in the presence of a strong signal while it increases it if the signal is relatively small.

Though, the AGC is helpful in optimizing the gain depending on the input signal, it requires extra hardware and therefore, extra power consumption. If no AGC is used, a trade-off exists between the required ADC SNR and the maximum gain

¹⁴The dynamic range is defined as the ratio between the maximum signal that can be detected without causing a large distortion and the minimum detectable signal.

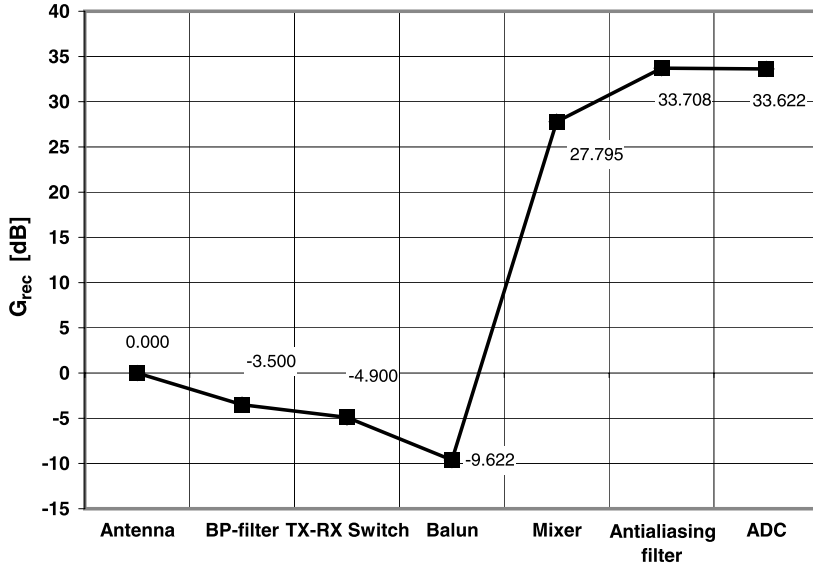


Fig. 5.21 Gain distribution through the receiver chain

achievable in the receiver chain. Indeed, the receiver gain strongly depends on the maximum signal the system needs to handle. When the transmitter is at its minimum distance from the receiver, the signal at the ADC input is not allowed to exceed the FSR. The maximum achievable gain from the input of the mixer-LNA block to the ADC input (see Fig. 5.17) is given by the following equation:

$$G_{rec}^2 = \frac{FSR^2}{2V_{mix-in,rms}^2} \quad (5.16)$$

where the factor 2 converts the FSR from peak value to rms value, and $V_{mix-in,rms}$ is the rms voltage at the mixer input which is equal to:

$$V_{mix-in,rms} = V_{in,max} - L_{BPF} - L_{TX-RX} - L_{Balun} \quad (5.17)$$

where $V_{in,max}$ is the maximum signal at the receiver antenna and L_{BPF} , L_{TX-RX} and L_{Balun} are the attenuations due to the BPF, the TX-RX switch and the balun and all the values are expressed in dBs. Using the attenuation values shown in Fig. 5.17, the maximum possible receiver voltage gain is limited to 34 dB. The gain distribution along the receiver chain is shown in Fig. 5.21.

If a receiver gain larger than 34 dB is required (for correct demodulation of the wanted signal), either the ADC should be designed for a larger FSR (with an increase in the required ENOB at a given SNR) or an AGC system has to be implemented.¹⁵ Given the previous specifications for the ADC and supposing no AGC

¹⁵It is important to notice that, if the AGC decreases the gain because a strong interferer is present, the end result is a decrease in the receiver sensitivity. This is different from the case in which the

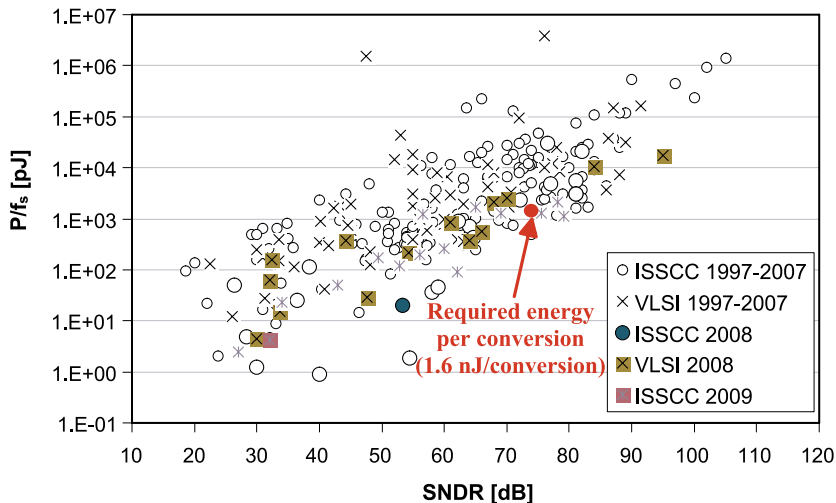


Fig. 5.22 ADC performance survey 1997–2009 [125]

system is used, a minimum distance of 40 cm between TX and RX nodes can be achieved without saturating the ADC.

To evaluate the feasibility of an ultra-low power converter from a system point of view, in terms of power dissipation, we need a figure of merit that links the estimated power consumption to the sampling rate and the ENOB. The following FOM can be used to estimate the required ADC power consumption based on the sampling rate and the ENOB calculated previously in this section:¹⁶

$$FOM' = \frac{P}{f_{\text{sample}}} \quad (5.18)$$

where P is the power consumption and f_{sample} is the sampling frequency. Supposing a target power consumption of 250 μW , 10.7 effective bits and 192 kS/s we obtain a FOM equal to about 1.6 nJ/conversion. The SNDR is equal to $6.02 \times ENOB + 1.76 = 66$ dB and therefore, from Fig. 5.22 it is possible to see that the proposed power consumption is within the state-of-the-art ADCs.

gain is reduced because the wanted signal has a higher amplitude. In this case there is no loss of sensitivity. On the other hand, if no AGC is used, a strong interferer can saturate the first gain stage jamming, in this way, the whole receiver. Therefore, an AGC system besides improving the system dynamic range avoids a strong interferer from jamming the receiver chain improving in this way the reliability of the communication link.

¹⁶This FOM is not often found in literature. The most common FOM in literature considers the ADC ENOB and has an extra term at the denominator of (5.18) equal to 2^{ENOB} . In this case we adopted the FOM of (5.18) because in [125] a comprehensive survey of ADC performance is made using the FOM of (5.18). The widely used FOM can be easily related to the FOM used in this section via the SNDR, which is the x -axis in Fig. 5.22.

Inter-modulation Distortion and Receiver Linearity

In a real environment, there is a finite probability that in-band interferers can produce, due to the receiver non-linearities, inter-modulation products that can fall in the band of interest. If the interferers are very strong, it is possible that the resulting inter-modulation products overwhelm the wanted signal. Interferers can come from other nodes communicating on different channels or other wireless sources using the same bandwidth. The 2.4 GHz band is also used for standards like Bluetooth, Zigbee and WiFi, which can act as sources of interference for the network. In the following analysis only interferers coming from nodes of the same network will be considered. Indeed, if coexistence between different standards has to be taken into account, the required selectivity of the ultra-low power node will increase considerably. This increment in the receiver selectivity is not consistent with the ultra-low power target above a certain limit and therefore, an ultra-low power network will be supposed to work properly mainly in a scenario in which other standards or wireless networks do not constitute the main sources of interference.

For the inter-modulation distortion the first passive blocks in the chain will not be considered. Therefore, the three blocks to be considered are the LNA-Mixer, the AA-filter and the ADC. For the third order inter-modulation distortion calculation, the two interferers to be considered are placed in the adjacent and in the alternate channels. When they mix due to the non linearities in the receiver chain, they generate a third order inter-modulation product (IM_3) which falls in the wanted channel causing degradation in the SNR.

The power levels of these two interferers (interf1 is the adjacent channel interferer and interf2 is the alternate channel interferer) along the receiver chain are depicted in Fig. 5.23. The two interferers are first attenuated along the chain by the passive blocks (BPF, TX-RX switch and balun) and then they are amplified by the active mixer stage. Therefore, at the mixer output, the two signals are boosted consistently. This means that the AA-filter will be the most constrained block in terms of linearity.

The linearity performance of each receiver block can be described by its input third order input intercept point (IIP_3). Then, given the IIP_3 , the third order inter-modulation product at the output of each block can be derived using the following equation

$$IIP_3 = \frac{P_{out} - P_{IM3}}{2} + P_{IN} \quad (5.19)$$

where P_{out} is the alternate channel interferer output power, P_{IM3} is the power of the third order inter-modulation product and P_{IN} is the adjacent channel interferer input power.

The required per block IIP_3 and the generated $IM3$ level are summarized in Table 5.3. As previously mentioned the most stringent requirement is placed on the AA-filter. The ADC, instead, has a lower requirement because at its input the interferer levels have been partially filtered out by the AA-filter. With Table 5.3 it is possible to highlight which blocks constitute the system bottleneck in a power constrained design.

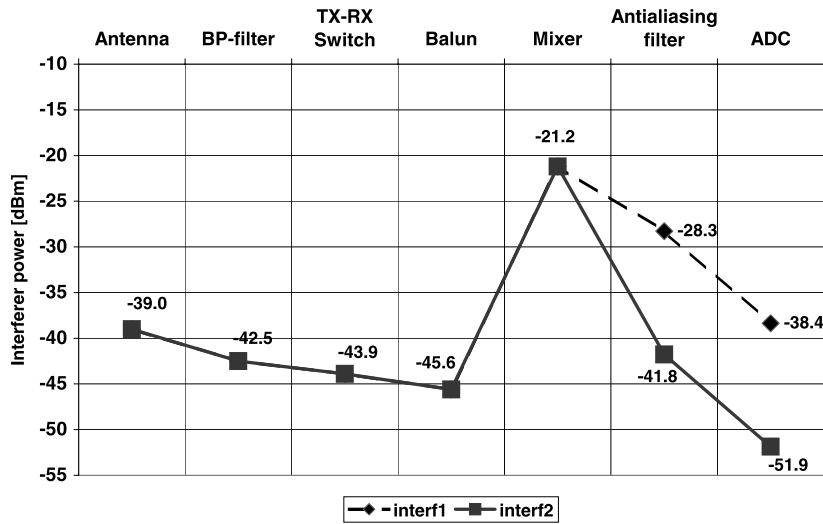


Fig. 5.23 Adjacent and alternate channel interferer levels through the receiver chain

Table 5.3 Per block generated IM_3 and required IIP_3

	IM_3 [dBm]	IIP_3 [dBm]
Mixer	-100.4	-3
AA-filter	-88.7	20
ADC	-109.9	4

For example, given the required gain, NF and IIP_3 , the use of a merged mixer LNA-mixer topology requires improvements either in the linearity performance or in the power consumption. Regarding the AA-filter, three configurations can be explored:

- g_m - C
- Switched capacitor
- Active RC

Active filters based on OPAMPs are more linear than the ones using open loop active transconductors (g_m - C). The linearity will be constrained by the linearity of the passive components in the opamp external feedback loop. A switched capacitor topology is not a suitable choice, because the filter must be able to handle signal frequencies well above the channel frequency. Therefore, either the sampling rate becomes unacceptable or the order of the AA-filter has to increase. Both choices lead to an increase in the filter power consumption. Concluding, the best solution in terms of linearity and power consumption is the active-RC filter in which the OPAMP can be designed with transistors biased in weak inversion to reduce the power consumption. In a zero-IF architecture, special care has to be taken to avoid that the flicker noise degrades the filter performance. Finally, a calibration is needed

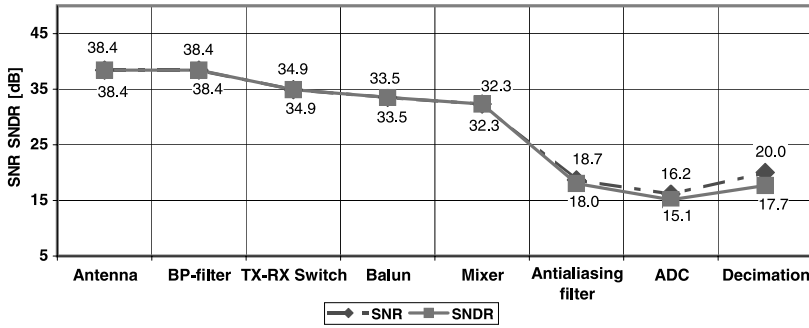


Fig. 5.24 SNDR and SNR along the receiver chain with the input signal at the sensitivity level

to correct for the spread in the position of poles and zeros of the filter due to process spread in resistances and capacitances. Nevertheless, this architecture assures very good linearity at reasonable cost in terms of power consumption.

In the evaluation of the SNDR, to avoid severe constraints on the linearity of the receiver, the SNR degradation due to the third order inter-modulation product has been fixed to less than 3 dB. This means that to keep the same SNR (20 dB) the signal cannot be at its minimum but 3 dB above. Generally in low power standards like Bluetooth the signal is allowed to be 6 dB above the sensitivity level when inter-modulation distortion is considered. This gives a way to relax further the linearity requirement of the receiver chain in case the low power target becomes difficult to achieve.

The SNR and the SNDR along the receiver chain are depicted in Fig. 5.24 when the signal is at its sensitivity level. As can be seen from Fig. 5.24, the presence of the third order inter-modulation product decreases the SNR after decimation by roughly 2.5 dB. If the signal is allowed to be 6 dB above the sensitivity level, like in the Bluetooth standard, the IP_3 requirements for the LNA-mixer block and the AA-filter can be relaxed to -7 dBm and 17 dBm respectively, which allows to reduce eventually the power consumption of these blocks.

5.5.2 Receiver Building Blocks State-of-the-Art

This section gives an overview of the state-of-the-art of some of the fundamental blocks in the receiver chain. The section focuses mainly on the following three blocks:

- Active mixer¹⁷
- AA-filter
- ADC

¹⁷It is practically a merged LNA-mixer in which the bias current is reused.

The aim of the section is to determine a circuit topology, which can achieve the block specifications mentioned in Fig. 5.17 at the minimum power level. When state-of-the-art topologies do not yet achieve the specifications at a reasonable power level, we will try to highlight a possible way to meet the specifications without increasing the power consumption too much.

Active Mixer

The receiver needs to handle low level signals. For this reason, a certain degree of signal amplification is required before downconverting it to a lower frequency. The most common way to achieve the target is to cascade an LNA and a mixer. Indeed, the LNA amplifies the signal while reducing by a few dBs¹⁸ the SNR. The mixer is generally a noisy block. Amplifying the signal before any downconversion is performed, reduces the input referred noise of the mixer by approximately the LNA gain. This allows the receiver to have a fairly low NF but it comes at the expense of a high power consumption.

To meet the low power target, the simple cascaded LNA-mixer topology is, therefore, not efficient. Looking at Fig. 5.17 we can see that the mixer NF can be as large as 10.3 dB. This gives us some margin to explore different topologies that can optimize the use of bias currents to obtain at the same time the required signal amplification and NF.

A common way to optimize the use of bias currents is known as current reuse. Following this approach in [126] a current reused LNA-mixer front-end with 31.5 dB voltage gain and 8.5 dB NF at 500 μ A current consumption (1.0 V power supply) in 0.18 μ m technology has been recently proven. The schematic block diagram is depicted in Fig. 5.25. The LNA and the g_m stage of the mixer use the same current ($I_{\text{bias-gm}}$). The capacitance C_{by} creates a good AC ground. In this way a single current source provides both the amplification via the LNA and the V - I conversion required in the mixer. To save power the design operates under a 1 V power supply. This means that it is not possible to stack too many transistors within a 1 V supply. For this reason, the mixer switches are folded and a bias current $I_{\text{bias-sw}}$ is used to bias them. This current can be made much smaller than ($I_{\text{bias-gm}}$), thus negligibly contributing to the overall power consumption. Finally, I - V conversion is performed using the resistors R .

Concluding, the LNA-mixer topology proposed in [126] achieves a larger gain than required and also a better noise figure. Therefore, though consuming only 500 μ A from a 1 V supply there is still some margin to further reduce the power consumption while meeting the specifications given in Fig. 5.17.

¹⁸In a well designed state-of-the-art LNA.

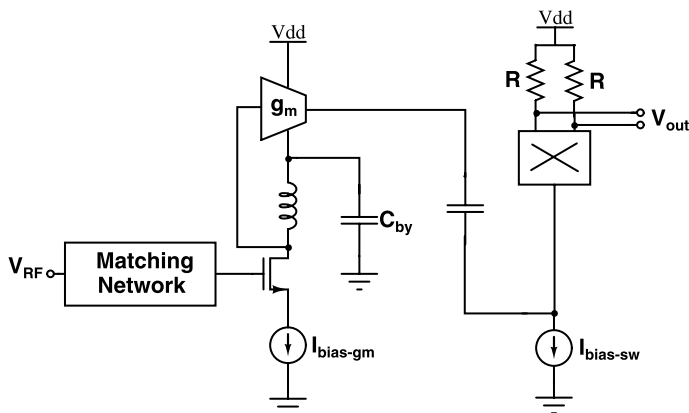


Fig. 5.25 Current-reused folded-cascode LNA-mixer (from [126])

Anti-aliasing Baseband Filter

As already mentioned previously in this section, though g_m - C filters require generally a lower power, they can become power hungry when high linearity is required. Given the required 20 dBm IIP_3 (see Fig. 5.17) we think it is more efficient to choose for an R - C active filter, where the linearity is governed by the passive elements used to implement the feedback around the OPAMPs.

Noise can be a major concern especially in zero-IF receiver topologies. Different solutions can be used in order to limit the degradation due to flicker noise at low frequencies:

- No signal energy around DC (wideband FSK)
- OPAMPs with PMOS input stage
- Large input transistors in the OPAMPs

The use of a wideband FSK modulation avoids that signal energy lies too close to DC. This partially alleviates the effect of flicker noise as well as DC offset. Unfortunately, for low data rate the frequency at which the signal is located lies not too far from DC (in the order of few tens of kilohertz) and therefore, it is still subjected to the detrimental effect of flicker noise especially in sub-100 nm technologies.

Using a PMOS input stage helps reducing the noise contribution of the OPAMPs because PMOS transistors present a lower noise than NMOS transistors. Moreover, the flicker noise is proportional to the transistor area ($W \times L$). For a given $\frac{W}{L}$ the W and the L are scaled accordingly to achieve the required transistor area that minimizes the flicker noise contribution from the OPAMPs.

Increasing the transistor length reduces the f_t of the transistor. However, this is not of much concern in baseband filters as long as the f_t remains at least ten times larger than the maximum frequency the filter needs to handle.

Though the required linearity is large, it can be achieved with less than 375 μ A per pole of the filter transfer function [127]. This value can be reduced because the

filter designed in this work has a higher gain (18 dB compared to 6 dB required) and a higher IIP_3 . However, the reported noise is better than required in the proposed receiver architecture. This definitely allows a reduction in the power consumption by properly scaling the values of the resistors around the OPAMP.

We have considered, as an example in this chapter, a signal bandwidth of around 24 kHz. To have some flexibility in the data-rate a channel filter with a 50 kHz pass-band is a good choice. In this chapter we have also shown that splitting the filter in a low-order analog filter followed by a sharp digital filtering can be considered an optimal solution from a power point of view. We set the order of the analog filter to a second order. Now given the noise and the linearity specifications shown in Fig. 5.17 and considering the work described in [127], we can conclude that a reasonable power target for the filter is 0.5 mW or less.

Analog-to-Digital Converters

In contrast to the high speed ADCs required by high demanding applications like for example digital TV streaming, sensor applications require ADCs with medium resolutions (up to roughly 12 bits) and sample rates up to a few MegaHertz. As already mentioned throughout the whole book, the most stringent specification for any integrated circuit that has to be used for WSN applications is the power consumption. Also the die area must be minimized to reduce the wireless node cost and this implies that each transceiver building block must be designed for the minimum area as well. Finally, a certain degree of flexibility is required for the ADC in order to reduce its sampling rate or accuracy to be able to save power when the maximum sampling rate or accuracy is not required. Summarizing, any ADC that aims to be used in a transceiver for WSNs must be designed according to the following characteristics:

- Low power consumption
- Small silicon area
- Low sampling frequency and medium resolution
- Programmable resolution and sampling frequency

Low power means commonly, simple hardware and current reuse when possible. Small silicon area can be achieved by reusing the same hardware. This means that, for example, flash ADCs are ruled out because they will occupy a large area especially if an 8 to 12 bit ADC is designed. Low sampling frequency means that we potentially do not need a fast architecture that make an extensive use of hardware (pipelined ADC). These complex architectures are very suitable if high sampling frequencies are required but they become a bottleneck from both silicon area and power consumption point of view if low to middle sampling frequencies are used. Finally, the architecture needs to be easily reconfigurable to achieve a good flexibility. A diagram showing the most important ADC architectures and their area of applications respect to the number of bits, the sampling frequency range and the application area is shown in Fig. 5.26.

Fig. 5.26 ADC architectures, application, resolution and sampling rates (from [128])

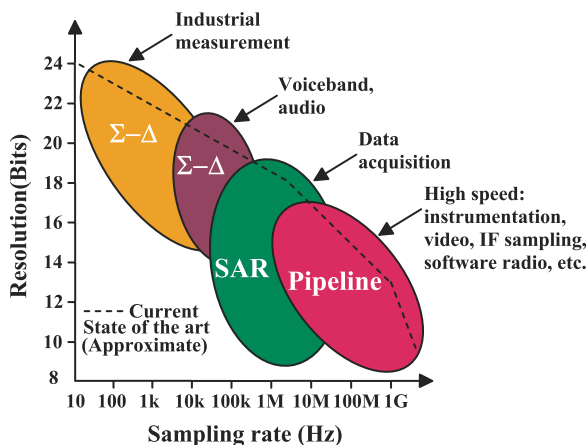
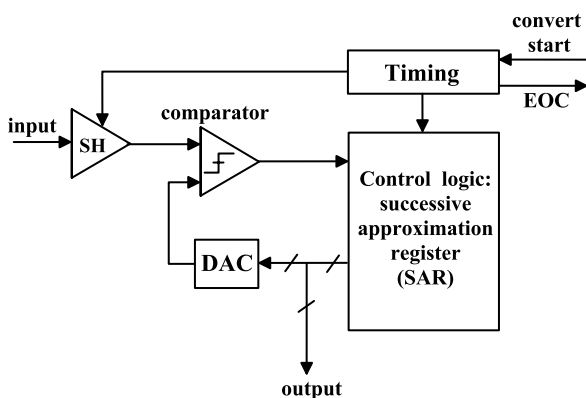


Fig. 5.27 SAR ADC architecture (from [128])



The dashed line represents the approximate state-of-the-art in 2005. From this figure we can see that pipeline ADC are mostly used for high speed applications and medium to low resolution, while the $\Sigma - \Delta$ converters are mainly used in applications requiring high to very high resolutions but low sampling rates. Given the required sampling rates and resolution the SAR ADC seems the architecture that best fits all the requirements.

The basic SAR ADC architecture is shown in Fig. 5.27.

The SH block allows to keep the signal constant during the conversion time. In this way the architecture can process rapidly varying signals, which otherwise would change the analog value to process during the conversion. This architecture uses an algorithmic process to convert the analog signal into a digital signal. To perform this operation the analog signal is compared to a reference value set by the DAC. If the analog signal is smaller than the DAC output a digital 0 is set for the bit otherwise a digital 1 is set. Generally, the conversion starts with the DAC set to the midscale value and the first output of the comparator sets the MSB value. After the value of the MSB is assessed, the DAC is set to 1/4 or 3/4 scale depending on the value

of the MSB. In this way, the second bit can be assessed looking at the comparator output. All the bits are stored in the SAR register and the process continues till all the bits are determined. At the end of the process an EOC signal is assessed.

The accuracy and the linearity of the SAR ADC mainly depend on the internal DAC. For this reason, the most common way to implement it is by using a charge redistribution DAC because the capacitance matching mainly depends on the photolithography accuracy, which is pretty high in CMOS technology. Moreover, the SAR ADC requires a number of iterations equal to the number of bits. Therefore, for a given sampling frequency, the reference clock driving the architecture has to be the number of bits faster. On the other hand, this architecture fully exploits hardware reuse achieving at the same time low power and small silicon area. Finally, programmability can be easily achieved by reducing the reference clock and therefore, reducing the sampling rate. It is also possible to reduce the ADC number of bits, trading in this way conversion speed for accuracy for a given clock frequency.

Different successful low power implementations of the SAR ADC have been reported in literature. In [129] a SAR ADC has been proposed that achieves 100 kS/s sampling frequency and consumes only 25 μ W. The designed ADC has a 12 bits resolution and achieves 10.5 bits ENOB. Flexibility is achieved by the possibility to change either the sampling frequency or the resolution. The resolution can be set to be 8 or 12 bits, while the sampling frequency can be varied between 0 and 100 kS/s. This work achieves most of the specifications required for the ADC as listed in Fig. 5.17. The sampling rate needs to be increased accordingly. An increase in the sampling frequency will linearly scale up the power consumption. Therefore, if a 200 kS/s SAR ADC is required, we can expect the power consumption to double with respect to the value given in [129]. This will lead to a 50 μ W power consumption for ADC. Given the fact that a zero-IF architecture is used, two ADC are required to correctly digitize the signal. Concluding, the mixed-signal interface between the analog RF receiver and the digital DSP will consume around 100 μ W, which is in line with the power budget of the receiver.

5.6 Simulation and Experimental Results

The frequency hopping synthesizer is implemented in a 90 nm CMOS process. The die photograph is shown in Fig. 5.28. A Serial Programmable Interface (SPI) is used to change between different hopping bins. The SPI is also implemented on-chip. The baseband synthesizer excluding the tunable LPF-notch filters is $500 \mu\text{m} \times 300 \mu\text{m}$, while the tunable LPF-notch filter occupies a die area equal to $400 \mu\text{m} \times 700 \mu\text{m}$.

We have implemented on-chip two versions of the synthesizer. The first version comprises the baseband synthesizer without the LPF-Notch filter. The second version is the combination of the baseband synthesizer plus the LPF-Notch filter. Moreover, we have implemented a test structure for the filter to be able to test its transfer characteristics as a stand alone block. Therefore, this paragraph will first describe the measurement results of the baseband stand alone block. Then, the filter measurement results will be discussed and finally measurement results for the whole system will be given.

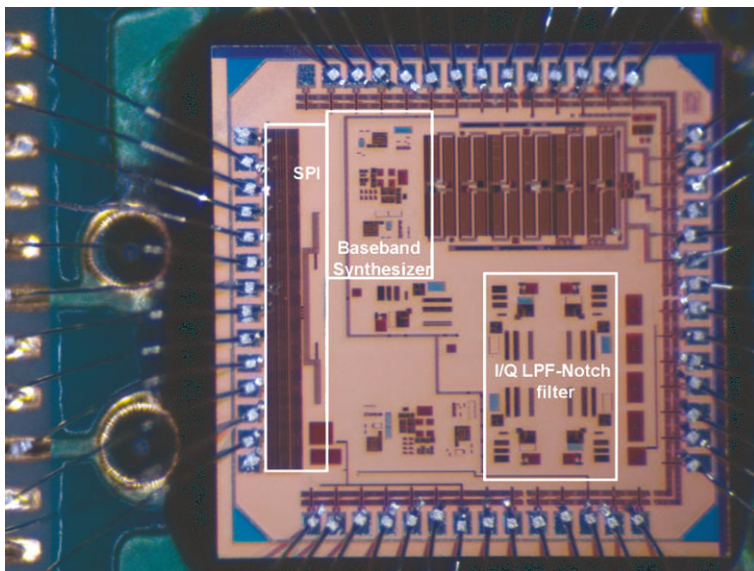


Fig. 5.28 Microphotograph of the FHSS synthesizer in 90 nm CMOS. The not highlighted parts contain a test structure for the baseband synthesizer without the LPF-notch, a single LPF-notch test structure and test-structures of a different system

5.6.1 Baseband Synthesizer Without the LP-notch Filter

In Fig. 5.29 the time waveform of a 0.5 MHz frequency, generated by the synthesizer, is shown. The analog levels with their time durations are very close to the ideal waveform described in Sect. 5.3.3. In Fig. 5.30 the spectrum of the 0.5 MHz synthesized frequency is shown. The 3rd and the 5th harmonics are rejected by more than 50 dBs while the 7th and the 9th are 17 dB and 19 dB below the fundamental, as expected.

The quadrature unbalance for the I and Q signals has also been measured for all the channel frequencies averaging over around 1000 acquisitions. The quadrature error was on the average below ± 0.3 degrees (worst case) over all frequencies. This value allows to achieve an image rejection after upconversion to the center frequency of the ISM band larger than 40 dB.

The measurement reported in Fig. 5.29 has been repeated for all the 28 synthesized channels. The level of the 3rd, 5th, 7th and 9th harmonics versus the channel frequency is reported in Fig. 5.31.

The harmonic rejection tends to lower at higher frequencies. The reason is mainly due to phase unbalances in the Walsh function based harmonic rejection block, which tend to be more visible at higher frequencies. Nevertheless, even at the highest synthesized frequency the 3rd and the 5th harmonics are rejected by more than 39 dBs. A slight filtering effect is visible for the 7th and 9th harmonics. Their amplitude is larger at lower frequencies and smaller by around 6 dB at the maximum

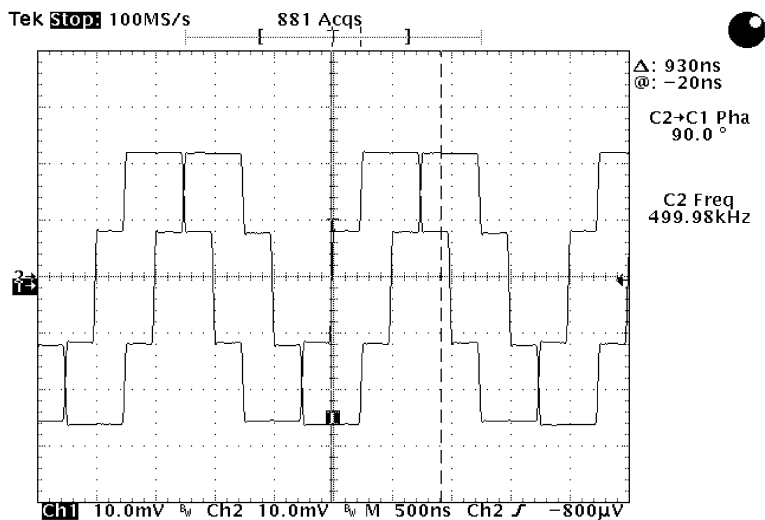


Fig. 5.29 Time waveform of a 0.5 MHz synthesized frequency

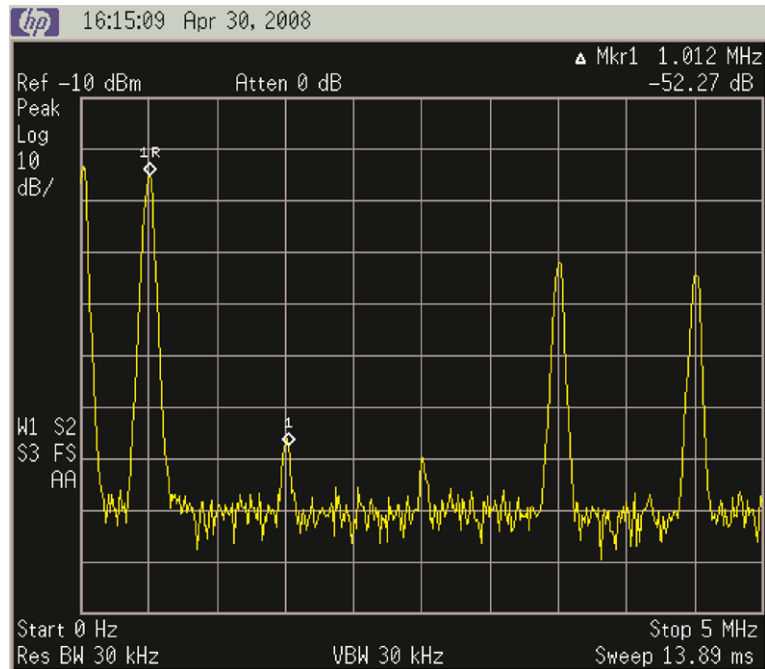


Fig. 5.30 Spectrum of a 0.5 MHz synthesized waveform

Fig. 5.31 3rd, 5th, 7th and 9th harmonic levels below the fundamental versus synthesized frequency (prior to LP-notch filter)

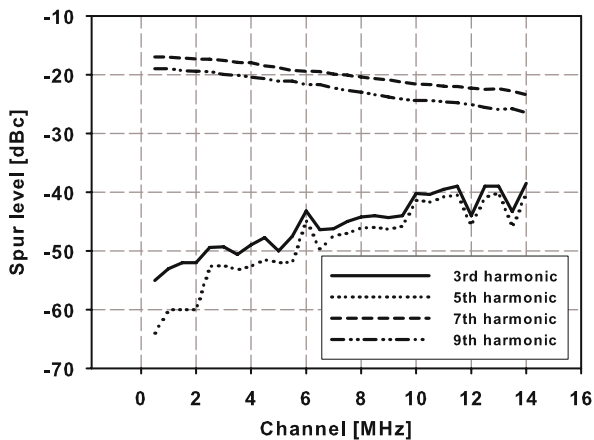
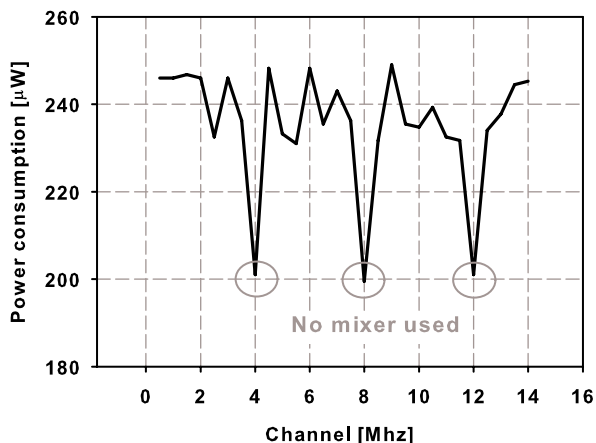


Fig. 5.32 Measured power consumption versus the synthesized frequency



operating frequency. This is due to the buffer cut-off frequency between the base-band synthesizer and the output pad.

The power consumption is constant with the synthesized channel and it is below $250\ \mu\text{W}$ from a $0.75\ \text{V}$ supply for the baseband synthesizer excluding the LP-notch filter. This is shown in Fig. 5.32. For three channels the power consumption decreases by around $50\ \mu\text{W}$. In these channels, as shown also in Table 5.2, the mixer is not used and the desired frequency is synthesized dividing the reference signal by an integer number.

5.6.2 Stand Alone Tunable LP-notch Filter

The measured filter characteristics of the tunable LP-notch filter are shown in Fig. 5.33. At high frequencies, signal coupling directly from the input to the output

can be seen. This phenomenon is connected to signal coupling on the PCB when a high frequency signal is coupled from the network analyzer via the PCB to the input of the filter. This supposition is confirmed from the measurements on the full system (baseband synthesizer plus LP-notch filter) as shown in Sect. 5.6.3. Moreover, a high pass effect is visible in the measured transfer characteristics of the filter. This is due to our measurement setup. Indeed, to operate properly, the filter requires a DC bias. On the other hand the network analyzer used to measure the filter transfer characteristics requires no DC voltage at its input. Therefore, a DC blocker needs to be used. The cut-off frequency of such a DC blocker starts at 3 MHz and therefore, measurements at frequencies below 3 MHz present such a high pass behavior.

This high-pass effect has been corrected in the measured filter attenuation shown in Fig. 5.34. A first order roll-off has been supposed in the correction formula used to derive Fig. 5.34. It can be seen that the worst case extra attenuation is larger than 20 dB (Figs. 5.33 and 5.34). This attenuation is calculated within the ISM bandwidth that is 83.5 MHz for the 2.4 GHz ISM band. Therefore, the maximum synthesized frequency that has the 7th harmonic falling inband is 12 MHz, while the maximum synthesized frequency having the 9th harmonic falling inband is the 9.5 MHz. These two frequencies set the borders where the SFDR specification must be met. Finally, the extra attenuation values in Fig. 5.34 have been obtained tuning the filter with the pre-defined settings for each channel frequency. If the channel settings of the filter are optimized for the maximum possible attenuation of the unwanted harmonics, around 6 dB extra attenuation can be obtained as shown in Table 5.4.

Finally, as explained in Sect. 5.3.3 it is important to assess the harmonic distortion introduced by the buffer used to decouple the LPF from the notch filter. The 3rd order harmonic distortion of the LP-notch filter versus the synthesized frequency is plotted in Fig. 5.35.

As it can be seen from Fig. 5.35, the 3rd order distortion is below 45 dBc and therefore, does not affect the system. The filter power consumption comes from the decoupling buffer placed between the LPF and the notch filter (see Fig. 5.14). This buffer consumes 25 μ A from a 0.75 V supply. The number of required buffers is 4

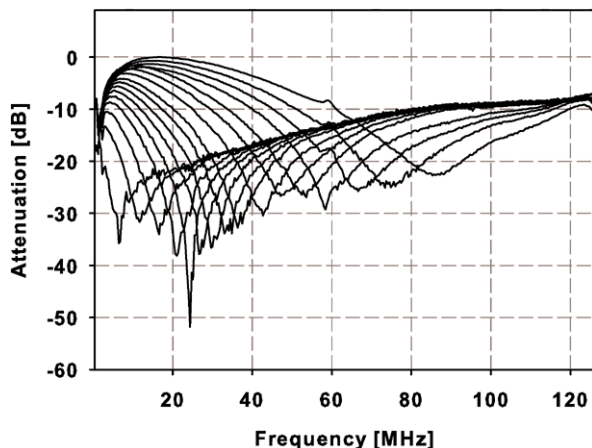


Fig. 5.33 LP-notch filter transfer characteristics

Fig. 5.34 Measurement results of the stand alone LP-notch filter

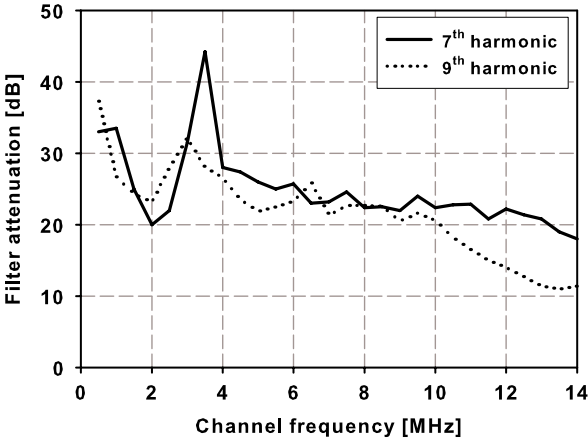
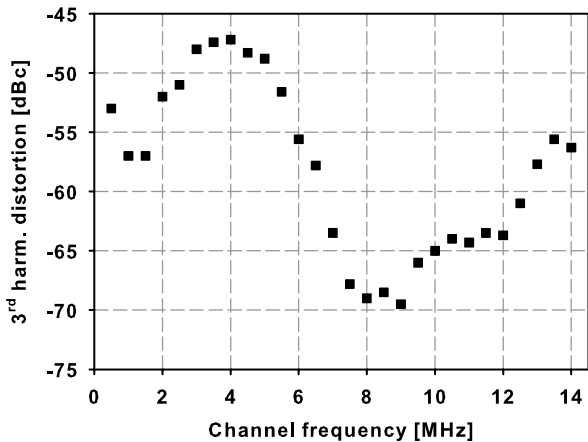


Fig. 5.35 3rd order LP-notch filter harmonic distortion



because of the differential quadrature signals used and therefore, the tunable LPF-notch filter consumes 75 μ W.

5.6.3 Complete Frequency Hopping Baseband Synthesizer

The transient waveforms for a 2 MHz synthesized frequency at the output of the complete frequency hopping synthesizer are shown in Fig. 5.36. As it can be seen from Fig. 5.36 the even order harmonic is greatly reduced and an SFDR larger than 44 dB is obtained. Moreover, the effect of the LPF is visible also at higher frequencies, confirming that for the filter stand-alone measurements shown in Fig. 5.33 we were limited by PCB coupling at high frequencies.

The quadrature signals for the combination of synthesizer and LP-notch filter are shown in Fig. 5.37 for a synthesized frequency of 2 MHz. The quadrature error is

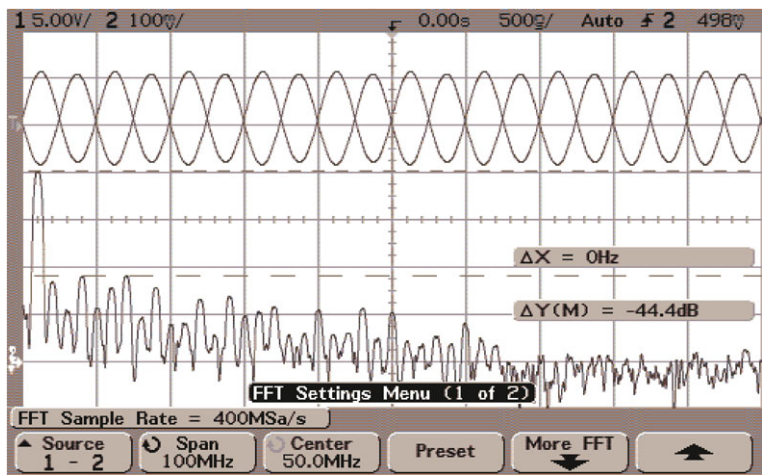


Fig. 5.36 Spectrum of the differential signal output for a synthesized frequency of 2 MHz

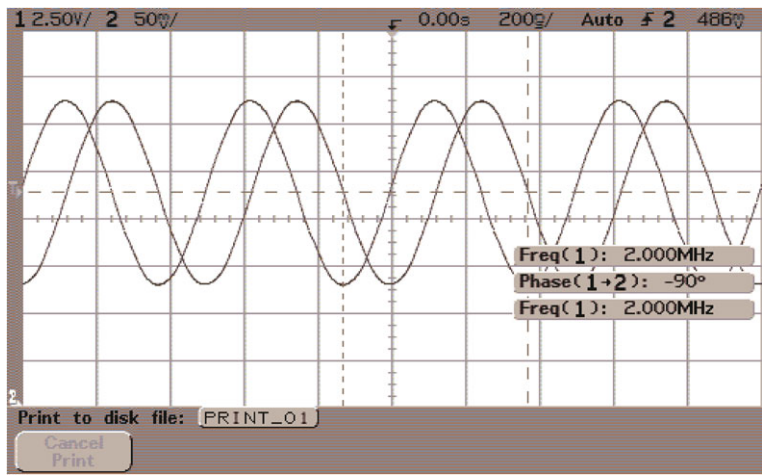


Fig. 5.37 Quadrature unbalance for a 2 MHz output signal

negligible at low frequencies and it goes up to a maximum of ± 0.3 degrees for a 14 MHz synthesized frequency.

The hopping speed of the synthesizer is limited by the internal time constants of the silicon implementation. Simulations show that a hopping speed larger than 500 khop/s is achievable without, in first approximation, any increase in the power consumption.¹⁹ Measurement results confirmed the very high speed of the synthesizer as shown in Fig. 5.38. The time required to change from a 1 MHz initial fre-

¹⁹The slight increment in the power consumption is due to the SPI working at a higher frequency.

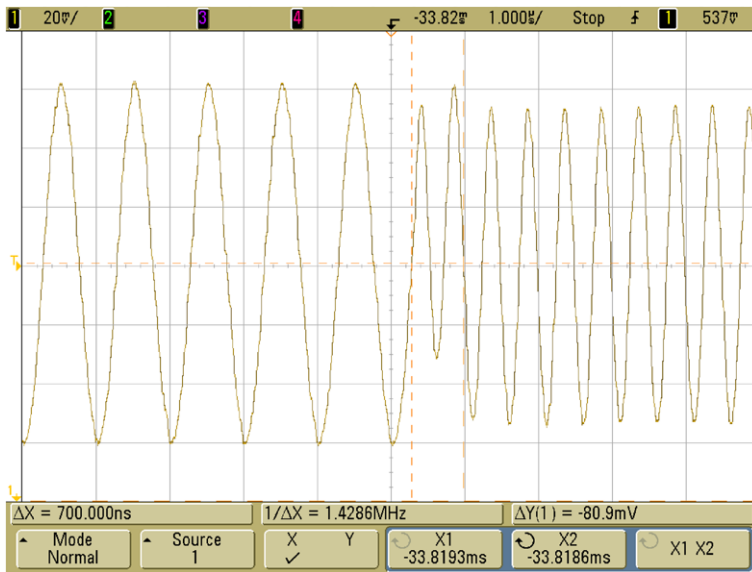


Fig. 5.38 Measured system hopping speed. Frequency is changed from 1 MHz to 2 MHz

quency to a 2 MHz final frequency is around 700 ns, which allows for an hopping speed larger than 1 Mhop/s.

5.6.4 Benchmarking

The system performance is summarized in Table 5.4. The baseband frequency hopping synthesizer meets all the requirements listed in Sect. 5.3.1. This allows to integrate the baseband synthesizer into a frequency hopping wireless node improving the transmitter efficiency and reducing also the receiver power consumption.

In Table 5.5 the performance of the proposed baseband frequency hopping synthesizer is compared with state-of-the-art baseband synthesizers. These synthesizers are based on a DDFS architecture and do require an upconversion stage similarly to the proposed baseband synthesizer architecture. All the referred synthesizers do not include the upconversion stage and therefore, only the power consumption of the baseband part is considered. The proposed solution compares very well against all the state-of-the-art solutions shown in Table 5.5 in terms of power consumption per MegaHertz. The estimated power consumption in Table 5.5 is calculated as follows. For every design listed in Table 5.5, the power consumption per MegaHertz is calculated. Then, we suppose that to synthesize a maximum frequency $f_{\max}$, the DDFS clock frequency should be at least three times $f_{\max}$. Therefore, the estimated power consumption is the power consumption per MegaHertz times $3 \times f_{\max}$. We compared our work with the state-of-the-art works listed in Table 5.5 both in terms of power per MegaHertz and predicted power to synthesize frequencies up to 14 MHz.

Table 5.4 Proposed frequency hopping synthesizer performance summary

Hopping range	≈ 5	Octaves
Number of channels	56	–
Channel separation	0.5	MHz
3rd harmonic	< -42	dBc
5th harmonic	< -54	dBc
7th harmonic	< -44	dBc
9th harmonic	< -45	dBc
Power consumption	325	μ W
Supply voltage	0.75	V
SFDR	> 42	dBc

The best design from literature in terms of power per MegaHertz is the one disclosed in [57]. Even though its power consumption per MegaHertz is already an order of magnitude higher than the one achieved by our proposed solution, its power consumption refers only to the digital core. When the power consumption of the DAC and of the image filter are considered, the overall synthesizer power consumption will increase further. Moreover, the solution in [57] cannot provide a quadrature output. It achieves a 60 dB SFDR, but the frequency at which the SFDR is measured is not mentioned. The SFDR can change a lot versus frequency as it can be seen in [133]. The proposed baseband frequency hopping synthesizer achieves better than 42 dBc of SFDR over the whole frequency range up to a 14 MHz synthesized frequency.

When looking at the estimated power consumption, the proposed solution again performs better than the state-of-the-art solutions listed in Table 5.5. Again the work that achieves the lowest power consumption is the work proposed in [57]. However, this is a poor comparison because that design does not have a quadrature output, and it does not include the required DAC and image filter. Nevertheless, the proposed solution has a power consumption that is half the one of the work in [57] while giving an analog quadrature output. The work disclosed in [133], which contains a DAC has an estimated power consumption an order of magnitude higher than the one in the proposed baseband frequency hopping synthesizer. Moreover, above 100 kHz its SFDR drops dramatically and it does not have a quadrature output. The work disclosed in [130], on the other hand, has a good SFDR up to higher frequencies, but its estimated power consumption is around 7 times higher than the proposed solution and it does not include the image filter power consumption. All the other works listed in Table 5.5 have an estimated power consumption one or more orders of magnitude larger than the solution proposed in this book. We can, therefore, conclude that the proposed baseband frequency hopping synthesizer allows to reduce the power consumption, both in terms of power per MegaHertz and

Table 5.5 Proposed frequency hopping synthesizer performance summary

Reference	Power/MHz [mW/MHz]	Estim. Power [mW]	DAC incl.	SFDR [dBc]	Quadrature out	Technology
[130]	0.05	2.1	yes	55	no	0.35 μ m CMOS
[57]	0.015	0.63	no	60	no	0.18 μ m CMOS
[131]	0.8	33.6	yes	51	no	0.25 μ m CMOS
[132]	0.4	16.8	no	90	no	0.25 μ m CMOS
[133]	0.08	3.4	yes	51	no	0.5 μ m CMOS
[134]	4	168	yes	52	no	0.8 μ m BiCMOS
[55]	0.32	13.4	no	58	no	0.8 μ m CMOS
[135]	1.33	55.9	yes	90.3	yes	0.25 μ m CMOS
This work	0.0011	0.325	–	> 42	yes	90 nm CMOS

estimated power consumption,²⁰ by an order of magnitude. This target is achieved while still keeping an SFDR better than 42 dBc over the whole frequency range.

5.7 Conclusions

In this chapter we have seen that the constraints of a two-way link wireless node are different with respect to the constraints for a one-way link wireless nodes described in Chap. 4. The new constraints required a novel hopping synthesizer design with respect to the one described in Chap. 4. Given the frequency accuracy requirements for a successful demodulation of BFSK modulated signals an accurate reference signal has been considered. This increases cost and form factor but these drawbacks are compensated by the absence of a network infrastructure, like was used for the one-way link described in Chap. 4.

The major obstacle to the use of frequency hopping spread spectrum in power constrained systems is the hopping synthesizer. In this Chapter we described a novel frequency hopping synthesizer that can meet at the same time the required agility and frequency accuracy of a DDFS but at much lower power consumption. We compared the proposed solution with state-of-the-art DDFS based synthesizers. This comparison shows that the proposed baseband frequency hopping synthesizer reduces the power consumption per MegaHertz, as well as the overall power consumption, by an order of magnitude compared to state-of-the-art solutions. Moreover, the hopping synthesizer achieves an SFDR better than 42 dBc over the whole synthesized frequency range (0.5 MHz to 14 MHz), and outputs a quadrature signal. The last feature is not very common among DDFS based synthesizers because it generally further increases the power consumption of the synthesizer.

Concluding, with a measured power consumption of 325 μ W and a settling time better than 700 ns the proposed frequency hopping baseband synthesizer can be used as a core for a fast hopping autonomous transceiver for wireless sensor networks.

²⁰For synthesized frequencies up to 14 MHz.

Chapter 6

Summary and Conclusions

Wireless sensor networks have the potential to become the third wireless revolution after wireless voice networks in the 80s and wireless data networks in the late 90s. This revolution will finally connect together the physical world of the human and the virtual world of the electronic devices. Though in the recent years large progress in power consumption reduction has been made in the wireless arena in order to increase the battery life, this is still not enough to achieve a wide adoption of this technology. Indeed, while nowadays consumers are used to charge batteries in laptops, mobile phones and other high-tech products, this operation becomes infeasible when scaled up to large industrial, enterprise or home networks composed of thousands of wireless nodes.

Wireless sensor networks come as a new way to connect electronic equipments reducing, in this way, the cost associated with the installation and maintenance of large wired networks. A network of 2000 nodes, where each node requires replacing the battery every 10 years, would require the replacement of a battery every two days. Therefore, to open the door of the third wireless revolution it is necessary to overcome the maintenance problem.

To accomplish this task, it is necessary to reduce the energy consumption of the wireless node to a point where energy harvesting becomes feasible and the node energy autonomy exceeds the life time of the wireless node itself. Besides this requirement, cost represents another important requirement for wireless sensor networks. Installing a wireless sensor network must be cheaper than installing an ad-hoc wired network. Therefore, small form factor and short bill of materials are important requirements as well. Finally, maintenance cost must be zero and this requires the use of unlicensed ISM frequency bands for data transmission. This brings us to the last requirement for a wireless sensor network: the robustness of the link which can be very challenging in an interference overcrowded scenario like the one present in the ISM bands. Therefore, in this book we tried to put the basis toward the design of a wireless sensor node that can fulfill the aforementioned requirements:

- Robustness
- Autonomous operation
- Low cost

We have analyzed a broad range of applications that can benefit from the adoption of a wireless sensor network. Given the very broad range of applications we have divided them into two big sub-classes

- One-way link networks
- Two-way link networks

These two classes present unique challenges and do require careful optimization at system and circuit level. We have showed that similarities do exist at system level between the two scenarios and we demonstrated that a viable, low-cost way to have a robust link passes through the use of spread spectrum techniques. Among different spread spectrum techniques we have chosen a frequency hopping spread spectrum system because it showed, globally, a higher potential of power down-scaling when carefully optimized. Frequency modulation formats and instantaneous data rates have been also studied in order to assess the best system level choices for an overall system power optimization. A BFSK modulation format together with a low (1 to 10 kbps) data rate are the most suitable choices in a power constrained scenario. We have demonstrated that the PN code synchronization overhead connected with any spread spectrum system can be minimized designing a fast-hopping system.

In this book we proposed a novel approach to the design of ultra-low power radios for WSNs. We started from a proven robust frequency hopping spread spectrum system like Bluetooth and we removed all the features, which are not required for the wireless sensor network scenario trying to meet the required ultra-low power operation.

State of the art frequency hopping systems are still power hungry and therefore, cannot be readily used in an autonomous wireless system. Our study showed that the power consumption bottleneck in a frequency hopping based transceiver resides in the frequency hopping synthesizer. The combination of frequency agility, accuracy and reasonably low spur levels require severe trade-offs which undermine the low power operation. A power analysis of currently available solutions for frequency hopping synthesizers has been carried out. The scope of this study was to evaluate if any of the current architectural solutions could achieve the required combination of frequency agility and accuracy at a power level low enough to embed the synthesizer in an autonomous wireless node. Power models of PLL based synthesizers and DDFS based synthesizers showed that the power targets cannot be achieved using current technologies even if the synthesizer specifications are scaled down to fit the requirements of a wireless sensor network.

This study suggested that a different approach to the frequency synthesis is required in order to design a frequency hopping spread spectrum system for wireless sensor networks. The two different wireless scenarios required, at this point, a different approach, exploiting the differences between them in order to minimize the system power consumption and the cost of the wireless network.

The one way link is an asymmetric scenario in which the transmitter node is power constrained while the receiver node is mains supplied and, therefore, has a virtually unlimited power budget. To squeeze the power consumption in the transmitter, its architecture has been kept simple, not very accurate and with very simple

calibration steps to be performed prior to any transmission. This translated in a crystal-less transmitter, that employs a direct modulation of a high-frequency VCO using frequency pre-distortion. The frequency pre-distortion allows having a linear frequency grid of transmitting channels for an optimal occupation of the available frequency spectrum. The low accuracy of the reference signal (1%) is good enough to fulfill the FCC rules during transmission. All the other required synchronization steps (TX-RX frequency alignment, PN code synchronization, etc.) are performed at the receiver side. With this architecture, a wireless link using the 915 MHz ISM band with a raw BER smaller than 1.1%, has been demonstrated in an indoor environment and NLOS condition. The use of a bipolar technology combined with the availability of high- Q passives allowed to reduce the current consumption of the transmitter below 2.5 mW from a 2 V power supply and to keep the transmitter node cost very low. Though many applications will benefit from this scenario, not all possible applications can be covered by a one-way link. In many cases, the deploying of an infrastructure (the RGs) can be cumbersome, if not impossible.

A two way-link scenario is meant to cover all those applications in which an infrastructure cannot be deployed. In this scenario both the transmitter and the receiver are power constrained and therefore, a novel approach is required in order to optimize the power consumption and the cost of the wireless sensor network.

The first difference, with respect to the one-way link scenario, is the presence of a crystal-based reference signal which constitutes the pulsing heart of the transceiver. This increases the wireless node cost but this extra cost can be compensated by the fact that no infrastructure is required to support the network. Though all the transceiver blocks require careful optimization given the small available power budget, the bottleneck of this system in terms of power consumption is, again, the hopping synthesizer. This time the synthesis of the frequency bins is not performed directly at high frequency like in the one-way link scenario. Differently, in the two-way link scenario the channels are synthesized at baseband and then upconverted to high frequencies. This technique allows decoupling the high frequency operation required for wireless transmission from the required frequency agility.

The baseband synthesizer employs a novel architecture that considerably reduces the power consumption with respect to state-of-the-art frequency hopping baseband synthesizers. With a power consumption of 325 μ W from a 0.75 V supply, a settling time smaller than 700 ns and an SFDR larger than 42 dBc up to 14 MHz the proposed synthesizer, designed in a standard CMOS 90 nm process, tops state-of-the-art baseband synthesizers by more than a order of magnitude in terms of power consumption.

The realized prototype shows for the first time the feasibility of designing a fast frequency hopping spread spectrum system for wireless sensor networks. Though for this scenario, more challenges need to be solved in order to demonstrate a low power, robust and low cost transceiver, one of the biggest obstacles to the implementation of a fast hopping spread spectrum technique into a wireless sensor network has been solved. This opens a new scenario for wireless sensor networks in which it is not required anymore to trade-off between link robustness and power consumption.

Concluding, in this book, we have demonstrated that there is no unique architecture that is able to cover the wide range of applications foreseen for wireless sensor networks. By exploiting the unique characteristics of a particular scenario for wireless sensor networks and by a careful optimization of different parameters at system and circuit level we demonstrated that link robustness and low power are not necessarily exclusive. For both scenarios we have demonstrated that a fast frequency hopping system can be designed with a reasonable power consumption. This opens a new path in the design of radio systems for wireless sensor networks in which robustness is guaranteed by techniques that were previously exclusively used in radio systems for middle or high end applications like Bluetooth and military communications.

Appendix

Walsh Based Harmonic Rejection Sensitivity Analysis

The silicon implementation of both the passive interpolator as well as of the clock tree driving the interpolator is never ideal. The two most important unbalances are the mismatch in the interpolator resistances and the incorrect phase relation between the interpolator input signals. These unbalances can adversely affect the cancellation of the third and fifth harmonic in the Walsh shaper. In the appendix we focus on the effect of the aforementioned non-idealities on the cancellation of the third harmonic. The same procedure can be repeated for the fifth harmonic. In this way it is possible to evaluate the sensitivity of the 3rd and 5th harmonic rejection to the imperfections of a real silicon implementation. This information can be further used to optimize the design and reduce the effect of the unavoidable non-idealities in the silicon implementation.

In the following analysis we use the simplified model shown in Fig. A.1. In the model, ϕ represents the phase unbalance while ϵ is the resistor mismatch with respect to its ideal value. The fundamental tone can be expressed by the following equation

$$S_{\text{fund}} = \sqrt{2} \cos(\omega_0 t) + \cos(\omega_0 t + \phi_0) + \cos(\omega_0 t - \phi_0) \quad (\text{A.1})$$

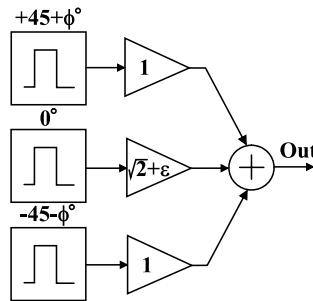
where ϕ_0 is the phase difference between the signals applied to the passive interpolator. Applying simple trigonometric relations, the previous equation can be simplified into the following equation:

$$S_{\text{fund}} = \sqrt{2} \cos(\omega_0 t) + 2 \cos(\omega_0 t) \cos(\phi_0) \quad (\text{A.2})$$

Now we can suppose that the interpolator resistances are mismatched by an amount equal to ϵ and therefore, the $\sqrt{2}$ weighted resistance has a weight equal to $\sqrt{2} + \epsilon$ and we suppose that no phase error is present. This means that $\phi_0 = \frac{\pi}{4}$. Equation (A.2) becomes:

$$S_{\text{fund}} = (2\sqrt{2} + \epsilon) \cos(\omega_0 t) \quad (\text{A.3})$$

Fig. A.1 Simplified model to evaluate the harmonic rejection sensitivity to resistor mismatches and signal phase unbalances



The same procedure can be applied to the third harmonic of the signals applied to the passive interpolator. The third harmonic can be written after the summation in Fig. A.1 as:

$$S_{3rd} = \frac{1}{3}(\sqrt{2}\cos(3\omega_0 t) + \cos(3\omega_0 t + 3\phi_0) + \cos(3\omega_0 t - 3\phi_0)) \quad (A.4)$$

and after some trigonometric manipulations

$$S_{3rd} = \frac{1}{3}(\sqrt{2}\cos(3\omega_0 t) + 2\cos(3\omega_0 t)\cos(3\phi_0)) \quad (A.5)$$

Supposing a mismatch in the interpolator resistances

$$S_{3rd} = \frac{\epsilon}{3}\cos(3\omega_0 t) \quad (A.6)$$

At this point it is possible to calculate the 3rd order harmonic rejection:

$$HR3_A = \frac{2\sqrt{2} + \epsilon}{\epsilon/3} \approx \frac{6\sqrt{2}}{\epsilon} \quad (A.7)$$

The same procedure can be applied when a phase error is present. No resistor mismatch is considered but $\phi_0 = \frac{\pi}{4} + \varphi$. In this case (A.2) can be written as follows:

$$S_{fund} = \sqrt{2}\cos(\omega_0 t)(1 + \cos\varphi - \sin\varphi) \quad (A.8)$$

while (A.5) becomes:

$$S_{3rd} = \frac{\sqrt{2}}{3}\cos 3\omega_0 t (1 - \cos 3\varphi - \sin 3\varphi) \quad (A.9)$$

Finally, the 3rd order harmonic rejection, in the presence of a phase error between the signals applied to the interpolator, can be written as follows:

$$HR3_\phi = 3 \frac{(1 + \cos\varphi - \sin\varphi)}{(1 - \cos 3\varphi - \sin 3\varphi)} \quad (A.10)$$

Supposing that the phase error is small we can expand the cosine and the sine functions in Taylor series as:

$$\cos x \approx 1 - \frac{x^2}{2} \quad (A.11)$$

$$\sin x \cong x - \frac{x^3}{6} \quad (\text{A.12})$$

Therefore, (A.10) becomes

$$HR3_\phi = \frac{2}{\varphi} \quad (\text{A.13})$$

References

1. K. Martinez, R. Ong, J.K. Hart, J. Stefanov, GLAC-SWEB: a sensor web for glaciers, in *First European Workshop on Wireless Sensor Networks* (2004)
2. R. Riem-Vis, Minimum sideband noise in oscillators, in *Workshop on Applications of Mobile Embedded Systems* (2004)
3. F. Michahelles, P. Matter, A. Schmidt, B. Schiele, Applying wearable sensors to avalanche rescue, in *Computers and Graphics* (2003), pp. 839–847
4. Z. Butler, P. Corke, R. Peterson, D. Rus, Networked cows: virtual fences for controlling cows, in *Workshop on Applications of Mobile Embedded Systems* (2004)
5. J. Burrell, T. Brooke, R. Beckwith, Vineyard computing: sensor networks in agricultural production, in *IEEE Pervasive Computing* (2004), pp. 38–45
6. *Nanoenergy—The Power for Micro Devices*, <http://www.frontedgetechnology.com/>, Front-edge Technology
7. J.P. Thomas, M.A. Qidwai, J.C. Kellog, Energy scavenging for small-scale unmanned systems. *J. Power Sources* **159**, 1494–1509 (2006)
8. J.A. Paradiso, T. Starner, Energy scavenging for mobile and wireless electronics. *IEEE Pervasive Comput.* **34**, 18–27 (2005)
9. V. Raghunathan et al., Design considerations for solar energy harvesting wireless embedded systems, in *4th International Symposium on Information Processing in Sensor Networks* (2005)
10. B. Leung, *VLSI for Wireless Communication* (Prentice Hall, Upper Saddle River, 2002)
11. P.G.M. Baltus, R. Dekker, Optimizing RF front ends for low power. *Proc. IEEE* **88**(10), 1546–1559 (2000)
12. J. Polastre, R. Szewczyk, D. Culler, Telos: enabling ultra-low power wireless research, in *4th International Symposium on Information Processing in Sensor Networks* (2005), pp. 364–369
13. C. Park, J. Liu, P.H. Chou, Eco: an ultra-compact low-power wireless sensor node for real-time motion monitoring, in *4th International Symposium on Information Processing in Sensor Networks* (2005), pp. 398–403
14. B.P. Otis et al., An ultra-low power MEMS-based two-channel transceiver for wireless sensor networks, in *Symposium on VLSI Circuits* (2004), pp. 20–23
15. C. Enz, A. El-Hoiydi, J.-D. Decotignie, V. Peiris, Optimizing RF front ends for low power. *Computers* **37**, 62–70 (2004)
16. D.E. Duarte, Clock network and phase-locked loop power estimation and experimentation. Ph.D. thesis, Penn State University (2002)
17. P. Heydari, A study of low-power ultra wideband radio transceiver architectures, in *Wireless Communications and Networking Conference*, Mar. 2005, pp. 758–763
18. I.D. O'Donnell, R.W. Brodersen, An ultra-wideband transceiver architecture for low power, low rate, wireless systems. *IEEE Trans. Veh. Technol.* **54**, 1623–1631 (2005)

19. R.H.T. Bates, G.A. Burrell, Towards faithful radio transmission of very wideband signals. *IEEE Trans. Antennas Propag.* **AP-20**, 684–690 (1972)
20. J. Ryckaert et al., Ultra-wide-band transmitter for low-power wireless body area networks: design and evaluation. *IEEE Trans. Circuits and Syst. I* **52**, 2515–2525 (2005)
21. L. Stoica et al., An ultrawideband system architecture for tag based wireless sensor networks. *IEEE Trans. Veh. Technol.* **54**, 1632–1645 (2005)
22. J.-P. Curty et al., Remotely powered addressable UHF RFID integrated system. *IEEE J. Solid-State Circuits* **40**, 2193–2202 (2005)
23. R.G. Vaughan, N.L. Scott, D.R. White, The theory of bandpass sampling. *IEEE Trans. Signal Process.* **39**, 1973–1984 (1991)
24. M.R. Yuce, W. Liu, A low-power multirate differential PSK receiver for space applications. *IEEE Trans. Veh. Technol.* **54**, 2074–2084 (2005)
25. B. Otis, Y.H. Chee, J. Rabaey, A 400 μ W-RX, 1.6 mW-TX super-regenerative transceiver for wireless sensor networks, in *IEEE International Solid-State Circuits Conf. (ISSCC)*, Feb. 2005, pp. 396–397
26. A. Vouilloz, M. Declercq, C. Dehollain, A low-power CMOS super-regenerative receiver at 1 GHz. *IEEE J. Solid-State Circuits* **36**, 440–451 (2001)
27. A. Kamerman, Spread-spectrum techniques drive WLAN performance. *Microwaves RF Sept.*, 109–114 (1996)
28. R.E. Ziemer, R.L. Peterson, D.E. Borth, *Introduction to Spread Spectrum Communications* (Prentice Hall, Upper Saddle River, 1995)
29. J.D. Oetting, A comparison of modulation techniques for digital radio. *IEEE Trans. Commun.* **27**, 1752–1762 (1979)
30. E. Lopelli, J.D. van der Tang, A.H.M. van Roermund, A 1 mA ultra-low-power FHSS TX front-end utilizing direct modulation with digital pre-distortion. *IEEE J. Solid-State Circuits* **42**, 2212–2223 (2007)
31. E. Lopelli, J. van der Tang, A. van Roermund, Ultra-low power frequency-hopping spread spectrum transmitters and receivers, in *15th Workshop on Advances in Analog Circuit Design*, Apr. 2006
32. B. Razavi, RF transmitter architectures and circuits, in *IEEE Custom Integrated Circuits Conf. (CICC)*, May 1999, pp. 197–204
33. L. Yang, C. Shuqing, R.W. Dutton, Numerical investigation of low frequency noise in MOS-FETs with high- k gate stacks, in *International Conference on Simulation of Semiconductor Processes and Devices*, Sept. 2006, pp. 99–102
34. C.A. Putman, S.S. Rappaport, D.L. Shilling, A comparison of schemes for coarse acquisition of frequency-hopped spread-spectrum signals. *IEEE Trans. Commun.* **31**, 183–189 (1983)
35. A. Rofougaran et al., A single-chip 900-MHz spread-spectrum wireless transceiver in 1- μ m CMOS—part II: receiver design. *IEEE J. Solid-State Circuits* **33**, 535–547 (1998)
36. H. Komurasaki et al., A 1.8-V operation RFCMOS transceiver for Bluetooth, in *Symposium on VLSI Circuits* (2002), pp. 230–233
37. Y.-J. Jung et al., A 2.4-GHz 0.25- μ m CMOS dual-mode direct conversion transceiver for Bluetooth and 802.11b. *IEEE J. Solid-State Circuits* **39**, 1185–1190 (2004)
38. T. Byunghak Cho et al., A 2.4-GHz dual-mode 0.18- μ m CMOS transceiver for Bluetooth and 802.11b. *IEEE J. Solid-State Circuits* **39**, 1916–1926 (2004)
39. A. Zolfaghari, B. Razavi, A low-power 2.4-GHz transmitter/receiver CMOS IC. *IEEE J. Solid-State Circuits* **38**, 176–183 (2003)
40. S. Byun et al., A low-power CMOS Bluetooth RF transceiver with a digital offset cancelling DLL-based GFSK demodulator. *IEEE J. Solid-State Circuits* **38**, 1609–1618 (2003)
41. P. van Zeijl et al., A Bluetooth radio in 0.18- μ m CMOS. *IEEE J. Solid-State Circuits* **37**, 1679–1687 (2002)
42. S.-C. Kim, L. Bertoni, M. Stern, Pulse propagation characteristics at 2.4 GHz inside buildings. *IEEE Trans. Veh. Technol.* **45**, 579–592 (1996)
43. C.C. Enz, N. Scolari, U. Yodprasit, Ultra low-power radio design for wireless sensor networks, in *International Workshop on Radio-Frequency Integration Technology*, Nov. 2005, pp. 1–17

44. D. Theil et al., A fully integrated CMOS frequency synthesizer for Bluetooth, in *IEEE Radio Frequency Integrated Circuits Symposium (RFIC)* (2001), pp. 103–106
45. M. Marletta et al., Fully integrated fractional PLL for Bluetooth application, in *IEEE Radio Frequency Integrated Circuits Symposium (RFIC)* (2005), pp. 557–560
46. S.-Y. Lee et al., A 1-V 2.4-GHz low-power fractional- N frequency synthesizer with sigma-delta modulator controller, in *IEEE International Symposium on Circuits and Systems (ISCAS)* (2005), pp. 2811–2814
47. W. Wang, H.C. Luong, A 0.8-V 4.9-mW CMOS fractional- N frequency synthesizer for RFID application, in *European Solid-State Circuits Conf. (ESSCIRC)* (2006), pp. 146–149
48. H. Huh et al., A CMOS dual-band fractional- N synthesizer with reference doubler and compensated charge pump, in *IEEE International Solid-State Circuits Conf. (ISSCC)* (2004), pp. 100–101
49. K. Woo, Y. Liu, D. Ham, Fast-locking hybrid PLL synthesizer combining integer & fractional divisions, in *Symposium on VLSI Circuits* (2007), pp. 1260–1261
50. S. Willingham et al., An integrated 2.5 GHz $\Sigma\Delta$ frequency synthesizer with 5 μ s settling and 2 Mb/s closed loop modulation, in *IEEE International Solid-State Circuits Conf. (ISSCC)* (2000), pp. 200–201
51. W.-Z. Chen, D.-Y. Yu, A dual-band four-mode $\Sigma\Delta$ frequency synthesizer, in *IEEE Radio Frequency Integrated Circuits Symposium (RFIC)* (2006), pp. 11–13
52. J.M.P. Langlois, D. Al-Khalili, Phase to sinusoid amplitude conversion techniques for direct digital frequency synthesis. *IEE Proc. Circuits Devices Syst.* **151**, 519–528 (2004)
53. J.S. Min, H. Samuelli, Analysis and design of a frequency-hopped spread-spectrum transceiver for wireless personal communications. *IEEE Trans. Veh. Technol.* **49**, 1719–1731 (2000)
54. S. Morteza pour, E.K.F. Lee, Design of low-power ROM-less direct digital frequency synthesizer using nonlinear digital-to-analog converter. *IEEE J. Solid-State Circuits* **34**, 1350–1359 (1999)
55. A. Bellaouar et al., Low-power direct digital frequency synthesis for wireless communications. *IEEE J. Solid-State Circuits* **35**, 385–390 (2000)
56. A.N. Mohieldin, A.A. Emira, E. Sanchez-Sinencio, A 100-MHz 8-mW ROM-less quadrature direct digital frequency synthesizer. *IEEE J. Solid-State Circuits* **37**, 1235–1243 (2002)
57. J.M.P. Langlois, D. Al-Khalili, Low power direct digital frequency synthesizers in 0.18 μ m CMOS, in *IEEE Custom Integrated Circuits Conf. (CICC)*, Sept. 2003, pp. 283–286
58. B.-D. Yang et al., An 800-MHz low-power direct digital frequency synthesizer with an on-chip D/A converter. *IEEE J. Solid-State Circuits* **39**, 761–774 (2004)
59. Y. Song, B. Kim, A 14-b direct digital frequency synthesizer with sigma-delta noise shaping. *IEEE J. Solid-State Circuits* **39**, 847–851 (2004)
60. D.D. Caro, A.G.M. Strollo, High-performance direct digital frequency synthesizers in 0.25 μ m CMOS using dual-slope approximation. *IEEE J. Solid-State Circuits* **40**, 2220–2227 (2005)
61. F.F. Dai et al., A direct digital frequency synthesizer with fourth-order phase domain $\Delta\Sigma$ noise shaper and 12-bit current-steering DAC. *IEEE J. Solid-State Circuits* **41**, 839–850 (2006)
62. E. Lopelli, J. van der Tang, A. van Roermund, A FSK demodulator comparison for ultra-low power, low data-rate wireless links in ISM bands, in *ECCTD 2005*, vol. 2, Sept. 2005, pp. 259–262
63. A technical tutorial on digital signal synthesis, <http://www.analog.com/en/> (1999). Analog Devices
64. J. Vanka, Direct digital synthesizers: theory, design and applications. Ph.D. thesis, Helsinki University of Technology, Helsinki, Finland (2000)
65. R. van de Plassche, *CMOS Integrated Analog-to-Digital and Digital-to-Analog Converters* (Kluwer Academic, Dordrecht, 2003)
66. G. Groenewold, Optimal dynamic range integrated continuous-time filters. Ph.D. thesis, Delft University of Technology, Delft, The Netherlands (1992)

67. M.A.T. Sanduleanu, Power, accuracy and noise aspects in CMOS mixed-signal design. Ph.D. thesis, Twente University of Technology (1999)
68. D. Liu, C. Svensson, Power consumption estimation in CMOS VLSI chips. *IEEE J. Solid-State Circuits* **29**, 663–670 (1994)
69. W. Lei, Y. Fukatsu, K. Watanabe, Characterization of current-mode CMOS R – $2R$ ladder digital-to-analog converters. *IEEE Trans. Instrum. Meas.* **50**, 1781–1786 (2001)
70. Y.P. Tsividis, J.O. Voorman, *Integrated Continuous-Time Filters: Principles, Design and Applications* (IEEE Press, Piscataway, 1993)
71. R.L. Oliveira Pinto, M.C. Schneider, C. Galup-Montoro, Sizing of MOS transistors for amplifier design, in *IEEE International Symposium on Circuits and Systems (ISCAS)*, vol. 4, May 2000, pp. 185–188
72. P. Kinget, M. Steyaert, Impact of transistor mismatch on the speed-accuracy-power trade-off of analog CMOS circuits, in *IEEE Custom Integrated Circuits Conf. (CICC)*, May 1996, pp. 333–336
73. C.C. Enz, F. Krummenacher, E.A. Vittoz, An analytical MOS transistor model valid in all regions of operation and dedicated to low-voltage and low-current applications. *Analog Integr. Circuits Signal Process.* **8**, 83–114 (1995)
74. D.S. Karadimas, K.A. Efstathiou, An R – $2R$ ladder-based architecture for high linearity DACs, in *Instrumentation and Measurement Technology Conference Proceedings*, May 2007, pp. 1–5
75. S.-W. Lee, H.-J. Chung, C.-H. Han, C – $2C$ digital-to-analogue converter on insulator. *IEEE Electron. Lett.* **35**, 1242–1243 (1999)
76. E. Lopelli, J. van der Tang, A. van Roermund, A sub-mA frequency synthesizer technique, in *IEEE Radio Frequency Integrated Circuits Symposium (RFIC)*, June 2006, pp. 495–498
77. E. Lopelli, J. van der Tang, A. van Roermund, A frequency offset recovery algorithm for crystal-less transmitters, in *Personal Indoor and Mobile Radio Communications Symposium*, Sept. 2006
78. A. Pouttu, J. Juntti, Performance studies of slow frequency hopping M -ary FSK with concatenated codes in partial band noise jamming, in *Military Communication Conference*, Nov. 1995, pp. 335–339
79. P.E. Chadwick, Design compromise in frequency synthesizers, in *Proceedings of the RF Technology Expo*, Anaheim, US (CA), Jan. 1986
80. J. van der Tang, S. Hahn, A monolithic 0.4 mW SOA LC VCO, in *European Solid-State Circuits Conf. (ESSCIRC)*, Sept. 1999, pp. 150–153
81. S. Finocchiaro, G. Palmisano, R. Salerno, C. Scalfani, Design of bipolar RF ring oscillators, in *International Conference on Electronics, Circuits and Systems*, Sept. 1999, pp. 5–8
82. R. Navid, T.H. Lee, R.W. Dutton, Lumped, inductorless oscillators: how far can they go?, in *IEEE Custom Integrated Circuits Conf. (CICC)*, Sept. 2003, pp. 543–546
83. M. Alioto, G. di Cataldo, G. Palumbo, Design of low-power high-speed bipolar frequency dividers. *Electron. Lett.* **38**, 158–160 (2002)
84. D. Kasperkovitz, D. Grenier, Traveling-wave dividers: a new concept for frequency division. *Microelectron. Reliab.* **16**, 127–134 (1977)
85. X. Wang, A. Fard, P. Andreani, Phase noise analysis and design of a 3-GHz bipolar differential colpitts VCO, in *European Solid-State Circuits Conf. (ESSCIRC)*, Sept. 2005, pp. 391–394
86. T.M. Nowatski, K.I. Zambrano, Method for automatically compensating for accuracy degradation of a reference oscillator, US5,552,749, June 1995
87. I. Minako, Apparatus for detecting frequency offset, EP1128620, Aug. 2001
88. F. Martin et al., Toward wireless receivers without crystals, in *IEEE Radio Frequency Integrated Circuits Symposium (RFIC)*, June 2005
89. E. Lopelli, J. van der Tang, Frequency detection method, NL1029668, Aug. 2005
90. H.M. Elissa et al., ARCTAN differentiated digital demodulator for FM/FSK digital receivers, in *The 2002 45th Midwest Symposium on Circuits and Systems*, vol. 2 (2002), pp. 200–203
91. J. Min et al., Low power correlation detector for binary FSK direct-conversion receivers. *Electron. Lett.* **31**, 1030–1032 (1995)

92. M. Saitou et al., Direct conversion receiver for 2- and 4-level FSK signals, in *Fourth IEEE International Conference on Universal Personal Communications* (1995), pp. 392–396
93. S. Hara et al., A novel FSK demodulation method using short-time DFT analysis for LEO satellite communication systems. *IEEE Trans. Veh. Technol.* **46**, 625–633 (1997)
94. H.M. Kwon, K.B.E. Lee, A novel digital FM receiver for mobile and personal communications. *IEEE Trans. Commun.* **44**, 1466–1476 (1996)
95. T.T. Tjhung et al., Error rates for narrow-band digital FM with discriminator detection in mobile radio systems. *IEEE Trans. Commun.* **38**, 999–1005 (1990)
96. S. Hinedi et al., The performance of noncoherent orthogonal M-FSK in the presence of timing and frequency errors. *IEEE Trans. Commun.* **43**, 922–933 (1995)
97. PIC16F627A Data Sheet, <http://www.microchip.com> (2007). Microchip
98. CoolFlux DSP the embedded ultra low power C-programmable DSP core. <http://www.coolfluxdsp.com> (2008). NXP semiconductor
99. AD7392 Data sheet, <http://www.analog.com/> (2003). Analog Devices
100. AD5341 Data sheet, <http://www.analog.com/> (2003). Analog Devices
101. Y.H. Chee, A.M. Niknejad, J.M. Rabaey, An ultra-low power injection locked transmitter for wireless sensor networks. *IEEE J. Solid-State Circuits* **41**, 1740–1748 (2006)
102. A. Molnar et al., An ultra-low power 900 MHz RF transceiver for wireless sensor networks, in *IEEE Custom Integrated Circuits Conf. (CICC)*, Oct. 2004, pp. 401–404
103. B.W. Cook et al., An ultra-low power 2.4 GHz RF transceiver for wireless sensor networks in 0.13 μm CMOS with 400 mV supply and an integrated passive RX front-end, in *IEEE International Solid-State Circuits Conf. (ISSCC)*, Feb. 2006, pp. 1460–1469
104. P. Choi et al., An experimental coin-sized radio for extremely low-power WPAN (IEEE 802.15.4) application at 2.4 GHz. *IEEE J. Solid-State Circuits* **38**, 2258–2268 (2003)
105. H.-M. Seo et al., A fully CMOS integrated transceiver for ubiquitous networks in sub-GHz ISM band, in *Asia-Pacific Conference on Communications*, Oct. 2005, pp. 700–704
106. I. Kwon et al., A fully integrated 2.4-GHz CMOS RF transceiver for IEEE 802.15.4, in *IEEE Radio Frequency Integrated Circuits Symposium (RFIC)*, June 2006, pp. 275–282
107. CC1100 Datasheet, <http://focus.ti.com/docs/prod/folders/print/cc1100.html>
108. CC2500 Datasheet, <http://focus.ti.com/docs/prod/folders/print/cc2500.html>
109. ADF7020 Datasheet, <http://www.analog.com/en/>
110. SX1223 Datasheet, <http://www.semtech.com>
111. H. Tanaka, T. Ieki, Y. Hirano, Y. Ishikawa, SAW oscillator multi-chip module for 300 MHz low power radio, in *IEEE International Frequency Control Symposium*, May 1997, pp. 836–840
112. B. Otis, J.M. Rabaey, A 300 μW 1.9 GHz CMOS oscillator utilizing micromachined resonators. *IEEE J. Solid-State Circuits* **38**, 1271–1274 (2003)
113. Y.H. Chee, A.M. Niknejad, J. Rabaey, A sub-100 μW 1.9-GHz CMOS oscillator using FBAR resonator, in *IEEE Radio Frequency Integrated Circuits Symposium (RFIC)* (2005), pp. 123–126
114. E. Lopelli et al., A 0.75 V 325 μW 40 dB-SFDR frequency-hopping synthesizer for wireless sensor networks in 90 nm CMOS, in *IEEE International Solid-State Circuits Conf. (ISSCC)*, Feb. 2009, pp. 228–229
115. J.A. Weldon et al., A 1.75-GHz highly integrated narrow-band CMOS transmitter with harmonic-rejection mixers. *IEEE J. Solid-State Circuits* **36**, 2003–2015 (2001)
116. A. SenGupta, F.L. Walls, Effect of aliasing on spurs and PM noise in frequency dividers, in *Frequency Control Symposium and Exhibition*, July 2000, pp. 541–548
117. V. Geffroy, G. De Astis, E. Bergeault, RF mixers using standard digital CMOS 0.35 μm process, in *International Microwave Symposium IEEE MMT-S*, vol. 1, May 2001, pp. 83–86
118. F. Behbahani, Y. Kishigami, J. Leete, A.A. Abidi, CMOS mixers and polyphase filters for large image rejection. *IEEE J. Solid-State Circuits* **36**, 873–887 (2001)
119. J.L. Walsh, A closed set of normal orthogonal functions. *Am. J. Math.* **45**, 5–24 (1923)
120. R. Kitai, Synthesis of periodic sinusoids from Walsh waves. *IEEE Trans. Instrum. Meas.* **24**, 313–317 (1975)

121. A. Bateman, Transmitter and receiver architectures, <http://www.avren.com/Courses>
122. A.H. Aly, D.W. Beishline, B. El-Sharawy, Filter integration using on-chip transformers, in *Microwave Symposium Digest IEEE MTT-S Digest*, vol. 3, June 2004, pp. 1975–1978
123. Y.-C. Wu, M.F. Chang, On-chip spiral inductors and bandpass filters using active magnetic energy recovery, in *IEEE Custom Integrated Circuits Conf. (CICC)*, May 2002, pp. 275–278
124. D. Shim et al., Ultra-miniature monolithic FBAR filters for wireless application, in *Microwave Symposium Digest IEEE MTT-S Digest*, June 2005, pp. 213–216
125. B. Murmann, ADC performance survey 1997–2009, <http://www.stanford.edu/~murmman/adcsurvey.html> (2009)
126. T. Song et al., A 2.4-GHz sub-mW CMOS receiver front-end for wireless sensors network. *IEEE Microw. Wirel. Compon. Lett.* **16**, 206–208 (2006)
127. H.A. Alzahr, H.O. Elwan, M. Ismail, A CMOS highly linear channel-select filter for 3G multistandard integrated wireless receivers. *IEEE J. Solid-State Circuits* **37**, 27–37 (2002)
128. Which ADC architecture is right for your application?, <http://www.analog.com/>, June 2005, Analog Dialogue
129. N. Verma, A.P. Chandrakasan, A 25 μ W 100 kS/s ADC for wireless microsensor application, in *IEEE International Solid-State Circuits Conf. (ISSCC)*, Feb. 2006, pp. 222–223
130. A. McEwan, S. Shah, S. Collins, A direct digital frequency synthesis system for low power communications, in *European Solid-State Circuits Conf. (ESSCIRC)*, Sept. 2003, pp. 393–396
131. J. Jiang, E.K.F. Lee, A ROM-less direct digital frequency synthesizer using segmented non-linear digital-to-analog converter, in *IEEE Custom Integrated Circuits Conf. (CICC)*, May 2001, pp. 165–168
132. X. Li et al., A direct digital frequency synthesizer based on two segment fourth-order parabolic approximation. *IEEE Trans. Consum. Electron.* **55**, 322–326 (2009)
133. A.N. Mohieldin, A. Emira, E. Sanchez-Sinencio, A 100 MHz, 8 mW ROM-less quadrature direct digital frequency synthesizer. *IEEE J. Solid-State Circuits* **37**, 1235–1243 (2002)
134. J. Vankka et al., A direct digital synthesizer with an on-chip D/A-converter, in *European Solid-State Circuits Conf. (ESSCIRC)*, Sept. 1997, pp. 216–219
135. A. Torosyan, D. Fu, A.N. Willson Jr., A 300-MHz quadrature direct digital synthesizer/mixer in 0.25- μ m CMOS. *IEEE J. Solid-State Circuits* **38**, 875–887 (2003)

Index

A

- AA-filter, 74
 - power consumption, 86
- Acquisition time, 28, 29
- ADC, 196, 205
 - architectures, 206
 - ENOB, 196, 198, 199
 - FOM, 199
 - full scale range, 196, 198
 - noise figure, 196
 - noise power, 196
 - SAR, 206, 207
 - SNR, 196, 197
- Additive White Gaussian Noise (AWGN), 27
- Antenna, 156, 157
 - characteristic gain, 12, 20
 - microstrip patch, 156–158
- AWGN, 30, 107, 145

B

- Back-scattering, 21
- Baseband filter, 105
- Battery
 - size estimation, 10
- BER, 31, 101, 108, 145, 147–149, 164
- Bluetooth, 51, 110, 202

C

- Channel capacity, 20
- Charge pump
 - current consumption, 66, 67
 - leakage current, 66
- Code acquisition
 - matched filter, 47, 48, 50
 - two-level, 47, 49, 50
- Correlator, 47, 48
 - active, 49

- non-coherent, 146
- passive, 48

- Crystal, 5, 190
 - oscillator, 112

D

- DAC, 72, 155, 206, 207
 - charge redistribution, 85
 - INL, 72, 102
 - number of bits, 72, 101
 - power consumption, 78, 86
 - quantization error, 101
 - R -2 R , 79
 - SFDR, 73
 - SNR, 72
- Data rate, 31, 33–35
- DDFS
 - low-pass filter, 59
 - phase accumulator, 59
 - power consumption, 61, 85
 - power model, 70
 - SFDR, 72, 90
 - specifications, 70
- Demodulation
 - algorithm, 143
 - FSK, 107
- Demodulator
 - ADM, 144
 - architectures, 143
 - correlation, 144, 145
 - DCDM, 144, 145
 - ST-DFT, 144, 147
- Direct Sequence Spread-Spectrum, 23
 - processing gain, 23
 - transmitter, 23
- Duty-cycle, 8, 9, 33

E

Energy scavenging, 8
techniques, 7

F

Federal Communication Commission
(FCC), 19, 51, 55

Filter

AA-filter, 73, 204
BPF, 192, 193
LP-notch, 179, 187, 188, 210, 212
SAW, 193
twin-T notch, 188, 189

Forward Error Correction (FEC), 108
Reed-Solomon, 108

Frequency, 99, 143

Frequency divider, 123, 124, 126
divide-by-three, 183
programmable, 173, 182
traveling-wave, 125

Frequency Hopping Spread Spectrum, 23

acquisition algorithm, 50
acquisition time, 48, 49
code acquisition, 46
efficiency, 54
frequency synthesizer, 98, 171
products, 52, 168
synchronization, 45
synthesizer, 54
transmitter, 24, 112, 130

Frequency hopping synthesizer
baseband synthesizer, 172, 179

Frequency offset, 102, 103, 134, 136, 138,
139, 144, 145, 147–150
acquisition, 135
acquisition time, 136, 138, 139
recovery algorithm, 137
recovery circuitry, 135
sign recovery, 142

Frequency pre-distortion, 97–98, 105, 143
algorithm, 98, 143, 148, 162

Frequency synthesizer, 173
DDFS based, 58, 61
harmonic rejection, 175
phase noise, 108
PLL based, 56, 58
SFDR, 174, 175, 178

H

Harmonic rejection, 176–178, 186, 187, 208
Walsh functions based, 176, 178, 183,
185

Harmonic rejection sensitivity, 223, 224

I

Impulse radio, 20

Inductors, 14

Inter-modulation distortion, 200

L

Link budget analysis, 192
discrete parts, 192
integrated parts, 193
propagation, 192

M

Microstrip line, 189

Mixer, 179
conversion gain, 179
passive, 179
SSB, 179–181, 185

Modulation, 31, 165

BFSK, 101, 107
BPSK, 25, 30, 31, 165
formats, 29, 31
FSK, 30, 31, 165, 170
data rate, 32
detector, 145
index, 31, 148, 193
OOK, 30, 31, 165
power efficiency, 24, 28, 29
scheme, 19, 30

Modulator

FSK, 169, 170

N

Near-far sensitivity, 26

NF, 194–196

P

Path loss

exponent, 12
NLOS, 12

Phase accumulator

number of bits, 60
power consumption, 75

PLL, 190

bandwidth, 67
baseband, 69

PLL (*cont.*)

- charge pump, 62, 64
- fractional- N , 57
- frequency divider, 65
- integer- N , 56
- lock-in time, 65
- loop filter, 56, 57, 62, 64
- PFD, 57, 65
- phase noise, 56, 57
- power model, 62, 65
- SFDR, 66

PN synchronization, 95

Power amplifier (PA)

- power efficiency, 52

Q

Quality factor

- filter, 193
- inductor, 14
- oscillator, 113
- switch, 119

R

Receiver

- architecture, 39, 133, 191
- direct conversion, 135
- low-IF, 41
- SNDR, 202
- SNR, 195, 202
- super-heterodyne, 40
- zero-IF, 39

ROM, 58–60, 71, 72, 97, 100, 103

- power consumption, 75

S

Sensitivity

- receiver, 12

Silicon on Anything (SOA), 14, 114, 116

Spread-Spectrum Systems, 22

Single-Sideband (SSB), 169

Sub-sampling, 21

- phase noise, 22

Super-capacitor, 9

Super-regenerative, 22

Synchronization

- time, 13, 28

T

Transmitter

- architecture, 36, 94, 112, 130, 169

- efficiency, 54, 165

- offset PLL, 38

- specifications, 106

V

VCO, 62, 111, 112

- colpitts, 114, 131, 132

- differential, 63

- direct modulation, 170

- gain, 63

- phase noise, 113–116, 121–123, 128, 160

- power consumption, 63

- power-VCO, 114, 130, 131, 160

- SAW stabilized, 170

- start-up time, 160, 161