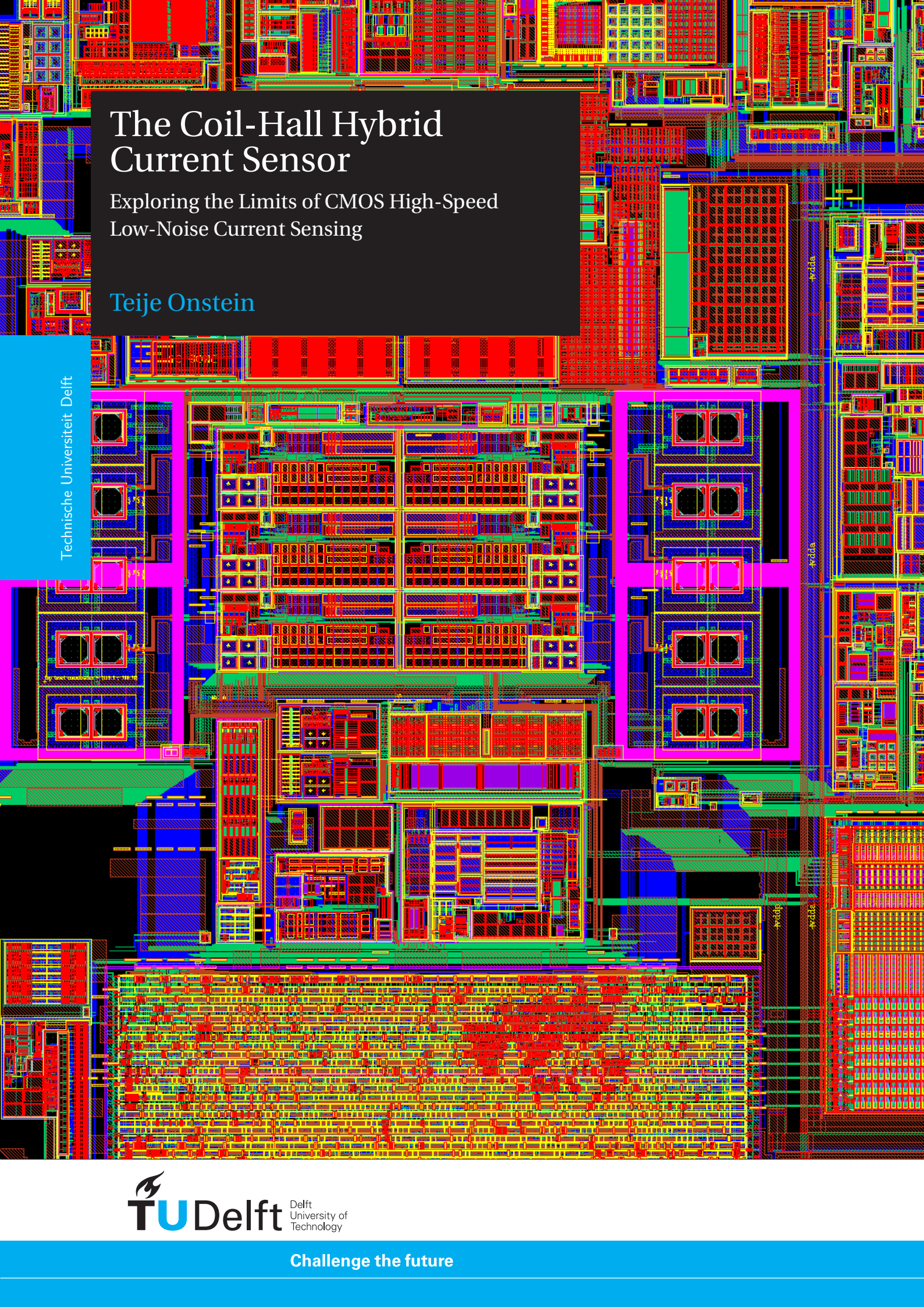




The Coil-Hall Hybrid Current Sensor

Exploring the Limits of CMOS High-Speed
Low-Noise Current Sensing

Teije Onstein

Technische Universiteit Delft

THE COIL-HALL HYBRID CURRENT SENSOR

EXPLORING THE LIMITS OF CMOS HIGH-SPEED
LOW-NOISE CURRENT SENSING

by

Teije Onstein

in partial fulfillment of the requirements for the degree of

Master of Science

in Electrical Engineering

at Delft University of Technology,
to be defended publicly on Tuesday July 23, 2024 at 2:00 PM.

Supervisors:	Ir. Dr. C.J.M. Verhoeven,	TU Delft
	Ir. A.J.M. Montagne,	TU Delft
	Dr. Ir. G. van der Horn,	SystematIC Design B.V.
	Ir. A. Jouyaeian	SystematIC Design B.V.
Thesis Committee:	Ir. Dr. C.J.M. Verhoeven,	TU Delft
	Ir. A.J.M. Montagne,	TU Delft
	Dr. Ir. Q. Fan,	TU Delft
	Dr. Ir. G. van der Horn	SystematIC Design B.V.

ACKNOWLEDGEMENTS

I would like to express my gratitude to the many people who have helped me over the course of this thesis project and my academic career so far. You have all helped me reach my current level of competence in one way or another. For this I'll be forever grateful

First of all, I would like to thank Chris Verhoeven and Anton Montagne for introducing me to the Structured Electronics Design methodology and for providing guidance throughout this project. I will never forget the exciting moments when you provided me with new insights that elevated my design skills to a new level. But I will certainly also not forget the fun conversation in between.

Furthermore, I could not have undertaken this journey without the help of everyone at SystematIC Design B.V. In particular, I would like to thank Gert van der Horn and Richard Visee for allowing me to do my thesis project at their fantastic company. Besides, I would also like to thank Gert for his daily supervision and the many clear explanations. Special thanks to Amir Jouyaeian for always being available to answer my questions and always making time for the many fruitful discussions. I will not forget the many moments when your head peaked past the monitor when I asked one of my many questions. Many thanks to Chris van den Bos for providing guidance during my first project at SystematIC. I really appreciate that you were always available to answer my questions; you taught me many things. I also want to thank Lars Bouman for being a fantastic "office mate" and great sparring partner, as well as Ciska van 't Hof for the many wonderful events, baked pancakes, decorated offices, and, in general, your great contribution to the amazing atmosphere at SystematIC. Of course also many, many thanks to everyone else at SystematIC. Vincent, Sanja, Samra, Ron, Roger, Mert, Mengfei, Lorenzo, Joep, Jelle, Frank, Elias, Britt, Bert, Baldo, and John, thanks for always being available to help and providing me with a wonderful year.

Furthermore, I would like to express my appreciation to my friends at university, Floris, Soji, Jorit, David, Carli, Nadine, Arthur, and Mees. Thanks for the many interesting and fun discussions. These have helped me see problems from a new perspective and contributed to my development as an Electrical Engineer. However, some of the discussions will also be a great contribution to the to-be-written "Long Answers to Stupid Questions." I hope many more of these discussions will follow.

And last but certainly not least, I would like to thank my family—it goes without saying—thanks for always being there for me.

*Teije Onstein
Delft, July 2024*

ABSTRACT

With the rising demand for high-bandwidth, high-resolution current sensors, commonly used Hall-effect devices fall short due to their relatively high wide-band noise. The coil-Hall hybrid architecture addresses this issue by combining the Hall plate with a pick-up coil, known for its high SNR at high frequencies. This CMOS-compatible architecture achieves excellent noise and bandwidth performance while maintaining the ability to sense DC signals. However, this performance comes at the cost of increased complexity in combining, calibrating, and stabilizing the two distinct signal paths.

This thesis thoroughly investigates the various challenges and trade-offs associated with this architecture and provides a comprehensive overview of potential system solutions. Additionally, the insights gained were used to design a prototype chip for SystematIC Design B.V. Seeking to reduce complexity, this led to the invention of the so-called "hybrid-hybrid" architecture. This new architecture effectively overcomes several challenging trade-offs that have hindered existing designs until now and is considered a promising approach for achieving even better performance in the future.

CONTENTS

1	Introduction	1
1.1	Thesis Goal	2
1.2	Thesis Scope & Outline	2
2	Coil-Hall Hybrid Current Sensor: Important Design Aspects	4
2.1	Core Functionality	4
2.2	Performance Aspects	4
2.2.1	Noise.	5
2.2.2	Input Range & Dynamic Range.	5
2.2.3	Bandwidth	5
2.2.4	Gain Stability.	5
2.2.5	Gain Flatness over Frequency	5
2.2.6	Offset	5
2.2.7	Switching Ripple	5
2.3	Cost Factors.	6
2.3.1	Power Consumption	6
2.3.2	Chip Area	6
2.3.3	Manufacturing Cost: Calibration Time.	6
2.3.4	Design Complexity.	6
2.4	System Requirements	7
3	Signal Path Combining: General Considerations	8
3.1	Transducer Models	8
3.1.1	Hall Plate Model	8
3.1.2	Coil Model	9
3.2	Why a Multipath Architecture?	10
3.3	Ideal Transfer Function	11
3.4	Noise Calculations	13
3.5	Effect of an Unwanted Pole	14
3.6	Coil Voltage versus Current Readout	15
3.6.1	Reading out the Coil Voltage	15
3.6.2	Reading out the Coil Current.	16
3.6.3	Prototype Chip Design Choice	16
4	Signal Path Combining: Implementations	18
4.1	Architecture 1: Shared Low-Pass Filter	19
4.1.1	Ripple versus Noise Performance	19
4.1.2	Offset Performance	20
4.2	Architecture 2: Integrator and High-Pass Filter	23
4.2.1	Second-Order Summing Filter	23
4.2.2	Integrator DC Gain.	25
4.2.3	Integrator Implementation Constraints and Challenges	26
4.3	Architecture 3: "Hybrid - Hybrid" Architecture	30
4.3.1	A Brief Remark on the Implementation	30
4.3.2	Noise Performance.	30
4.3.3	Gain Flatness	32
4.3.4	Offset and Clipping	33
4.4	Conclusion on the Combining Architectures	34

5	Gain Calibration	36
5.1	Sources of Gain Error	37
5.2	Foreground Calibration Procedure	38
5.3	Foreground Calibration Strategies	39
5.4	Conclusion on Foreground Calibration & Strategy Selection	41
6	Gain Stabilization	42
6.1	Gain Stabilization: Feedforward Compensation.	43
6.1.1	Feedforward Compensation in the Low-Frequency Hall Path	43
6.1.2	Feedforward Compensation in the High-Frequency Coil Path	44
6.2	Gain Stabilization: Background calibration	45
6.2.1	Signal Orthogonality	46
6.2.2	Gain Measurement Accuracy.	49
6.2.3	Interference Rejection: Spread Spectrum Signaling	50
6.3	Conclusion on Gain Stabilization & Strategy Selection	52
7	System Implementation & Verification	53
7.1	Transducer Parameters	54
7.2	Summing Filter Implementation	55
7.3	High-Frequency Path Readout Architecture.	55
7.4	Choosing the Component Sizes	56
7.5	Controller Implementation	59
7.5.1	Controller Requirements.	59
7.5.2	Controller Architecture	60
7.6	High-Frequency Path Performance Verification.	61
7.6.1	Bandwidth & Gain Flatness	62
7.6.2	Noise.	62
7.6.3	Gain Stability.	63
7.6.4	Offset & Output Swing	65
7.6.5	Coil Clipping & Slow Settling	65
7.6.6	Power- & Area Consumption	66
7.7	Conclusion & Future Design Recommendations	68
8	Conclusion	71
A	Prototype Chip Test Plan	72
A.1	Chip Startup & Entering Test Mode	72
A.2	Hall Plate Calibration & Matching.	73
A.3	High-Frequency Gain Trimming	73
A.4	Bandwidth Measurement	74
A.5	Noise Performance & Spinning Ripple	74
A.6	Frequency Response Measurement & Gain Flatness	74
A.7	Gain Drift over Temperature	75
A.8	Gain Drift over Lifetime	75
B	Coil Model	77
B.1	Plots	79
C	Dimensioning of the Coil	88
C.1	Voltage Readout.	88
C.2	Current Readout	90
D	On-Chip Calibration Coil	92
	Bibliography	94

NOMENCLATURE

A_C	DC gain in the high-frequency (coil) path
A_H	DC gain in the low-frequency (Hall) path
A_{co}	Coupling factor from current to magnetic field at the Hall plate (subscript h) or Coil (subscript c)
A_{sys}	The system gain from input current to output voltage
C_C	Coil capacitance
C_z	Feedback capacitance in the TIA
I_B	Hall plate bias current
I_N	Input referred integrated noise current
$I_{N,opt}$	Optimum input referred integrated noise current (associated with $f_{x,opt}$)
L_C	Coil inductance
R_C	Coil resistance
R_H	Hall plate resistance
R_p	Brute-force resistance at the input of the TIA
R_z	Feedback resistance in the TIA
$R_{C,n}$	Equivalent noise resistance at the input of the high-frequency path
$R_{H,n}$	Equivalent noise resistance at the input of the low-frequency path
S_C	Coil sensitivity
S_H	Hall plate sensitivity
S_I	Hall plate current-related sensitivity
S_V	Hall plate voltage-related sensitivity
T	Temperature
V_B	Hall plate bias voltage
γ	Parameter to indicate how much f_x deviates from $f_{x,opt}$ ($f_x = \gamma f_{x,opt}$)
ω_u	Integrator unity gain frequency
ζ	Parameter to indicate how much the parasitic pole p_p is distanced from p_x ($p_p = \zeta p_x$)
f_x	Cross-over frequency of the main summing filter ($p_x = 2\pi f_x$)
f_y	Frequency of a second pole if a second-order summing filter is used

f_z	Cross-over frequency of the hybrid sensor inside the hybrid-hybrid architecture high-frequency path ($p_z = 2\pi f_z$)
$f_{x,opt}$	Optimum cross-over frequency for noise
k	Boltzmann constant
p_p	Parasitic pole frequency (in radians)

1

INTRODUCTION

Current sensing plays a vital role in control, monitoring, protection, and power management in various applications such as switched-mode power converters, photovoltaic inverters, motor drivers, and battery management systems, to name a few. In many applications, relatively simple and cheap shunt-based sensors are employed for their excellent accuracy and noise performance [1, 2]. However, in high-voltage and high-power applications, the lack of galvanic isolation poses a major drawback. Magnetic field-based current sensors, on the other hand, do not have this issue; namely, by sensing the magnetic field generated by the current, the sensitive readout can be completely isolated from the high-voltage system.

There exist many different types of magnetic field sensors. The most widely used are fluxgate [3] sensors, magnetoresistive sensors [4], Hall-effect sensors [5], and pick-up coils. However, only the latter two are compatible with the standard CMOS process and are, hence, considerably cheaper to manufacture. This makes them very attractive.

Hall-effect sensors are based on the deflection of charge carriers under the influence of a magnetic field. Namely, when a current is applied to, for example, a slab of doped silicon (a Hall plate), the charge carriers will experience a Lorentz force. Opposite charges will accumulate at the slab's boundaries, an electric field will form, and a potential difference can be sensed. This potential difference is proportional to the magnetic field, the bias current, and some material- and geometry-related parameters.

There are two main downsides associated with such sensors. First of all, due to various imperfections in the device fabrication process [5], a Hall plate has a large offset voltage. To still obtain good DC accuracy, the spinning current technique [6] can be used. This technique effectively modulates the offset voltage to higher frequencies while keeping the signal of interest at DC. This, however, introduces significant tones at multiples of this so-called spinning frequency.

Secondly, as the Hall plate generally has a relatively large resistance and requires significant power for a decent sensitivity, it suffers from relatively high wide-band noise. As a result, a strict trade-off exists between bandwidth, power, and resolution. Combined with the aforementioned spinning tones, this makes Hall-effect sensors unattractive for high-bandwidth applications.

A pick-up coil's sensing mechanism, on the other hand, is based on Faraday's law of induction. This well-known law states that a *changing* magnetic flux gives rise to an electric field. As a result, a pick-up coil has a differentiating transfer characteristic: its sensitivity increases with frequency. Combined with a flat output noise spectrum (caused by the metal coil resistance), a large signal-to-noise ratio can be achieved at high frequencies. As opposed to Hall plates, pick-up coils are, therefore, superb for high-frequency sensing applications. However, they cannot be used to sense DC signals.

Especially in the control of power converters, with the rise of wide bandgap materials and the associated higher switching frequencies, there is an increasing demand for high-bandwidth, accurate current sensors [7, 8]. Neither Hall plates nor coils can provide this functionality by themselves. However, by combining these sensors (in a low- and high-frequency path), as conceptually shown in Figure 1.1, outstanding performance can be achieved. Previous work on this coil-Hall hybrid architecture [9–15] has indeed shown that a

simultaneous large bandwidth and high resolution can be obtained. A state-of-the-art implementation can reach integrated noise levels as low as 43mA_{RMS} while maintaining a bandwidth of 5MHz [13]. In a typical Hall-only solution, when trying to achieve a similar bandwidth, the resolution would generally be several hundreds of mA_{RMS} [16].

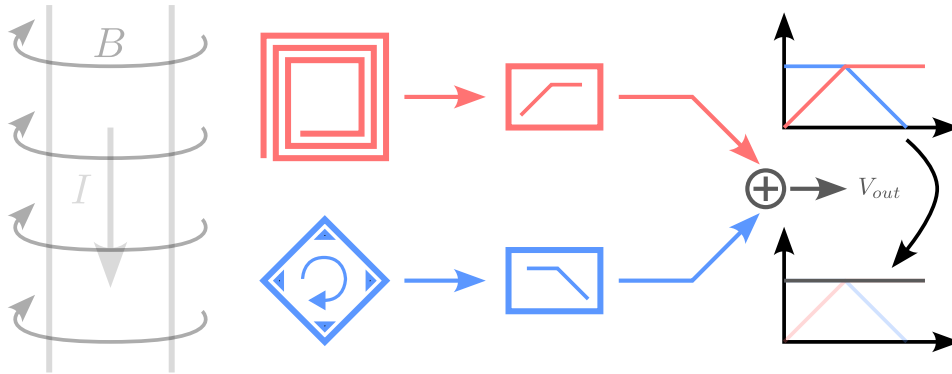


Figure 1.1: A conceptual representation of the coil-Hall hybrid architecture

1.1. THESIS GOAL

For SystematIC Design B.V., the company that has assigned this project, such a CMOS-compatible, high-resolution, and high-bandwidth current sensor is highly compelling. However, designing a hybrid sensor that seamlessly combines a Hall plate with a pick-up coil can be relatively complex. As will become evident, working with two completely different sensor types and accurately merging the two distinct signal paths introduces numerous challenges, demanding careful navigation of some tight trade-offs.

The goal of this thesis is twofold. First and foremost, it aims to thoroughly explore and comprehend the fundamental trade-offs involved. These are associated with combining the signal paths, as well as accurately calibrating and stabilizing the transducer gains — essential aspects for transforming this architectural idea into a viable product. This thesis offers a comprehensive overview of both existing and, over the course of this project, newly developed architectural solutions, highlighting their advantages and disadvantages. Since no single solution is perfect, the goal is to equip the reader with enough insight to determine their own approach to designing a coil-Hall hybrid current sensor. This thesis should, ideally, make the development of future iterations of this sensor type considerably easier.

Secondly, as tasked by SystematIC, this insight will be used to design a prototype chip. This chip is not meant to reach state-of-the-art performance. However, it is intended as a successor to SystematIC's previous Hall-Hall¹ design. This prototype is for them to show potential customers the significant performance increase associated with this new architecture. However, to save on design time and R&D costs, SystematIC highly values simplicity in the new design and, ideally, sees large parts of their previous chip design re-used. This simplicity constraint poses challenges but, at the same time, presents an excellent opportunity to explore and develop a new, less complex system solution that could lead to a more efficient way to achieve even better performance in the future.

1.2. THESIS SCOPE & OUTLINE

The two goals of this thesis complement each other well. Throughout the general discussion of the various challenges and solutions, the prototype chip will serve as a practical example to help the various trade-offs and design choices be understood even better. Most chapters are structured so that first, the general ideas are discussed, and second, a specific design choice for the prototype is made. Moreover, as SystematIC mainly wishes for a fully analog implementation² and to limit this thesis's scope, only analog implementations of the signal path will be considered. However, within this "analog subspace", a whole range of design considerations will be discussed.

¹Also a multi-path architecture like the coil-Hall architecture, but with only Hall plates in both paths

²Of course, one could work with an ADC and DAC to have an analog input and output, however, for SystematIC, this would mean an undesirable redesign of the entire system

First, in the next chapter, various design aspects relevant to the coil-Hall hybrid architecture will be introduced. Besides, the requirements imposed by SystematIC will be presented. In Chapter 3, some general considerations related to combining the two (low-frequency Hall-based and high-frequency coil-based) signal paths will be highlighted. In Chapter 4, various existing architectures will be discussed in detail, and the new "hybrid-hybrid" architecture will be introduced. In Chapter 5, especially challenges related to one-time high-frequency calibration of the coil-path will be discussed, while in Chapter 6, challenges associated with stabilizing this calibrated gain in both signal paths will be considered. In Chapter 7, the eventual prototype design for SystematIC will be presented, its performance will be verified, and some future recommendations will be given. And lastly, in Chapter 8, a conclusion will be drawn on whether the goals of this thesis have been successfully accomplished.

2

COIL-HALL HYBRID CURRENT SENSOR: IMPORTANT DESIGN ASPECTS

Before diving into the implementation details of the coil-Hall hybrid current sensor, it is essential to take a step back and consider what an ideal current sensing system would look like. And, more importantly, which phenomena could cause a deviation from this ideal behavior. Namely, by always keeping a clear overview of the core functionality and relevant performance aspects in mind, different design trade-offs can be more effectively navigated. To determine these design aspects, only three simple yet effective questions need to be answered:

1. What is the purpose/task of the system (core functionality)?
2. What error sources cause the system to deviate from this ideal behavior, and how much deviation is still acceptable (performance)?
3. What price must be paid to reach the required system performance (cost factors)?

Carefully evaluating these questions is believed to be an important step toward finishing any design to a satisfactory degree. Therefore, in this chapter, these questions will be answered for the coil-Hall hybrid current sensor architecture. Besides this chapter will conclude with a list of requirements for the to-be-designed prototype chip for SystematIC Design.

2.1. CORE FUNCTIONALITY

The first question aims at exposing the *ideal behavior* or *core functionality* of a system. For the current sensor, this can be captured in a single statement. Namely:

*The system should provide a **well-defined** and **stable** transfer
from the input current to a certain output quantity.*

The output quantity and system transfer could take on many forms. For example, one could choose an encoded digital bitstream or an up-modulated analog current. However, based on their customers' needs, SystematIC Design has expressed the desire for an analog output voltage proportional to the input current.

2.2. PERFORMANCE ASPECTS

Although the ultimate goal is to approach the aforementioned ideal behavior as closely as possible, fundamental, technological, and economic limitations [17] prevent us from doing so. The second question helps identify the different error sources that limit the system's performance. Furthermore, once identified, limits for each nonideality should be added to a list of requirements.

As will be shown in the next chapter, several performance aspects must be carefully considered for the coil-Hall hybrid current sensor. These aspects form the basis of some strict design trade-offs when attempting to determine the best conceptual architecture. Including some brief remarks, these aspects are as follows:

2.2.1. NOISE

Different noise sources limit the system's resolution. The performance will be expressed as an input referred root-mean-square noise current (I_{RMS}). Especially with the coil-Hall hybrid sensor, in which the noise is filtered at a relatively low frequency, thermal- and flicker noise should both be considered.

2.2.2. INPUT RANGE & DYNAMIC RANGE

The input range sets the upper bound for the maximum input signal. This is mainly limited by signal clipping to the supply or other (more soft) nonlinear effects. The ratio of the maximum input signal to the resolution (noise level) is the system's so-called dynamic range; maximizing this quantity as much as possible is desirable.

2.2.3. BANDWIDTH

If left uncompensated, any resistance in combination with a (parasitic) capacitance or inductance will limit the system bandwidth. The bandwidth is specified as the frequency at which the gain has dropped by 3dB. For example, the inherent coil bandwidth decreases when increasing its sensitivity: the larger the number of turns, the greater the coil's resistance and capacitance.

2.2.4. GAIN STABILITY

Ideally, the gain remains as constant/stable as possible over different operating conditions. However, many unwanted external factors¹ generally affect transducer and circuit properties. For SystematIC, gain variation over temperature and lifetime is regarded as the most important and will be mainly focused on. The gain variation will be expressed as a relative variation (percentage) w.r.t. the ideal gain.

2.2.5. GAIN FLATNESS OVER FREQUENCY

At the core of the coil-Hall architecture is the combining of two completely different signal paths. In light of providing a well-behaved and accurate sensing system, ideally, no artifacts of this path-combining should be visible in the overall frequency response. However, as will be shown later, especially around the cross-over frequency (i.e., the frequency at which one path takes over from the other), some relatively large gain variations could occur. These variations should be minimized. Like gain variation over temperature and lifetime, gain flatness over frequency will be expressed as a (maximum) percentual variation w.r.t. to the ideal gain.

2.2.6. OFFSET

Offset limits the system's DC accuracy. In a Hall-based current sensor, a significant amount of offset comes from contact misalignment and Hall plate impurities. However, since the invention of the spinning current technique [6], this offset contribution has been greatly reduced. Consequently, the offset added by the read-out circuitry has become even more critical. In the coil-Hall architecture, careful design must ensure that both signal paths do not add significant offset. The total offset is either expressed as an input referred current or an output voltage (3σ value).

2.2.7. SWITCHING RIPPLE

Lastly, an unwanted ripple signal could propagate to the output if switching is performed in the signal path. As switching is not always required, not all systems have to deal with such ripples. However, for any Hall-based current sensor nowadays, the Hall plates are spun to reach the required offset performance. This spinning ripple (up-modulated offset and low-frequency noise) should generally be well below the integrated noise level (as defined in Section 2.2.1) and is expressed as an input referred (RMS) current.

¹Think, for example, of temperature, mechanical stress, humidity, light, etc.

2.3. COST FACTORS

Effective design is not simply about maximizing performance. Effective design is about maximizing performance *while minimizing the cost*. Therefore, identifying a design's relevant performance aspects is only one side of the coin. Of equal importance is having a clear overview of the cost factors involved and a specified limit for each of them.

Cost factors exist in many different forms. To name a few, one could think of costs during system operation, such as power consumption, required supply voltage or contamination of the electrical environment, costs related to the price of a product, such as manufacturing cost and design labor, or costs related to the product's system integration, such as dimensions and mass. Although they are all important to consider, their relevance differs from project to project. For the design of the coil-Hall hybrid sensor, it turns out that this list can be reduced to only four cost factors/categories. These are at the heart of some important design trade-offs.

2.3.1. POWER CONSUMPTION

Power consumption is one of the most important cost factors to consider in virtually every electrical system. Whether people aim for higher system efficiency or increased battery life, they are certainly right to consider it. However, especially in microelectronics, sometimes too much focus goes towards reduced power consumption while overlooking the system's higher-level application. In this case, the coil-Hall current sensor will most likely be integrated into relatively high-power systems. For example, the to-be-designed prototype will be rated for currents up to 80A. Only considering the current rail resistance of roughly $1\text{ m}\Omega^2$, such a large current will already burn multiple watts. As a result, an increase in power consumption of several tens or even hundreds of milliwatts in the current sensor will most likely not make a significant difference to the total system efficiency.

That being said, this design will still aim for low power consumption. As power consumption is fundamentally related to important performance aspects such as noise and bandwidth, potential customers and academia often use it as part of a FOM (figure of merit) to compare designs. Therefore, low power consumption is still seen as an advantage when competing with different designs/products. Yet, in this design it will always be considered in the context of the potential (high-power) application, and therefore, lower power consumption will not be blindly strived for at the cost of, for example, a significant performance penalty.

2.3.2. CHIP AREA

Generally, decreased chip area is desirable. With a smaller chip area, more chips can be manufactured on a die, reducing the cost per chip. However, for the prototype chip for SystematIC Design, chip area is believed to be even more critical. Namely, as SystematIC's previous Hall-Hall design will be reused as much as possible, only limited space is available for the coil and associated readout. This makes chip area an even more limited resource.

2.3.3. MANUFACTURING COST: CALIBRATION TIME

Almost any company's ultimate goal is to sell their product with a profitable margin. Minimizing each product's manufacturing cost directly helps towards achieving this goal. Reducing the chip area is one way to accomplish this. However, other aspects, such as minimizing the number of process steps³, using a relatively low-cost technology and limiting each chip's required calibration time, also contribute towards this goal. SystematIC Design has already fixed everything technology-related during their previous design; for the current design, this ideally should not be changed. However, as will be shown later, the required calibration time is still highly design-dependent and is important to keep in the back of one's mind.

2.3.4. DESIGN COMPLEXITY

Lastly, one could consider the design's complexity as a cost factor. Although high design complexity could potentially improve performance at lower operational costs (e.g., power consumption), it generally also introduces more risks. As a result, more time (and thus money) must be spent on proper chip design and validation. Besides, there is a higher chance of overlooking critical design aspects and the concept not working anyway. Therefore, a more robust and simple design is generally desirable.

²This is roughly the minimum resistance of the current-rail used in SystematIC's chips

³SystematIC only uses, for example, 4 metal layers for routing and MIM (metal-insulator-metal) capacitors

2.4. SYSTEM REQUIREMENTS

Having a clear overview of the core functionality, relevant performance aspects, and cost factors related to the coil-Hall hybrid architecture, the following four chapters will introduce multiple system architectures for both **combining** (Chapters 3 and 4), **calibrating** (Chapter 5), and **stabilizing** (Chapter 6) the different signal paths. Ultimately, the concept that aligns best with SystematIC's needs will be further implemented.

SystematIC wishes this new coil-Hall architecture to be integrated into an existing Hall-Hall design to greatly enhance its noise and bandwidth performance. Ideally, this should be done with as few changes as possible; for example, the low-frequency signal path containing the spinning Hall plates (as will be shown later) should be re-used if possible. Besides, as mentioned in the introduction, to reduce costs, SystematIC highly values simplicity in the new design. Still, the noise should decrease from 180 mA_{RMS} to less than 80 mA_{RMS} , while the bandwidth should increase from 318 kHz to 2 MHz : an improvement of more than 5x in noise power and 6x in bandwidth compared to the previous design.

At the same time, the switching/spinning ripple should be further reduced to stay well below the new integrated noise level. According to measurements of this ripple from previous chips, the low-pass filtering at the output of the low-frequency path (see Chapter 4) should provide more than 56 dB of attenuation at 200 kHz (instead of 45 dB in the previous designs). All other specifications remain the same compared to the previous design and are summarized in the following Tables⁴. A test plan on how various specifications should later be verified can be found in Appendix A.

Requirement	Value
System Gain ⁵	25 mV/A

Table 2.1: Functional Requirements

Requirement	Value
Temperature Range	-40°C to $+125^\circ\text{C}$
Supply Voltage	3.3 V - 5.5 V (5 V typical)

Table 2.2: Operating Conditions

Requirement	Value
Noise (at 25°C)	80 mA_{RMS} (typical)
Input Range ⁶	$\pm 80 \text{ A}$
Bandwidth	2 MHz
Gain Stability ⁷ (temperature)	$\pm 3\%$ (ideally $\pm 1\%$)
Gain Stability (lifetime)	$\pm 3\%$ (ideally $\pm 1\%$)
Gain Flatness over Frequency	$\pm 1\%$
Offset	$\pm 70 \text{ mA}$ (3σ)
Switching Ripple	$\ll 80 \text{ mA}_{RMS}$

Table 2.3: Performance Specifications

Requirement	Value
Supply Current	20 mA
Supply Current (HFP ⁸)	5 mA
Chip Area	4 mm^2
Chip Area (HFP)	1 mm^2
Chip Technology	180 nm , 5 V

Table 2.4: Cost Specification

⁴It should be noted that only requirements relevant for this thesis are listed

⁵SystematIC actually wants multiple gain settings ranging from 25 mV/A to 100 mV/A in some discrete steps (25, 40, 50, 66.6, 80, 100)

⁶SystematIC also wants multiple input ranges (corresponding to the different gains): $\pm 20 \text{ A}$, $\pm 30 \text{ A}$, $\pm 50 \text{ A}$, $\pm 80 \text{ A}$, $0-40 \text{ A}$, $0-50 \text{ A}$, $0-80 \text{ A}$

⁷For this prototype chip, $\pm 3\%$ is sufficient, but ideally SystematIC would like to bring this down to $\pm 1\%$ for future designs (the same holds for lifetime drift)

⁸High-Frequency Path, as will become evident, this thesis will mainly be concerned with the design of this path

3

SIGNAL PATH COMBINING: GENERAL CONSIDERATIONS

One of the main challenges with the coil-Hall hybrid architecture is accurately combining two signal paths based on two completely different transducer types. This chapter is meant to introduce some general considerations associated with this signal path combining. To keep things general for now, no implementation is considered yet. Therefore, all equations and intuition introduced in this chapter effectively hold for any implementation using a coil and Hall plate in the two different signal paths. This knowledge will be used in the next chapter, in which multiple implementations will be presented.

This chapter commences by providing a brief explanation of the two used transducers: the Hall plate and (pick-up) coil. Next, the goal of each (Hall- and coil-based) signal path will be investigated, the ideal transfer functions for properly combining these paths will be considered, and some noise calculations will be performed. Section 3.5 will discuss the effect of unwanted poles in the low- or high-frequency path, while, lastly, section 3.6 will discuss which quantity (voltage or current) can be best used when reading out the coil.

3.1. TRANSDUCER MODELS

Before exploring the coil-Hall architecture further, it is helpful to have a general understanding of the different transducers and have an adequate electrical model to work with. In this section, these models will be provided for both the Hall plate and a coil. This section merely provides relevant design information, not an in-depth explanation of all physical phenomena.

3.1.1. HALL PLATE MODEL

A Hall plate is a four-terminal device that, if properly biased, can be used to sense magnetic fields. Simply put, a Hall plate works by the Lorentz force acting on moving charge carriers. Namely, under the influence of a magnetic field, the Lorentz force causes the moving charge carriers to be deflected to the "walls" of the plate. To reach an equilibrium state, the separation of charges creates a counteracting electrical field and, thus, a voltage to be sensed. This voltage is proportional to the magnetic field, the bias current (or voltage), and various material- and geometry-related parameters [5].

A Hall plate is often modeled as a Wheatstone-bridge-like structure as shown in Figure 3.1 [5]. The model consists of the two terminals to which the bias current (or voltage) is applied and two terminals at which the Hall voltage (or current) is read out. As a Hall plate is generally implemented in a symmetrical slab of (n-doped) silicon, each terminal is interconnected via an equal resistance R_H . This simple model can be expanded by adding capacitances to the substrate; however, this capacitance is commonly very small, and the dynamic behavior is normally dominated by the capacitance associated with the readout [18] and the Hall plate resistance. Therefore, it is chosen that the capacitances are omitted in this model.

Normally, the Hall plates are characterized using the resistance R_H , the voltage-related sensitivity S_V , and the current-related sensitivity S_I . The latter two indicate the obtainable sensitivity for a certain bias voltage or current, respectively ($S_H = S_V V_B = S_I I_B$). As obtained from measurement data, the Hall plates used by

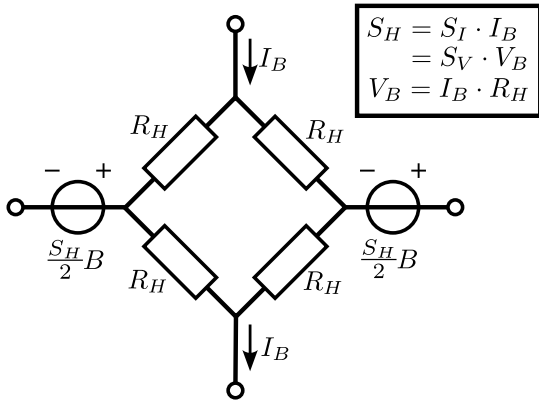


Figure 3.1: Electrical model of the Hall plate

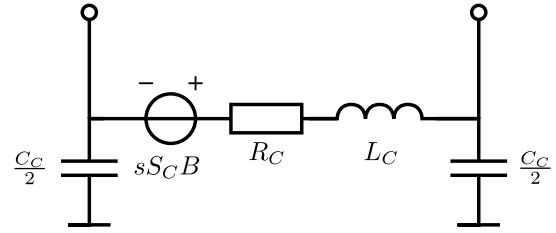


Figure 3.2: Electrical model of the coil

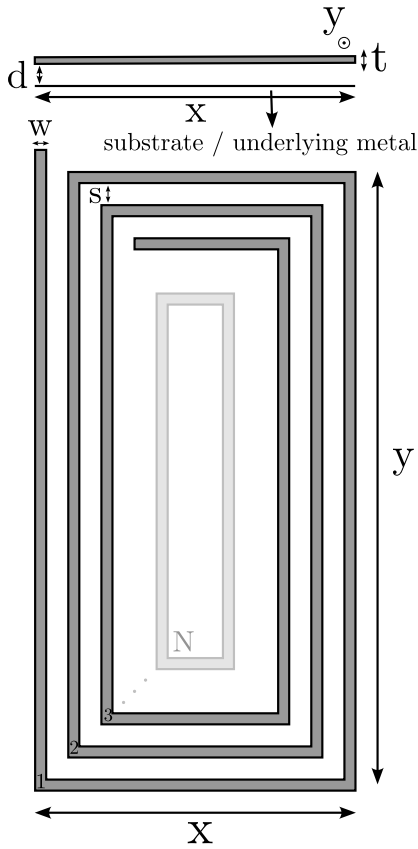


Figure 3.3: Coil representation with geometry-related parameter definitions for calculation purposes

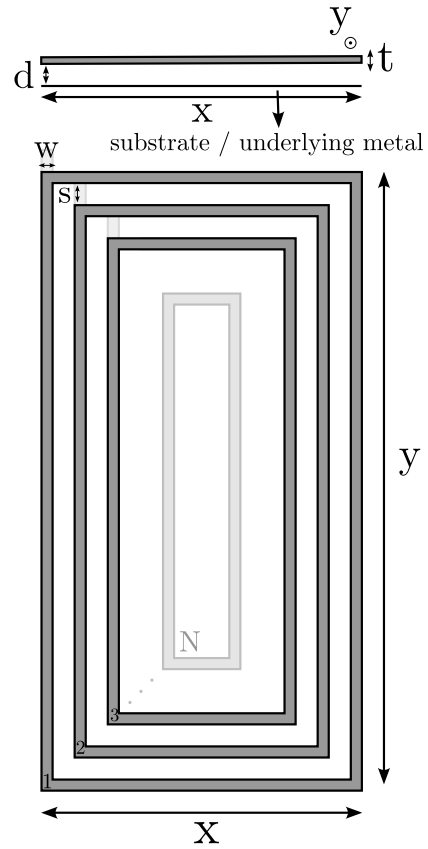


Figure 3.4: Coil of Figure 3.3 with each turn seen as a separate rectangle used for model parameter calculation purposes

SystematIC nominally have a voltage-related sensitivity S_V of $52 \text{ mV V}^{-1} \text{ T}^{-1}$, a current-related sensitivity S_I of $255 \text{ V A}^{-1} \text{ T}^{-1}$, and a resistance R_H of $5 \text{ k}\Omega$ at room temperature¹. This resistance could effectively be lowered by placing multiple Hall plates in parallel (to obtain, for example, a lower noise level). However, for a given sensitivity, this comes at the cost of an increase in bias current.

3.1.2. COIL MODEL

Next, the coil should be considered. The coil's magnetic field sensing mechanism is based on Faraday's law of induction. As shown by Equation 3.1, this law states that a changing magnetic flux gives rise to an electromotive force. Which, at the coil terminals, can be sensed as a voltage or, transformed over the coil impedance, as

¹These three parameters are highly process- and temperature-dependent

a current.

$$\mathcal{E} = -\frac{d\phi_B}{dt} \quad (3.1)$$

Assuming the magnetic field distribution at the planar coil is roughly constant, the magnetic flux captured by the coil can be written as $\phi_B = B_{\perp} \cdot A_{tot}$ in which $B_{\perp}$ (or throughout this thesis often written as B) is the magnetic field component perpendicular to the coil and A_{tot} the total area enclosed by all the coil turns. In the Laplace domain, neglecting the sign, the transfer function can be written as:

$$V_C = sS_C B \quad (3.2)$$

with

$$S_C = A_{tot} \quad (3.3)$$

After obtaining a general understanding of the coil's transduction mechanism, the next step is to devise an electrical model of the coil. This model will be used for the design and validation of the noise and dynamic behavior and, therefore, should contain both resistive and reactive elements. Although a planar coil is a complex distributed network, the 2-terminal lumped model shown in Figure 3.2 will be used. This model contains the (differentiating) coil sensitivity (S_C), a series resistance (R_C), the coil's self-inductance (L_C), and a capacitance to the ground (either the substrate or another metal layer) equally distributed on each side (C_C).

No fabricated coils were available to verify this model. However, several sources from the literature have reported similar models that correspond relatively accurately to measurement data [19–21]. Furthermore, simulations have given the impression that the simple model shown in Figure 3.2 gives a lower bandwidth estimate than a distributed model (made from multiple sections of this simple model). Therefore, as long as the bandwidth specifications are met with this model, the actual coil will likely also meet the bandwidth requirements. For estimating the noise performance, due to the low-frequency filtering, only the sensitivity and resistance are required.

To determine the different component values for this model, the rectangular planar coil shown in Figure 3.3 can be divided into a set of closed rectangles as shown in Figure 3.4. The different model parameters can be related to the total coil length (L_{tot}) and the total coil area (A_{tot}). With the coil winding inwards, the length and enclosed area of each rectangle decrease slightly for each additional turn. Therefore, the total length is determined by calculating the perimeter of each rectangle and then summing these for all N rectangles, while the total area is determined by first calculating the area enclosed by each rectangle and then summing these values for all N rectangles. As long as the trace width (w) is not too large compared to the outer dimensions (x and y), this approximation is believed to be accurate.

From these parameters, R_C can be approximated using the metal sheet resistance and the total number of squares ($\frac{L_{tot}}{w}$), C_C can be estimated as a parallel plate capacitance to ground (either the substrate or an overlapping metal layer), and the inductance L_C can be approximated using a modified version of Wheeler's formula [22]. The derivations, formulas, and several plots can be found in Appendix B. In Appendix C more details can be found on how to properly dimension the coil for optimal performance.

3.2. WHY A MULTIPATH ARCHITECTURE?

Having obtained a general understanding of the various transduction mechanisms, it is important to consider why multiple signal paths, making use of these different transducers, are used in the first place. With this understanding, the following design steps can be navigated more effectively. An excellent way to obtain such insight is by looking at the history of Hall-based sensors and, especially, at the challenges that arose. These challenges form the basis for the multipath architecture and can, therefore, be used to determine the particular functionality required from this architecture.

In early Hall effect-based designs, only a single signal path was used. While Hall plates could be perfectly used in a relatively wide-band system, a problem arose when measuring DC signals; the inherent Hall plate

offset turned out to be extremely large. To combat this problem, the spinning current technique [6] was invented, which, still to this day, is the best way to reach low offset performance. However, this technique has its downside. Namely, as this modulation technique mixes the offset to a higher frequency, it reduces the system's accuracy at higher frequencies.

Normally, the resulting spinning ripple is filtered by a low-pass filter. However, as a higher spinning frequency generally deteriorates the offset performance [18, 23, 24], the filter requires a relatively low corner frequency and thus greatly reduces the system bandwidth (to maximally several hundreds of kilohertz). Later, so-called "ripple reduction loops" [24] were employed instead of a low-pass filter. Instead of filtering all frequencies above a certain range, this innovative solution uses several feedback loops to remove only tones at multiples of the spinning frequency, effectively implementing a notch filter. Although this solution allows for accurate measurement of high-frequency signals, the notches still create "blind spots" inside the frequency band of interest.

The initial idea behind adding an additional (high-frequency) signal path was to "fill in" these notches² and, thus, flatten the frequency response. This allows for a large bandwidth without sacrificing performance due to the large ripple components: it effectively breaks the aforementioned trade-off between offset and bandwidth performance.

First, this was done using a high-frequency path based on a non-spinning Hall plate [9]. However, due to the relatively large Hall plate resistance and power required to obtain a decent sensitivity, another trade-off between resolution, bandwidth, and power consumption started to form. As with any system with wide-band noise, this trade-off caused the resolution to decrease quickly for an increase in bandwidth.

Luckily, using multiple signal paths again offered a way out of this tight trade-off. Namely, instead of a Hall plate, a coil could be used in the high-frequency path [9, 10]. As a coil does not require any power and effectively causes the signal-to-noise ratio³ (SNR) to increase with frequency, it is the perfect substitute for the high-frequency Hall plate to overcome this noise-bandwidth trade-off.

A multipath architecture is therefore used for two distinct functions:

1. Remove spinning ripple while allowing a low spinning frequency for good offset performance (coil-Hall or Hall-Hall)
2. Decrease overall noise while still maintaining a wide bandwidth (coil-Hall)

3.3. IDEAL TRANSFER FUNCTION

The previous section determined *what* functionality the multipath architecture should provide. The next step is to determine *how* such functionality can be obtained. This section will consider this by looking at the required transfer functions in each path to obtain the desired system response. By deriving these equations for a general system without considering any implementation, these equations can be used to explore the general behavior of the coil-Hall hybrid sensor and find (in this and the next chapter) some fundamental limits concerning various design aspects.

For this purpose, Figure 3.5 shows a high-level conceptual overview of the full system. It shows the transfer of the input current I on the current rail to the output voltage V_{out} via the magnetic field domain and the two signal paths. One path contains the spinning Hall plate for measuring DC and low-frequency signals, while the other path contains the (differentiating) coil to measure high-frequency signals with a large SNR. These two paths are filtered using the transfer functions $H_{LFP}(s)$ and $H_{HFP}(s)$ respectively, and are finally summed to form a single output signal.

While complex internally, externally, the system should look like a gain A_{sys} that accurately relates the input current (I) to the output voltage (V_{out}) over a wide range of frequencies. This behavior can be captured as shown in Equation 3.4:

$$V_{out} = (sS_C A_{Co,C} \cdot H_{HFP}(s) + S_H A_{Co,H} \cdot H_{LFP}(s)) \cdot I = A_{sys} \cdot I \quad (3.4)$$

²Or, if a low-pass filter is used, extend the bandwidth while filtering the ripple

³Due to its increasing sensitivity with frequency and flat noise spectrum

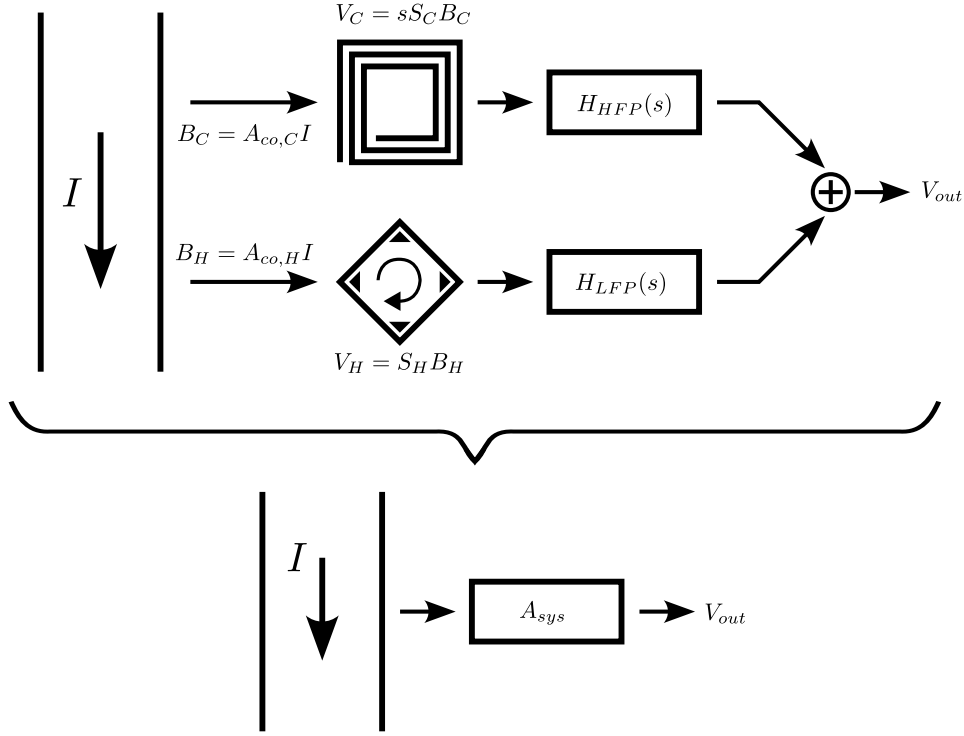


Figure 3.5: General system overview for combining the two signal paths with an arbitrary transfer function, from a current rail input to an output voltage

In this equation, S_C is the coil sensitivity (in $\text{VT}^{-1} \text{rad}^{-1} \text{s}$), $A_{co,C}$ and $A_{co,H}$ the coupling factors from the input current to the magnetic field at the coil and Hall plate respectively (in TA^{-1}) and S_H the sensitivity of the Hall plate (in VT^{-1}). This equation, in turn, can be solved to yield the following result:

$$H_{HFP}(s) = \frac{1}{s} \cdot \frac{1}{S_C A_{co,C}} \cdot (A_{sys} - S_H A_{co,H} \cdot H_{LFP}(s)) \quad (3.5)$$

This equation describes the high-frequency path transfer function, $H_{HFP}(s)$, required to get a flat response for a given low-frequency path transfer function, $H_{LFP}(s)$. Theoretically, an infinite number of transfer functions $H_{LFP}(s)$ can be chosen. As long as the transfer $H_{HFP}(s)$ obeys the relation presented in Equation 3.5, the total response should be flat.

Needless to say, not all transfer functions need to be considered. Namely, as described in the previous section, the low-frequency path transfer must provide filtering to reduce the spinning ripple and noise coming from the low-frequency Hall plate and readout. Therefore, only a low-pass characteristic for $H_{LFP}(s)$ is meaningful to consider.

In this section, the first-order case will be investigated. This is educational as the performance of more complex cases (see, for example, the second-order case in Section 4.2.1) can be related to this first-order case. Besides, the first-order case is the least complex to implement accurately and has already been used in several designs. For this purpose, $H_{LFP}(s)$ can be written as:

$$H_{LFP}(s) = \frac{A_H}{1 + \frac{s}{p_x}} \quad (3.6)$$

The transfer has a DC gain, A_H , and a single pole at $s = -p_x$. Furthermore, to obtain an insightful expression for $H_{HFP}(s)$, equal magnetic field coupling factors are assumed⁴ ($A_{co,C} = A_{co,H} = A_{co}$) and the system gain is set equal to the DC gain in the low-frequency path ($A_{sys} = A_{co} S_H A_H$). With these assumptions, the required high-frequency path transfer can be written as:

⁴A relatively fair assumption if the Hall plates and coils are of similar dimensions and located close to each other

$$H_{HFP}(s) = \frac{S_H A_H}{S_C} \cdot \frac{1}{s} \cdot \frac{\frac{s}{p_x}}{1 + \frac{s}{p_x}} = \frac{S_H A_H}{S_C p_x} \cdot \frac{1}{1 + \frac{s}{p_x}} = \frac{A_C}{1 + \frac{s}{p_x}} \quad (3.7)$$

This equation accurately depicts the advantage of using a coil instead of a Hall plate in the high-frequency path. Namely, the integrating characteristic required to flatten the coil response causes the high-frequency path transfer to be bandwidth-limited. Therefore, just like in the low-frequency path, all noise and interference in the high-frequency path see a low-pass filter with a (relatively low-frequency) pole at $s = -p_x$. This would not be the case when using a Hall plate in the high-frequency path.

This equation also shows that $H_{HFP}(s)$ has a different DC gain compared to $H_{LFP}(s)$. Namely, the DC gain is $\frac{S_H}{S_C p_x}$ times larger than the low-frequency path DC gain. If, for example, S_H is larger than $S_C p_x$, any noise, interference, or offset at the coil output sees a larger gain than the Hall plate output and, as a result, becomes more dominant. This term ($\frac{S_H}{S_C p_x} = \frac{S_H}{2\pi f_x S_C}$) containing the transducers sensitivities (S_H and S_C) and the cross-over frequency⁵ ($f_x = \frac{p_x}{2\pi}$) effectively describes how well the two paths are balanced.

1. If f_x is chosen too small, the high-frequency coil only has a small gain when taking over the low-frequency path. This "weak" high-frequency path could, therefore, insert a lot of noise, offset, and interference.
2. On the other hand, if f_x is chosen too large, the low-pass filter is less effective, and the coil benefits are not optimally exploited by allowing more high-frequency noise and interference from the low-frequency path to propagate to the output.

This balancing act will especially become more evident when calculating the different noise contributions from each path in the next section.

3.4. NOISE CALCULATIONS

To obtain some insightful expressions for the noise performance, it is assumed that each path's (white⁶) noise contribution can be written as an equivalent noise resistance: $R_{H,n}$ and $R_{C,n}$ ⁷ (in Ω). These resistances capture the equivalent noise resistance of both the transducers (R_H and R_C) and the subsequent amplifiers at the input of the low- and high-frequency paths, respectively. Then, under the same assumptions used to derive Equation 3.7 and using the equivalent noise bandwidth of a first-order filter, the total integrated input-referred noise current (I_N in A_{RMS}) can be written as:

$$I_N = \sqrt{4kT \cdot (R_{H,n} + (\frac{A_C}{A_H})^2 \cdot R_{C,n}) \cdot \frac{1}{A_{co}^2} \cdot \frac{1}{S_H^2} \cdot \frac{\pi}{2} f_x} \quad (3.8)$$

In which T is the absolute temperature (in K) and k is the Boltzmann constant (approximately $1.381 \cdot 10^{-23} \text{ m}^2 \text{ kg s}^{-2} \text{ K}^{-1}$). After some simplification, this equation reduces to:

$$I_N = \sqrt{kT \cdot \frac{1}{A_{co}^2} \cdot (\frac{R_{H,n}}{S_H^2} \cdot 2\pi f_x + \frac{R_{C,n}}{S_C^2} \cdot \frac{1}{2\pi f_x})} \quad (3.9)$$

It can clearly be seen that this "balancing act" described in the previous section also applies to noise performance. A large cross-over frequency causes the noise of the low-frequency Hall path to become dominant, while a small cross-over frequency causes the noise of the high-frequency coil path to dominate. Optimum noise performance occurs when the two noise contributions are equal. This happens at a specific cross-over frequency $f_{x,opt} = \frac{p_{x,opt}}{2\pi}$ ⁸:

⁵i.e., the frequency f_x associated with the pole at $s = -p_x$, at which the high-frequency coil path takes over from the low-frequency Hall path

⁶To keep the equations more fundamental and insightful 1/f-noise is not considered in this derivation; during implementation, however, one should certainly consider its effects

⁷ R_H and R_C will be reserved to indicate the physical resistance of the Hall plate and coil respectively

⁸Indeed, this cross-over frequency shows that, for example, a "strong" low-frequency path (with a large S_H and small $R_{H,n}$) causes the optimum cross-over to occur at a higher frequency. The inverse happens for a "strong" high-frequency path in which S_C is large, and $R_{C,n}$ is small.

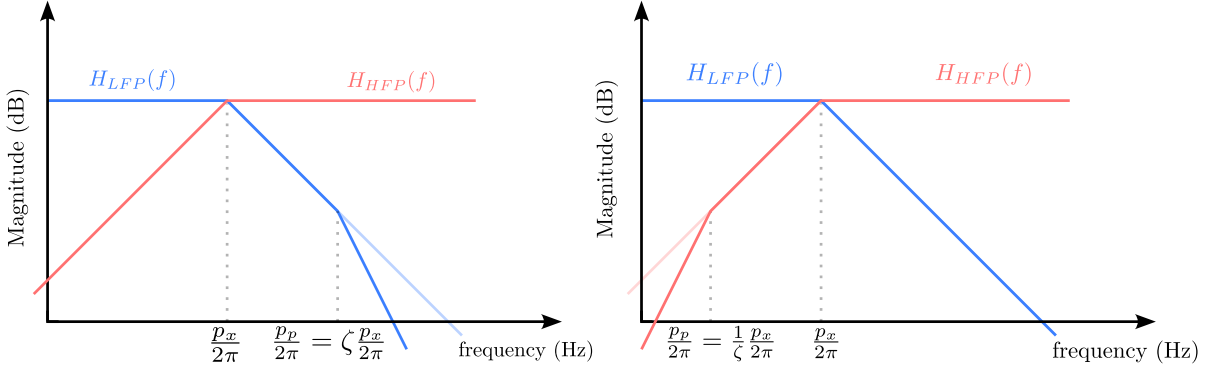


Figure 3.6: The two first-order combining transfer functions (excluding the differentiating coil characteristic's effect) with the additional pole defined in each path as a function of the cross-over frequency.

$$f_{x,opt} = \frac{1}{2\pi} \frac{S_H}{S_C} \sqrt{\frac{R_{C,n}}{R_{H,n}}} \quad (3.10)$$

Resulting in the total noise:

$$I_{N,opt} = \sqrt{2kT \cdot \frac{1}{A_{co}^2} \cdot \frac{\sqrt{R_{H,n}R_{C,n}}}{S_H S_C}} \quad (3.11)$$

As one would expect, better noise performance occurs for smaller resistances ($R_{C,n}$ and $R_{H,n}$), larger transducer sensitivities (S_C and S_H), lower temperatures (T), and a larger current to magnetic field coupling factor (A_{co}). Generally, the designer only has a significant influence on the transducer and readout related factors $\frac{\sqrt{R_{H,n}}}{S_H}$ and $\frac{\sqrt{R_{C,n}}}{S_C}$ in the low- and high-frequency paths, respectively. To minimize these terms, low-noise amplifiers, small transducer resistances, and large sensitivities are required.

Lastly, it is important to realize that f_x cannot always be chosen at the optimum cross-over frequency. Gain flatness, ripple reduction, and area constraints are some of the potential factors that could force a deviation from the optimum frequency. When deviating from the optimum cross-over frequency by a factor γ (i.e., $f_x = \gamma f_{x,opt}$). It can be proven that the noise penalty can be written as $\sqrt{\frac{1}{2} \cdot (\gamma + \frac{1}{\gamma})}$; the input-referred noise changes to:

$$I_N = \sqrt{(\gamma + \frac{1}{\gamma}) \cdot kT \cdot \frac{1}{A_{co}^2} \cdot \frac{\sqrt{R_{H,n}R_{C,n}}}{S_H S_C}} \quad (3.12)$$

This equation shows that the penalty is small for small deviations from the optimum. However, in the limit, each doubling or halving of cross-over frequency effectively causes a $\sqrt{2}$ (3 dB) noise penalty.

3.5. EFFECT OF AN UNWANTED POLE

As mentioned in Chapter 2, an important design aspect is gain flatness over frequency. As will become evident in the next chapter, several mechanisms could cause deviations from an ideal flat response. However, in most cases, the gain flatness is affected by an additional (unwanted) pole in either the low- or high-frequency path. The associated error will be investigated in this section; this result will be used in the next chapter.

For this purpose, the total normalized response after combining the two paths with either an additional pole in the low- or high-frequency path can be written respectively as:

$$H_{tot,LF}(s) = \frac{1}{1 + \frac{s}{p_x}} \cdot \frac{1}{1 + \frac{s}{p_p}} + \frac{\frac{s}{p_x}}{1 + \frac{s}{p_x}} = \frac{1 + \frac{s}{p_x} + \frac{s^2}{p_x p_p}}{1 + s(\frac{1}{p_x} + \frac{1}{p_p}) + \frac{s^2}{p_x p_p}} \quad (3.13)$$

$$H_{tot,HF}(s) = \frac{1}{1 + \frac{s}{p_x}} + \frac{\frac{s}{p_x}}{1 + \frac{s}{p_x}} \cdot \frac{\frac{s}{p_p}}{1 + \frac{s}{p_p}} = \frac{1 + \frac{s}{p_p} + \frac{s^2}{p_x p_p}}{1 + s(\frac{1}{p_x} + \frac{1}{p_p}) + \frac{s^2}{p_x p_p}} \quad (3.14)$$

If the location of the additional pole is defined as indicated in Figure 3.6 ($p_p = \zeta p_x$ for a non-ideal low-frequency path or $p_p = \frac{1}{\zeta} p_x$ for the non-ideal high-frequency path with $\zeta > 1$), it can be seen (from the similarity between both transfer functions) that both transfer functions will effectively yield the same result: a dip between p_x and p_p in the combined frequency response. With the given definitions, the maximum gain drop as a function of ζ is equal in both cases. Therefore, only one case has to be further elaborated on.

For this purpose, the squared magnitude of $H_{tot,LF}(j\omega)$ can be considered:

$$|H_{tot,LF}(j\omega)|^2 = \frac{(1 - \frac{\omega^2}{p_p p_x}) + \frac{\omega^2}{p_x^2}}{(1 - \frac{\omega^2}{p_p p_x}) + \omega^2 (\frac{p_x + p_p}{p_x p_p})^2} \quad (3.15)$$

And, after substituting $p_p = \zeta p_x$ and noticing that the largest dip occurs at the geometric mean of the two pole frequencies $\omega = \sqrt{p_x p_p}$, the squared magnitude can be written as:

$$|H_{tot,LF}(j\omega)|^2 = (\frac{\zeta}{1 + \zeta})^2 \quad (3.16)$$

Finally, the error ϵ with respect to the ideal gain of 1x (0 dB) can be shown to be:

$$\epsilon = 1 - |H_{tot}(j\omega)| = \frac{1}{1 + \zeta} \quad (3.17)$$

As expected, this expression indicates that the further away the pole is from the cross-over frequency, the smaller the gain flatness error becomes. In the limit, it scales inversely proportional to the relative pole distance. For example, for an unwanted pole, $s = -p_p$ two decades (i.e., $\zeta = 100$) above the cross-over frequency in the low-frequency path or two decades below the cross-over frequency in the high-frequency path, an error of maximally 1% in the gain flatness. This result has also been verified by simulations.

3.6. COIL VOLTAGE VERSUS CURRENT READOUT

As with many sensors, either the current or voltage can be chosen as the information-carrying quantity. As one of the two quantities commonly conveys the information more accurately, this is usually a relatively straightforward decision. However, reading out either quantity comes with its own advantages and disadvantages when considering a coil. These will be briefly investigated in this section⁹. In Appendix C, more detailed explanations and a strategy for dimensioning the coil can be found.

3.6.1. READING OUT THE COIL VOLTAGE

First, one could consider reading out the coil voltage. The main advantage of this approach is its stability over temperature and lifetime in combination with its expected low variation over process corners. Namely, the magnetic field to voltage transfer of the coil is purely related to the coil geometry, which, in turn, can be controlled very accurately by lithography. Besides, this geometry should only change marginally over temperature and lifetime. Even though the differentiating coil transfer should be compensated with an integrating transfer (which generally results in the product of resistance and capacitance in the gain), temperature stability can be maintained by using components with a low-temperature coefficient. These are generally available in most IC processes.

While this temperature stability is a nice feature to have and omits the need for additional temperature compensation, reading out the coil voltage, unfortunately, also has a major disadvantage. Namely, due to the differentiating coil characteristic, the voltage at the coil's output can become relatively large for high-speed

⁹The Hall plate will not be considered as the low-frequency signal path of SystematIC's previous Hall-Hall design will be re-used. In that design, the Hall voltage is read out.

signals (e.g., several volts for a moderate sensitivity and a full-scale MHz signal). To prevent subsequent transistor stages from breaking down, this signal must be limited by, for example, a clamping circuit¹⁰.

Excessive distortion occurs when the coil signal is clamped. As a result, the maximum signal the system can handle at a certain frequency is limited; in other words, the full-power bandwidth is limited. This is likely not an issue in applications where the range of signal amplitudes and frequencies are known (e.g., think of a motor control system). However, for unknown signals, this sudden clipping could be a problem. In case of a short-circuit event, for example, the clipping effectively causes a part of the signal energy to be lost, resulting in the high-frequency path only partially reacting (simulation results are presented in Section 7.6.5). As a result, the response time is determined by the time constants related to the cross-over frequency f_x and not related to the full system bandwidth (e.g., 80 μ s versus 80ns), while, particularly for a short-circuit event, a fast response is required. To reduce the chances of this happening, the coil sensitivity should be limited, degrading the system's noise performance. Some more detailed information and some numerical examples can be found in Appendix C.1.

3.6.2. READING OUT THE COIL CURRENT

On the other hand, when reading out the coil current [12], one has to deal with large currents instead of voltages. Although clipping is still possible, the limit is now determined by the amplifier's current drive capability¹¹, not by the strict transistor breakdown voltage. In this case, the chances of clipping can be reduced by proper circuit design (e.g., a class-AB output stage with sufficiently large transistors) instead of only by reducing S_C . The current readout can, therefore, allow for a larger coil sensitivity and, thus, better noise performance. However, it should be noted that for this to work properly, the first amplifier stage should immediately introduce a bandwidth limitation; otherwise, extremely large voltages can still occur elsewhere in the signal chain.

Another advantage of reading out the coil current is an increased bandwidth. Namely, for voltage readout and reasonable coil sizes, it turns out that the bandwidth is mainly limited by R_C and C_C . When reading out the coil current, the capacitance C_C is, to a first-order approximation, shorted. This means that the dominant pole is now determined by R_C and L_C , which generally lies significantly higher¹². As explained in detail in Appendix C, a larger coil (and thus sensitivity) goes hand in hand with a decrease in bandwidth; therefore, also from a bandwidth point of view, S_C can be made larger compared to the voltage readout case.

Unfortunately, reading out the coil current also has its disadvantages. Most importantly, current readout introduces a large temperature dependency in the gain: the temperature-independent voltage is converted to a highly temperature-dependent current via the metal resistance (with a temperature coefficient of roughly 0.4% K⁻¹). This temperature dependence should be compensated somehow.

To conclude, reading out the coil current introduces some additional design challenges and restrictions related to temperature compensation, drive capability, and the immediate bandwidth limitation of the coil. This is likely to increase system complexity compared to voltage readout. However, if these challenges can be overcome, current readout could result in large performance benefits due to the potential use of larger and, thus, more sensitive coils.

Voltage versus current readout seems to introduce a trade-off between
system **complexity** and **noise/bandwidth** performance

3.6.3. PROTOTYPE CHIP DESIGN CHOICE

For the SystematIC prototype, the main purpose is to show that by converting the Hall-Hall architecture to a coil-Hall architecture with minimum changes, considerably better noise and bandwidth performance can be obtained. As mentioned earlier, this should be done with minimum changes possible and preferably in a relatively short time to save on R&D costs. Therefore, low complexity is highly valued for this first prototype. Besides, reading out either coil voltage or current will both outperform the Hall-Hall architecture by a

¹⁰In a closed-loop sensing configuration, the magnetic field is nullified, and this problem can probably be mitigated

¹¹If the coil is read out using a transimpedance or current amplifier

¹²In simulations with a more distributed model (multiple sections of the model shown in Figure B.1), this large increase in bandwidth was also verified. Not as large as set by only R_C and L_C , but still significantly larger

	Voltage Readout	Current Readout
Maximum Sensitivity	-	++
Gain Stability	++	-
Low Readout Complexity	+	-
Inherent Bandwidth	-	+

Table 3.1: Comparison table for coil voltage- versus current readout

significant margin. Consequently, the coil **voltage** will be chosen as the information-carrying quantity. These results are summarized in comparison Table 3.1.

4

SIGNAL PATH COMBINING: IMPLEMENTATIONS

Having obtained some general knowledge on various aspects of the coil-Hall hybrid architecture in the previous chapter, this chapter will focus in more detail on how to combine the two signal paths. As will become clear, the highlighted performance aspects and cost factors from Chapter 2 form tight trade-offs that are difficult to balance optimally. This chapter aims to expose and understand these trade-offs as clearly as possible and try to navigate them by looking into three different signal path combining architectures.

These architectures will be referred to as the **shared low-pass architecture**, the **integrator architecture**, and the **hybrid-hybrid architecture**. These architectures are categorized by the (more fundamental) conceptual ideas behind them, not by the exact circuit implementation. As will become evident, no particular architecture is the best; each has its own clear advantages and disadvantages. Therefore, the main goal of this chapter is to, ideally, provide someone with sufficient knowledge to make their own decision on which battles to pick when implementing a coil-Hall hybrid system.

The chapter is structured so that each section introduces a new architecture. Several equations will be derived, and to get a feel for the numbers, some numerical examples will be given. Lastly, after discussing all architectures, one architecture will be selected for the SystematIC prototype chip.

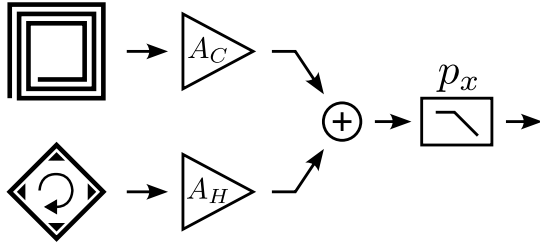


Figure 4.1: Combining architecture making use of a first-order low-pass filter in each path

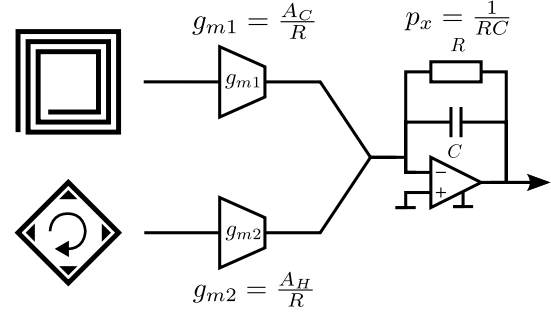


Figure 4.2: (Simplified) potential implementation of the architectural idea shown in Figure 4.1

4.1. ARCHITECTURE 1: SHARED LOW-PASS FILTER

As calculated in Section 3.3, a flat response can be obtained by (first-order) low-pass filtering both signal paths. This shared mathematical functionality immediately begs whether this can be exploited on the implementation level. Namely, to reduce area and power consumption, ideally, as many circuit blocks of the two signal paths can be shared by summing the two transducer signals as early as possible.

While theoretically possible, it is believed to be practically difficult. Consider, for example, the chopping circuitry in the low-frequency path and the potential large coil signals that must be dealt with in the high-frequency path. Combining these two paths too soon means dealing with these problems simultaneously. This would increase design complexity and reduce design flexibility, which is not considered optimal.

However, the idea behind the shared functionality is still believed to be promising when applied later in the signal chain. Namely, as f_x is generally relatively low (several kHz) it requires relatively large components to be implemented. The required amount of area can be reduced by sharing the low-pass filter between both paths. This results in the shared low-pass filter architecture, the first architecture to be considered in this chapter. This architecture is conceptually shown in Figure 4.1 with a potential implementation in Figure 4.2.

4.1.1. RIPPLE VERSUS NOISE PERFORMANCE

As explained in Section 3.4, the cross-over frequency $f_x = \frac{p_x}{2\pi}$ should be chosen such that the noise contribution of each path is equal. This would result in optimal noise performance. However, excellent noise performance is relatively pointless when the spinning ripple from the low-frequency path is to dominate it¹.

To intuitively investigate the effect of ripple, the ripple can be approximated by a sine-wave² with magnitude V_{os} , filtered by a first-order low-pass filter with a cut-off frequency f_x , and referred to the input. This attenuated ripple should be well below the integrated noise level and, thus, can be summarized by the following inequality:

$$\frac{4}{\pi} \frac{V_{os}}{A_{co}S_H} \cdot \frac{f_x}{f} \ll I_N \quad (4.1)$$

In this equation, V_{os} is the expected offset voltage at the Hall plate output (in V), f is the ripple frequency (in Hz), f_x is the cross-over frequency of the first-order filter (in Hz) and I_N the input-referred noise (in A_{RMS}). The ripple frequency f is assumed to be well above f_x such that the first-order filter can be approximated by its asymptote. In turn, using Equations 3.10 and 3.11 (describing $f_{x,opt}$ and $I_{N,opt}$), this inequality can be rewritten to yield:

$$f_x \ll \alpha \cdot \frac{1}{V_{os}} \cdot f \cdot f_{x,opt} \quad (4.2)$$

with α being a relatively complex parameter depending on temperature, the Boltzmann constant, and some transducer- and readout-dependent parameters:

¹Unless the user is willing to compensate for it externally

²First harmonic of the square-wave signal

$$\alpha = \sqrt{8\pi^2 kT} \cdot \sqrt{\frac{S_C}{S_H}} \cdot \left(\frac{R_{H,n}^3}{R_{C,n}}\right)^{\frac{1}{4}} \quad (4.3)$$

This equation states that for a ripple signal below the *optimum* integrated noise level, the cross-over frequency should be at least a factor $\alpha \cdot \frac{1}{V_{os}} \cdot f$ below the optimum cross-over frequency. Only if this factor is (considerably) larger than one, the optimum crossover frequency would already give sufficient ripple reduction, and thus, optimum noise performance can be achieved without a ripple penalty.

Assume, for example, an offset voltage V_{os} of 0.5mV (already a relatively low value) and transducer parameters $S_C = 2.5 \mu\text{VT}^{-1} \text{Hz}^{-1}$, $S_H = 80 \text{mVT}^{-1}$, $R_{H,n} = 1.2 \text{k}\Omega$ and $R_{C,n} = 1.4 \text{k}\Omega$ ³. At a temperature T of 300 K, this yields $\alpha \approx 1.07 \cdot 10^{-10}$ and results in a required chopping frequency f of at least 4.7 MHz to make the factor $\alpha \cdot \frac{1}{V_{os}} \cdot f$ larger than one. Such a high spinning frequency will almost certainly yield an offset penalty [18, 23–25].

SystematIC, for example, already uses a relatively high spinning frequency of maximally 200 kHz. For this frequency, Equation 4.2 would yield $f_x \ll 0.0428 f_{x,opt}$. According to Equation 3.12, this would result in a noise penalty of at least 10.7 dB, increasing the minimum noise from 23.6 mA_{RMS} (for a coupling factor A_{co} of $310 \mu\text{TA}^{-1}$) to roughly 81 mA_{RMS} ⁴. Besides, the filter would require considerably larger components to implement a cross-over frequency of less than 235 Hz⁵ compared to the optimum $f_{x,opt}$ of 5.5 kHz associated with the aforementioned parameter values.

There seems to exist a strict trade-off between
ripple and noise performance

It can be concluded that only the first-order low-pass filter is not an effective solution for achieving optimum noise performance. If this architecture is to be used anyway, employing additional techniques to reduce the ripple in the low-frequency path is advisable. Ripple reduction loops, notch filtering, or additional low-pass filters, for example, are just a few solutions that have proven effective. Of course, when applying additional filtering in the low-frequency path, due to additional poles, gain flatness over frequency should be carefully considered (see Section 3.5). Especially ripple reduction loops (as successfully employed in [11–14, 24]) are believed to be promising.

4.1.2. OFFSET PERFORMANCE

Assuming the ripple problem is solved, or its noise penalty is acceptable, this low-pass filter architecture still introduces a major challenge. Namely, the high-frequency path is DC-coupled to the output; this injects additional offset into the main signal path, which must be accounted for.

As shown by the equations in section 3.3, for a flat response, the high-frequency path requires a DC gain of $\frac{S_H A_H}{S_C p_x}$. Therefore, offset at the input of the high-frequency path can be referred to the input of the low-frequency path by a gain of $\frac{S_H}{S_C p_x}$. This gain is generally relatively close to one when trying to reach optimum noise performance. As a result, the high-frequency path would require similar levels of offset reduction as the low-frequency path. For the design of SystematIC, this means that the input of the high-frequency path should equivalently have less than 1.7 μV (3σ) of offset⁶. The raw offset of CMOS amplifiers is generally in the milli-volt range.

To mitigate this offset, it is believed that two strategies can be employed. First, just like in the low-frequency path, dynamic offset compensation techniques can be employed. However, as will become evident, employing such techniques in the high-frequency path might be more difficult. The second option would be to block the DC signals from the high-frequency path. This is based on the fact that the high-frequency path does not

³These values are relatively typical values associated with the transducers (coil and Hall plates), especially for the design for SystematIC. Throughout this chapter, these values will also be used in all other calculations.

⁴In this case of course, the ripple can be attenuated at bit less due to the larger amount of noise. Therefore, the corner frequency can be chosen slightly higher to yield better, but still far from optimal, noise results

⁵Resulting in roughly 58.5 dB of ripple reduction at 200 kHz

⁶See the requirements in Section 2.4, 1.7 μV results in 70mA for a coupling factor A_{co} of $310 \mu\text{TA}^{-1}$ and Hall plate sensitivity of 80mVT^{-1} (realistic values assumed throughout this chapter)

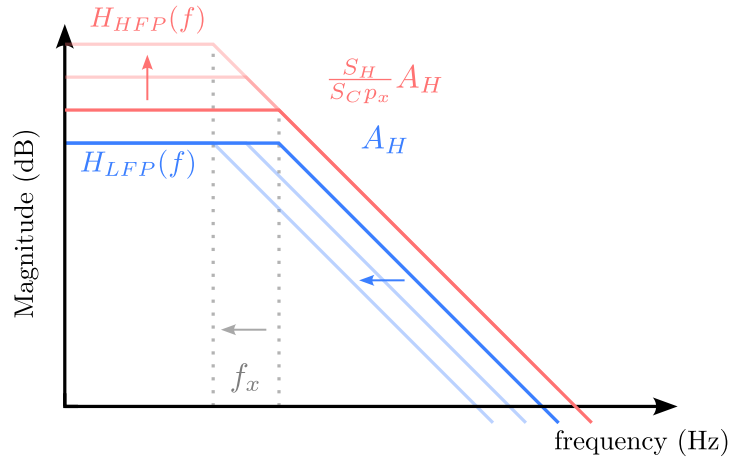


Figure 4.3: The effect of changing the cross-over frequency on the low- and high-frequency path transfers. It can be seen that a lower cross-over frequency does not help with filtering in the high-frequency path.

necessarily have to sense DC signals. However, this approach would contribute to an error in the gain flatness. The two approaches will be elaborated on in the next two sections.

DYNAMIC OFFSET COMPENSATION TECHNIQUES

Like in the low-frequency path, dynamic offset compensation techniques can be used to reduce the offset. However, unlike the low-frequency path, the effective low-pass filtering of the summing filter cannot be improved by lowering the cross-over frequency f_x . This can be shown by referring the ripple at the input of the high-frequency path back to the main input (i.e., the current domain). This input-referred ripple must be below the noise level and should thus adhere to the following inequality:

$$V_{os} \cdot \frac{S_H A_H}{p_x S_C} \cdot \frac{f_x}{f} \cdot \frac{1}{A_H S_H A_{co}} = \frac{V_{os}}{2\pi f A_{co} S_C} \ll I_N \quad (4.4)$$

This equation describes the input-referred ripple when applying a ripple with magnitude V_{os} and frequency f at the input of the high-frequency path. This equation was derived by using the required high-frequency path transfer for a flat response (as calculated in Section 3.3) to refer the ripple to the output of the high-frequency path and the system gain ($A_{sys} = A_H S_H A_{co}$) to refer the ripple back to the main input. This new inequality shows that reducing the cross-over frequency f_x does not have any effect on the filtering response at high frequencies. This is because the required high-frequency path DC gain is inversely proportional to f_x (as can be seen in Figure 4.3). Ripple can, therefore, only be further reduced by obtaining a better coupling factor A_{co} , a better coil sensitivity S_C , or a higher switching frequency f . Even worse, compared to the low-frequency path, ripple reduction loops or any other kind of filtering in the high-frequency path will not be possible, as this would directly affect the system frequency response and bandwidth.

This limited filtering is especially problematic when implementing chopping. Namely, for the transducer characteristics introduced in Section 4.1.1 and a coupling factor A_{co} of $310 \mu\text{T A}^{-1}$, a chopping frequency of at least 4.35 MHz would be required to get the ripple below the noise level of 23.6 m A_{RMS} . Chopping at this frequency will likely result in an equivalent residual offset of several hundred μV [25, 26] at the input of the high-frequency path⁷. This is too much compared to the maximum of $1.7 \mu\text{V}$. Therefore, to get the offset cancellation to work properly, likely, some more complicated techniques such as "offset-stabilization" or "ping-pong" [25] must be employed.

- The offset-stabilization technique [27] employs a high-gain, low-frequency path parallel to the controller's main signal path. In a feedback configuration, if the DC gain is high enough, this low-frequency path will reduce the main path's offset. However, this additional path must be offset-compensated itself. This can either be done through chopping (frequency domain modulation) or auto-zeroing (time domain modulation). To eliminate the chopping ripple, some relatively complex filtering must be implemented inside this path (as the external filtering is insufficient). Conversely, auto-zeroing introduces

⁷Due to effects such as charge injection and clock-feedthrough

problems associated with increased noise due to folding, potential ripple due to switching transients, and signal aliasing.

- The ping-pong technique also uses auto-zeroing but, if properly implemented, does not have the issue of signal aliasing [28]. With this technique, two amplifiers are being used alternately (space and time domain modulation). While one is used, the other is offset-compensated using the auto-zeroing technique. As long as the two amplifiers are well-matched, there is always an amplifier available in the signal path and the continuous-time operation is not compromised. However, noise folding due to the required auto-zeroing and potential ripple due to switching transients could still be an issue. Besides, double the area and power consumption are required for the two amplifiers being used alternately with this technique.

All in all, dynamic offset compensation could be a valid option in the high-frequency path. However, due to this path's inherently limited filtering capability, great care must be taken with potential ripple due to switching in the signal path. This adds significant complexity.

Dynamic offset compensation in the high-frequency path adds significant **complexity** due to severely limited **design freedom** in this path's filtering capability

BLOCKING THE DC SIGNAL

As concluded in the previous section, applying dynamic offset compensation techniques in the high-frequency path can become relatively complex. Besides, as DC current is already accurately measured in the low-frequency path, there is no particular need to do it in the high-frequency path. Therefore, to mitigate the offset problem, the second solution would be to block all DC signals in the high-frequency path.

Of course, this comes with its own challenges and costs. Blocking low-frequency signals would result in a high-pass transfer in the high-frequency path and, as explored in Section 3.5, the resulting low-frequency pole affects the gain flatness of the total system's frequency response. As shown by Equation 3.17, for an error of less than 2%, the pole frequency must be at least a factor 50x below the cross-over frequency f_x of the first-order summing filter. In this case, for optimum noise performance, the cross-over frequency should be at roughly 5.5 kHz, meaning that the high-pass transfer is only allowed to have a pole at maximally 55 Hz, quite a low frequency.

If this pole were to be somehow implemented using an RC time constant (for example, using an AC coupling capacitor), a large resistance of, for example, 10 M Ω would be required in combination with a 290 pF capacitance. This requires a serious amount of area (see, for example, the implementation in [12]). Luckily, this area can be largely reduced using a low-frequency feedback loop (i.e., DC servo loop), nullifying any DC (and low-frequency) signal at the output of the high-frequency path. The bandwidth of this loop should be limited to, in this case, maximally 55 Hz. This likely still requires relatively large components. However, additional attenuation (e.g., resistive division) in the loop can help reduce the required component sizes compared to the AC coupling case. This was indeed verified by the implementation shown in [13].

The downside of this DC servo loop is the additional active circuitry, which would add noise and offset. However, in this case, the offset can be more easily mitigated as the feedback loop provides additional filtering at higher frequencies, making chopping, for example, a valid option again.

Blocking DC signals in the high-frequency path ultimately introduces a trade-off between **gain flatness** and **area**

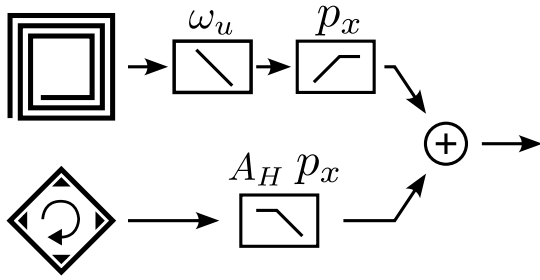


Figure 4.4: Combining architecture making use of an integrator and a low-pass high-pass summing filter

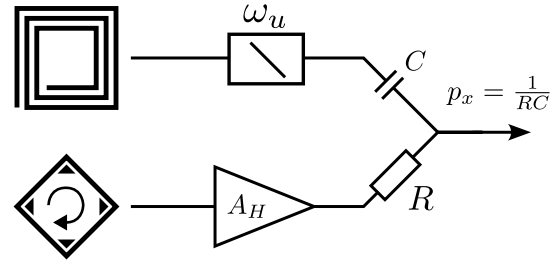


Figure 4.5: (Simplified) potential implementation of the architectural idea shown in Figure 4.4

4.2. ARCHITECTURE 2: INTEGRATOR AND HIGH-PASS FILTER

The second architecture consists of an integrator in the high-frequency path to flatten the coil's response and a summing high-pass/low-pass filter to combine the two signal paths. The main idea behind this architecture is the high-pass filter's inherent DC blocking capability: the high-frequency path no longer contributes any offset. Besides, this architecture allows for a simple implementation of a second-order filter, which is believed to slightly alleviate the trade-off between ripple and noise performance. The architecture is shown in Figure 4.4 with a potential implementation of the summing filter in Figure 4.5.

4.2.1. SECOND-ORDER SUMMING FILTER

Until now, all calculations have assumed a first-order low-pass transfer in both the low- and high-frequency path. However, as explained in Section 3.3, any transfer function can be chosen in the low-frequency path as long as the high-frequency transfer adheres to Equation 3.5. As previously concluded, a first-order implementation creates a strict trade-off between ripple and noise performance. This section will investigate if a (more complex) second-order transfer could alleviate this trade-off.

For this purpose, the low-frequency path transfer will be written as shown in Equation 4.5:

$$H_{LPF}(s) = \frac{A_H}{(1 + \frac{s}{p_x})(1 + \frac{s}{p_y})} \quad (4.5)$$

This transfer has a DC gain, A_H , and two poles at $s = -p_x$ and $s = -p_y$ with $|p_x| < |p_y|$. Furthermore, to simplify the high-frequency path transfer, just like in all previous calculations, equal magnetic field coupling factors are assumed between the coil and Hall plates ($A_{co,C} = A_{co,H} = A_{co}$), and the system gain is set equal to the DC gain in the low-frequency path ($A_{sys} = A_{co}S_H A_H$). With these assumptions, the required high-frequency transfer can be written as:

$$\begin{aligned} H_{HFP}(s) &= \frac{S_H A_H}{S_C p_x} \cdot (1 + \frac{p_x}{p_y}) \cdot \frac{1}{s} \cdot \frac{s(1 + \frac{s}{p_x + p_y})}{(1 + \frac{s}{p_x})(1 + \frac{s}{p_y})} \\ &= \frac{S_H A_H}{S_C p_x} \cdot (1 + \frac{p_x}{p_y}) \cdot \frac{1 + \frac{s}{p_x + p_y}}{1 + \frac{s}{p_y}} \cdot \frac{1}{1 + \frac{s}{p_x}} \end{aligned} \quad (4.6)$$

Compared to the first-order case (i.e., $p_y = \infty$), as shown in Equation 4.7⁸

$$H_{HFP(1st\ order)}(s) = \frac{S_H A_H}{S_C p_x} \cdot \frac{1}{1 + \frac{s}{p_x}}, \quad (4.7)$$

it can be seen that the required second-order transfer is very similar except for the additional $(1 + \frac{p_x}{p_y}) \cdot \frac{1 + \frac{s}{p_x + p_y}}{1 + \frac{s}{p_y}}$ term.

⁸This equation is repeated here from Section 3.3 for the sake of clarity

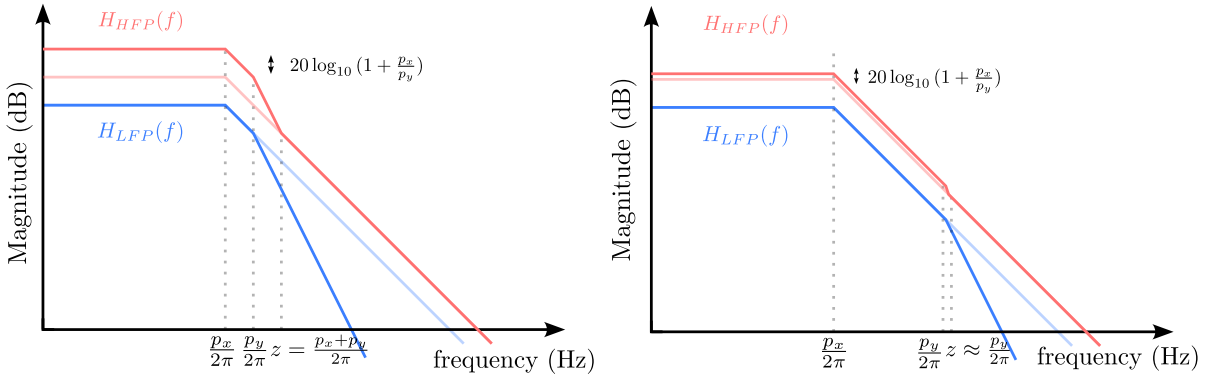


Figure 4.6: Second-order low-pass transfer in the low-frequency path with the associated high-frequency path transfer to obtain a flat response (further pole spacing shows that the response, as expected, converges to first-order behavior)

This term essentially shows that the second-order high-frequency path transfer requires a slightly larger gain at low frequencies compared to the first-order case. This can be quite intuitively understood by considering that the two poles result in sharper filtering in the low-frequency path, and thus, to get an overall flat response, the high-frequency path should provide more gain at low frequencies to compensate. The closer the poles are, the larger the gain becomes (until a two times increase for $p_x = p_y$). On the other hand, when the poles move away from each other, the required transfer converges to the first-order case, and no additional gain is required. This effect is graphically depicted in Figure 4.6.

It can be concluded that the increased filtering in the low-frequency path comes at the cost of an increased low-frequency gain in the high-frequency path. This causes the optimum noise cross-over frequency $f_{x,opt}$ to shift to a higher frequency. Besides, simulations have shown that it causes a slight increase in optimum noise of at most 0.5 dB (for $p_x = p_y$). However, the increased filtering could help to reduce the ripple in the low-frequency path more efficiently.

Obtaining the theoretical transfer function is one thing; implementing it is another. Especially the high-frequency path transfer function in Equation 4.6 seems considerably harder to implement than a first-order transfer. Different poles and zeros need to be matched accurately in both signal paths. This complexity is why no higher-order transfer functions were considered in the previous architecture.

However, in this architecture, flattening the coil response first, before summing it to the low-frequency path, simplifies the second-order implementation greatly. Namely, in this case, the passive summing network shown in Figure 4.7 can be used. As long as there are no parasitic capacitances or resistances to the ground, it can be seen that the same signal applied at both inputs must result in a flat response at the output. Yet, each path independently sees a different transfer to the output. If the desired second-order low-pass transfer is implemented in the low-frequency path, this summing filter will "automatically" implement the required high-pass response to flatten the overall response (see Figure 4.8). Together with the integrator in the high-frequency path, the required transfer functions from Equations 4.5 and 4.6 are implemented perfectly. The only downside is that any parasitic elements to the ground (e.g., see C_{par} in Figure 4.7) could cause an error in the gain flatness; this must be carefully considered when deciding to use this architecture.

The advantages of this architecture over the first-order case are best described by studying the same example used in Section 4.1.1. Namely, in that section, it was concluded that to reduce the ripple caused by a 0.5 mV offset below the noise of 23.6 mA_{RMS} (for some typical transducer parameters), a first-order filter with

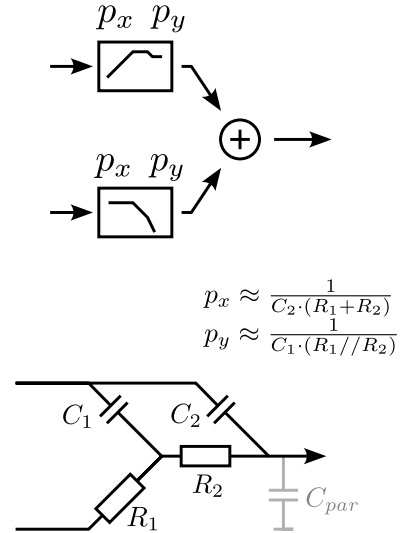


Figure 4.7: Simple implementation of the second-order summing filter

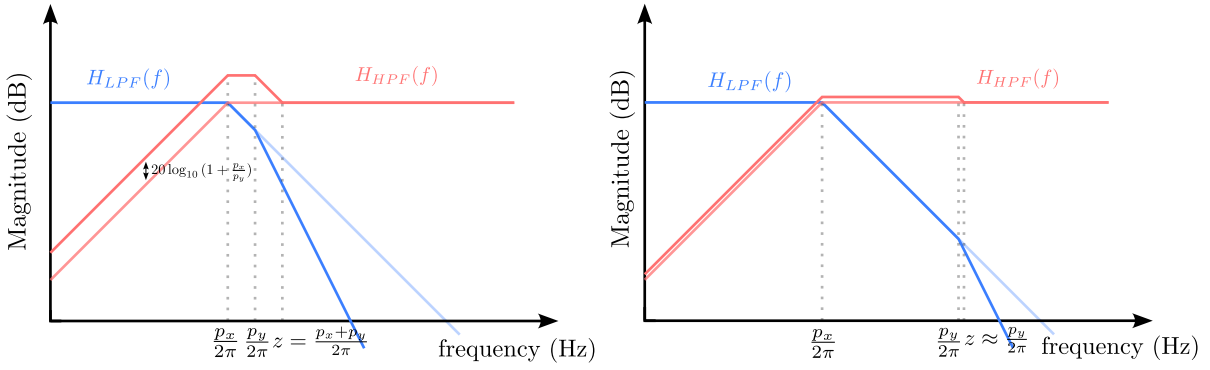


Figure 4.8: The low- and high-pass frequency response from the filter shown in Figure 4.7

a cross-over frequency of at most 235 Hz would be required. This filter results in an attenuation of roughly 58 dB at a frequency of 200 kHz. If this filter were to be implemented using a resistance of 2.5 M Ω , it would require an enormous capacitance of 270 pF. However, with a second-order filter implementation, poles at roughly 1.3 kHz and 40 kHz result in the same attenuation at 200 kHz (see Equation 4.8). Compared to the 10.7 dB noise penalty for the first-order case, this new cross-over frequency f_x of 1.3 kHz results in only a penalty of 3.5 dB⁹. Moreover, to implement these poles, with a total resistance of 2.5 M Ω ($R_1 = 500$ k Ω and $R_2 = 2$ M Ω), only a total capacitance of 60 pF ($C_1 = 10$ pF and $C_2 = 50$ pF) would be required.

It seems that a second-order filter **alleviates**
the harsh trade-off between **ripple** and **noise** performance and introduces an **area** benefit

This simple example shows that a second-order filter implementation can, in certain cases, provide large-area and noise benefits. However, to reach optimum noise performance, ripple reduction loops or notch filtering are still considered a better option. This is because noise and ripple performance can not yet be designed completely orthogonally with this second-order filter approach. There are simply not enough degrees of freedom. Still, an advantage of this approach is its relatively low complexity if implemented as shown in Figure 4.7.

$$\frac{V_{os}}{A_{co}S_H} \cdot \frac{f_x f_y}{f^2} < I_N \quad (4.8)$$

4.2.2. INTEGRATOR DC GAIN

An integrator can never be implemented in an ideal fashion. Namely, an ideal integrator would imply infinite gain at DC, which is impossible. As a result, the integrator has its DC pole shifted to a higher frequency. This is an unwanted pole in the high-frequency path, and, as explained in Section 3.5, this influences the gain flatness of the overall system response.

With the integrator transfer written as $\frac{\omega_u}{s}$ and $\omega_u = \frac{S_H A_H}{S_C}$ to make the low- and high-frequency path gains equal, the minimum required DC gain can be written as:

$$A_{DC} > \frac{\omega_u}{p_x} \cdot \zeta = \frac{S_H A_H}{S_C p_x} \cdot \zeta \quad (4.9)$$

This equation is simply deduced from the notion that the integrator must be able to at least provide the gain required at a frequency of $\frac{p_x}{\zeta}$ to move the unwanted pole below that frequency. p_x is the summing filter cross-over frequency (in rad s⁻¹) while ζ indicates how much the parasitic pole must be distanced from this

⁹This penalty due to deviating from the optimum cross-over frequency can still be estimated using Equation 3.12 even though now a second-order filter is used. This can only be done if the spacing between the poles is relatively large (i.e., $p_y \gg p_x$). If the pole-spacing is smaller, the optimum cross-over frequency moves to a higher frequency and the benefits are smaller

cross-over frequency. As explained in Section 3.5, the error in the overall gain flatness is roughly inversely proportional to ζ (e.g., $\zeta = 100$ results in a 1% error).

Furthermore, by substituting the cross-over frequency $f_{x,opt} = \frac{p_{x,opt}}{2\pi}$ required for optimum noise performance (see Section 3.4), the DC gain requirement can be written as:

$$A_{DC} > A_H \cdot \zeta \cdot \sqrt{\frac{R_{H,n}}{R_{C,n}}} \quad (4.10)$$

With a 2% gain flatness requirement ($\zeta = 50$), roughly equal noise resistances $R_{H,n}$ and $R_{C,n}$ and a low-frequency path gain A_H of 100x (20 dB¹⁰), a minimum DC gain of at least 5000x (74 dB) would be required. A small offset of only ± 0.5 mV would already result in a ± 2.5 V output signal. This will likely result in clipping of the integrator, especially for supply voltages below 5 V.

A flatter frequency response requires a larger DC gain

4.2.3. INTEGRATOR IMPLEMENTATION CONSTRAINTS AND CHALLENGES

For reasonable gain flatness and gain in the low-frequency path, the integrator requires such a large DC gain that its output likely clips on its own offset. Although it does not propagate to the system output, it is still important to mitigate this offset. The main difference compared to the previous architecture is that the residual offset can be considerably larger as long as the output does not clip. This allows for some slightly different solutions/strategies. These solutions and associated trade-offs will be elaborated on in this section.

OFFSET TRIMMING

A simple solution would be to do a one-time trim of the offset. While this would initially mitigate the clipping problem, the offset can still drift over temperature and lifetime. Moreover, with the large integrator DC gain, only a small offset variation could cause the system to clip again. To ensure system robustness, it is likely wise to use a more advanced technique to mitigate the clipping problem.

DYNAMIC OFFSET COMPENSATION TECHNIQUES

The next potential strategy would be reducing the offset such that the integrator no longer clips. To also account for offset drift over lifetime and temperature, this can be done with dynamic offset compensation techniques. While this option is valid, it poses similar challenges as discussed for the previous architecture (see Section 4.1.2). This is because the extremely limited filtering constraints in the high-frequency path are independent of the implementation, and therefore, the integrator architecture does not introduce any benefits in that regard. Also, it should be noted that the second-order filter implementation in the low-frequency path does not provide higher-order filtering in the high-frequency path and, therefore, does not help either.

DISCRETE AMPLITUDE FEEDBACK LOOP

As the exact voltage at the output of the integrator does not matter, one could also consider a very rough feedback loop whose sole purpose is to move the integrator output within a certain range. Such a feedback loop is conceptually depicted in Figure 4.9. The exact implementation is not important but the main idea is that two comparators sense whether the output of the integrator goes out of bounds. If that happens, a digital-to-analog converter (DAC) can generate a compensating signal.

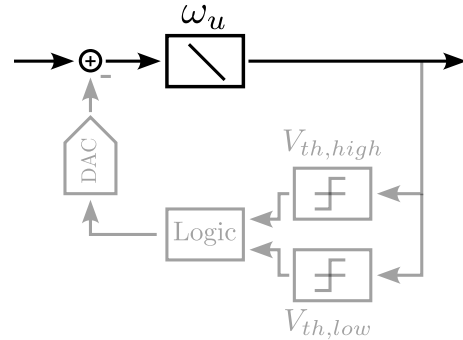


Figure 4.9: Conceptual representation of the discrete amplitude feedback loop used to very crudely set the output voltage and prevent the integrator from clipping

¹⁰ A smaller gain puts stricter requirements on offset, (wide-band) noise and interference requirements of the common signal path after combining. Generally, it can be assumed that 20 dB is a safe minimum

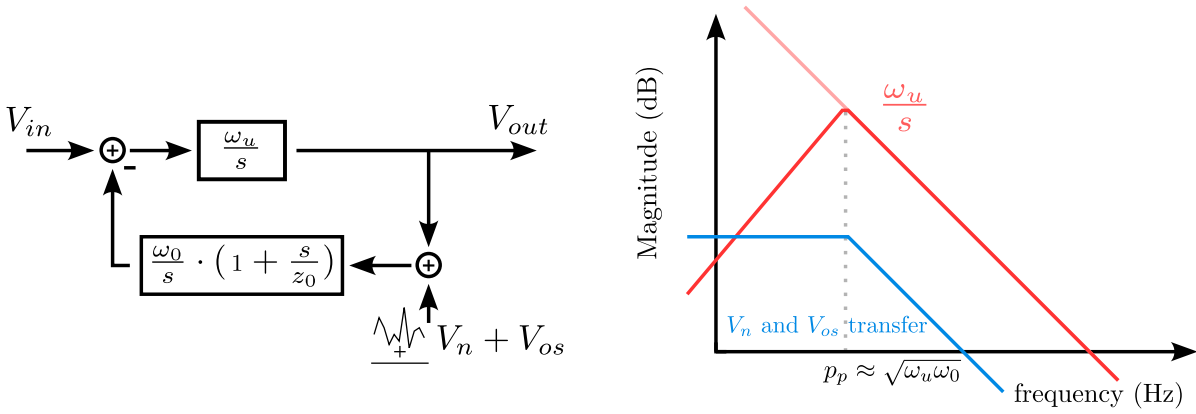


Figure 4.10: DC Servo loop concept showing the input signal V_{in} , output signal V_{out} , and a potential interfering noise V_n or offset V_{os} signal injected at the input of the feedback network. On the right the resulting signal transfers are shown. Noise is heavily filtered and offset sees (depending on potential attenuation in the loop) a gain of 0 dB to the integrator output.

This simple solution does have a downside however. Namely, the sharp discrete amplitude steps from the DAC will be visible as spikes at the output. Considering that any signal at the input of the high-frequency path sees a low-pass transfer given by Equation 4.6 or 4.7, a step (with magnitude V_{step}) at the input will essentially see the DC gain of the high-frequency path¹¹. Referred to the main input (current on the current rail), the magnitude of the spikes can be estimated as:

$$V_{step} \cdot \frac{S_H A_H}{S_C p_x} \cdot \frac{1}{S_H A_H A_{co}} = V_{step} \cdot \frac{1}{S_C p_x A_{co}} \ll I_N \quad (4.11)$$

Ideally, these spikes should be below the noise level.

Furthermore, one could substitute $f_x = f_{x,opt}$ for optimum noise performance. This results in the following equation:

$$V_{step} \cdot \frac{1}{S_H A_{co}} \cdot \sqrt{\frac{R_{H,n}}{R_{C,n}}} \ll I_N \quad (4.12)$$

For the realistic parameter values given in the previous sections, the factor $\frac{1}{S_H A_{co}} \cdot \sqrt{\frac{R_{H,n}}{R_{C,n}}}$ becomes approximately 37300 A V^{-1} . To get the spike below the noise level of 23.6 mA_{RMS} , the step size is not allowed to exceed 630 nV . Furthermore, to compensate for the offset of the integrator, the DAC must be able to span a full range of a few millivolts. For example, for a range of $\pm 3 \text{ mV}$, the DAC would require at least 13 bits to keep the spikes below the noise level¹². This adds complexity to this seemingly simple solution.

Due to the required low-pass characteristic in the high-frequency path transfer, it can generally be concluded that fast transients must be handled carefully. This became evident in this section but certainly also holds for any other solution requiring switching. Only if the switching results in a steady-state periodic signal (e.g., when chopping), the signal is largely attenuated by the low-pass characteristic. It does help to use a large coil sensitivity or a higher cross-over frequency (to reduce the high-frequency path DC gain), but the latter might not be optimal for noise performance.

Extreme care must be taken with non-periodic switching transients in the high-frequency path

¹¹The time constant associated with the unwanted integrator pole is considerably larger than the summing filter low-pass transfer time constant. Although the DC signal will eventually be filtered, the initial spike will not see much filtering due to the (unwanted) high-pass transfer.

¹²Some rough simulations also verified this result

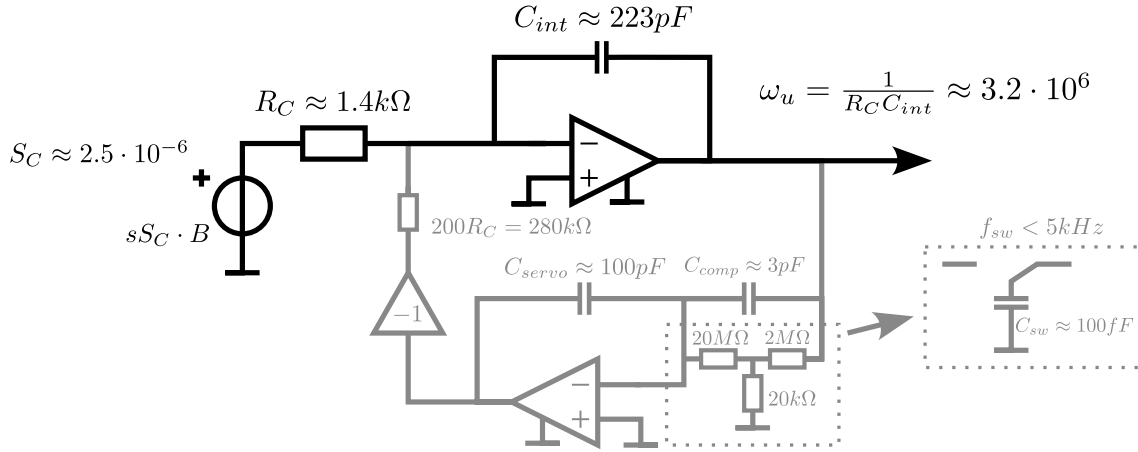


Figure 4.11: Potential DC servo loop implementation based on the parameter values used throughout this chapter (the response was also verified by some basic simulations)

CONTINUOUS-TIME DC SERVO LOOP

The last strategy would be to implement a low-frequency feedback loop to accurately regulate the DC voltage at the output of the integrator. This can be done by placing another integrator in a loop around the main integrator. To stabilize the loop, also a zero needs to be implemented. This configuration is conceptually depicted in Figure 4.10. This figure also shows the overall transfer function and the transfer of any noise or offset injected at the input of the feedback network to the overall output¹³.

The main challenge with this approach is to ensure sufficient gain flatness. Namely, the feedback loop will introduce two low-frequency poles, as shown by the response in Figure 4.10. As calculated in Section 3.5, this results in a dip in the frequency response¹⁴.

These unwanted poles must be located below $\frac{p_x}{\gamma}$ if a gain flatness error of maximally $\frac{100}{\gamma}\%$ is allowed. Furthermore, the poles will appear at the loop gain unity gain frequency, which occurs at $\sqrt{\omega_u \omega_0}$ (rad) if the effect of the zero is neglected. If the zero were to be placed at $z_0 = \frac{\sqrt{\omega_u \omega_0}}{\alpha}$ with $\alpha > 1$ for sufficient stability, it can be proven that the highest frequency pole moves to roughly $\alpha \cdot \sqrt{\omega_u \omega_0}$. Using this knowledge, the following constraint can be written from a gain flatness point of view:

$$\omega_0 < \frac{p_x^2}{\gamma^2} \cdot \frac{1}{\alpha^2} \cdot \frac{1}{\omega_u} \quad (4.13)$$

Which, after substituting $p_x = p_{x,opt}$ and $\omega_u = \frac{S_H A_H}{S_C}$, results in:

$$\omega_0 < \frac{1}{\gamma^2} \cdot \frac{1}{\alpha^2} \cdot \frac{S_H}{S_C} \cdot \frac{R_{C,n}}{R_{H,n}} \cdot \frac{1}{A_H} \quad (4.14)$$

For $\gamma = 50$, $\alpha = 2$, $A_H = 100$, and the previously used transducer parameters, this equation results in a requirement of $\omega_0 < 0.037$ rad. Considering that ω_0 is inversely proportional to a product of resistance and capacitance, this extremely small value requires large components to implement.

To get a grasp of the component sizes, a potential implementation for the obtained values is shown in Figure 4.11. Even by introducing a 200x attenuation via the $200R_C$ resistance and a 100x attenuation via the 20 kΩ and 2 MΩ resistive divider, large components are still required. The attenuation factors cannot be made much larger as they are limited by drive capability and input referred noise and offset¹⁵ constraints respec-

¹³This is highly filtered as the loopgain quickly drops at low frequencies

¹⁴It was simulated that two poles instead of one do not introduce a significant difference on the calculated dip, as long as the frequency compensation is adequate and no significant peaking occurs.

¹⁵ V_{os} resulting from the integrator controller effectively sees a gain of 100x to the output of the main integrator (i.e., V_{os} in Figure 4.10, now sees a gain larger than 0 dB)

tively.

Implementing a DC servo loop around the integrator requires a significant amount of **area**

DISCRETE-TIME DC SERVO LOOP

Lastly, one could consider implementing the large resistors from the continuous-time implementation with a switched capacitance (see Figure 4.11). With a sufficiently low switching frequency (f_{sw}), the feedback loop should not affect gain flatness much. As a result, no extremely large components are required anymore, and the required area can be reduced. However, the downside is that this approach again introduces switching and sampling into the high-frequency path. Consequently, the sampling causes signal aliasing and, as a result, signal-dependent spikes. Essentially, the same spike problem from the "discrete amplitude feedback loop" strategy arises again¹⁶. A relatively large capacitance C_{servo} and additional attenuation in the loop are likely still required to mitigate this problem. This approach will probably still result in a large area advantage compared to a continuous time implementation. However, the designer must carefully consider whether the aliased signal-dependent voltage spikes are an acceptable penalty to pay for the reduced area.

¹⁶One could even consider this approach to be an implementation of the discrete amplitude feedback loop concept

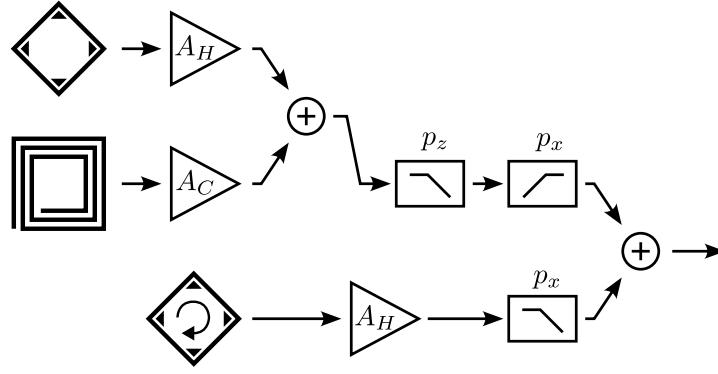


Figure 4.12: Conceptual representation of the so-called "hybrid-hybrid" architecture. This architecture uses a coil-Hall hybrid sensor in the high-frequency path and uses a high-pass low-pass summing filter with the same advantages as the integrator architecture described in the previous section.

4.3. ARCHITECTURE 3: "HYBRID - HYBRID" ARCHITECTURE

After thoroughly investigating the previous two architectures, it can be concluded that various tight trade-offs emerge due to the close dependency of many different design aspects. Simultaneously managing constraints related to noise performance, ripple reduction, offset, and gain flatness greatly increases system complexity. The idea behind this last architecture, the so-called "hybrid-hybrid" architecture, is to break these tight trade-offs and allow for better orthogonal design of these aspects.

For this purpose, this last and newly invented architecture uses a coil **and** a Hall plate instead of only a coil in the high-frequency path (see Figure 4.12). In other words, another coil-Hall hybrid current sensor is used as part of the overall hybrid sensor. The key advantage? This new high-frequency hybrid sensor is not used to process DC signals anymore, and therefore, offset is no longer important. This also means that the new Hall plate does not have to be spun, and no associated ripple has to be dealt with. This architecture essentially trades design complexity (and the associated problems) for power consumption (required for biasing the additional Hall plate).

4.3.1. A BRIEF REMARK ON THE IMPLEMENTATION

Although the high-frequency path of the hybrid-hybrid architecture can be implemented using either of the two architectures described in the previous sections, it is believed that the shared low-pass filter approach (see Section 4.1) is the best. This is because its challenges are mainly associated with offset and ripple, which are unimportant in the high-frequency path of this architecture. Therefore, this architecture's implementation becomes significantly less complex and can be implemented more efficiently. On the other hand, if the integrator architecture would be used, the same issues associated with clipping would need to be dealt with.

Furthermore, it should be noted that the general equations derived in Chapter 3 no longer hold for this architecture. This is because these equations were based on the fact that only a coil would be used in the high-frequency path. Still, these equations do describe how the hybrid sensor inside the high-frequency path behaves. Namely, for a flat response, the gain A_C must adhere to $A_C = \frac{S_H A_H}{S_C 2\pi f_z}$, with $f_z = \frac{p_z}{2\pi}$ being the cross-over frequency in the high-frequency path (see Figure 4.12).

4.3.2. NOISE PERFORMANCE

To evaluate the noise performance of the hybrid-hybrid architecture, the equations derived in Section 3.4 can not be used anymore. This is because the high-frequency path now sees a bandpass filter (low-pass filter with cross-over f_z and high-pass filter with cross-over f_x) instead of a low-pass filter. This makes deriving some general equations considerably more complex. However, the noise performance can still be approximated intuitively. For this purpose, it is educational to consider the noise contribution of each path independently.

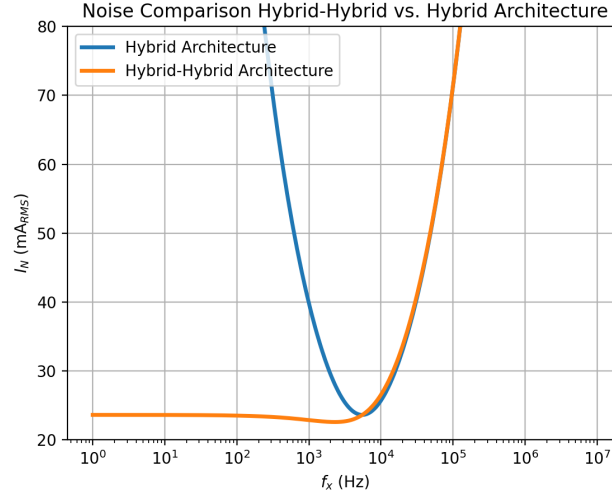


Figure 4.13: Comparison of noise as a function of cross-over frequency f_x between the original and hybrid-hybrid architecture (described by Equation 4.17). As expected, for the hybrid-hybrid architecture (with $f_z = f_{z,opt}$), lowering the cross-over frequency does not cause an increase in noise, while the original architecture has a narrow optimum.

In the high-frequency path, the noise can still be calculated using the equations obtained in Section 3.4. For the optimum cross-over frequency $f_{z,opt} = \frac{1}{2\pi} \frac{S_H}{S_C} \sqrt{\frac{R_{C,n}}{R_{H,n}}}$ the noise contribution from the high-frequency path can be written as:

$$I_{N,HFP}^2 = \frac{2kT}{A_{co}^2} \cdot \frac{\sqrt{R_{C,n}R_{H,n}}}{S_C S_H} \cdot F(f_x) \quad (4.15)$$

This is the equation for optimum noise in a hybrid architecture multiplied by a function $F(f_x)$. The function $F(f_x)$ describes the effect of the additional high-pass filtering from the summing filter. For a small f_x , this term is roughly 1 (i.e., no filtering), while it decreases for larger values of f_x .

For the low-frequency path, using the equivalent noise bandwidth of the first-order filter, the noise contribution can be written as:

$$I_{N,LFP}^2 = \frac{4kT}{A_{co}^2} \cdot \frac{R_{H,n}}{S_H^2} \cdot \frac{\pi}{2} f_x \quad (4.16)$$

In this case, a larger f_x increases noise due to the decreased low-pass filtering of the summing filter. The total noise (as shown in Equation 4.17), therefore, consists of a term that increases and a term that decreases with f_x . This noise "balancing act" looks familiar to the original hybrid current sensor (of which the noise is described by Equation 3.9). However, in this case, the term that decreases with f_x is no longer inversely proportional to f_x but is described by the function $F(f_x)$, which is limited to 1 for small values of f_x (instead of going to infinity). Therefore, f_x can now be decreased without a noise penalty.

A numerical evaluation of the transfer functions validated this conclusion. The results in Figure 4.13 indeed show that the hybrid-hybrid architecture (orange curve) allows for f_x to be decreased without a noise penalty; this is not the case for the original architecture (blue curve)¹⁷.

$$I_N^2 = \frac{2kT}{A_{co}^2} \cdot \left(\frac{\sqrt{R_{C,n}R_{H,n}}}{S_C S_H} \cdot F(f_x) + \frac{R_{H,n}}{S_H^2} \cdot \pi f_x \right) \quad (4.17)$$

The main advantage of decreasing f_x without a noise penalty is that this architecture breaks the tight trade-off between ripple and noise performance. This is essentially done by adding a degree of freedom in the design:

¹⁷The plots were made with $f_z = f_{z,opt}$ and the transducer parameter values used throughout this chapter

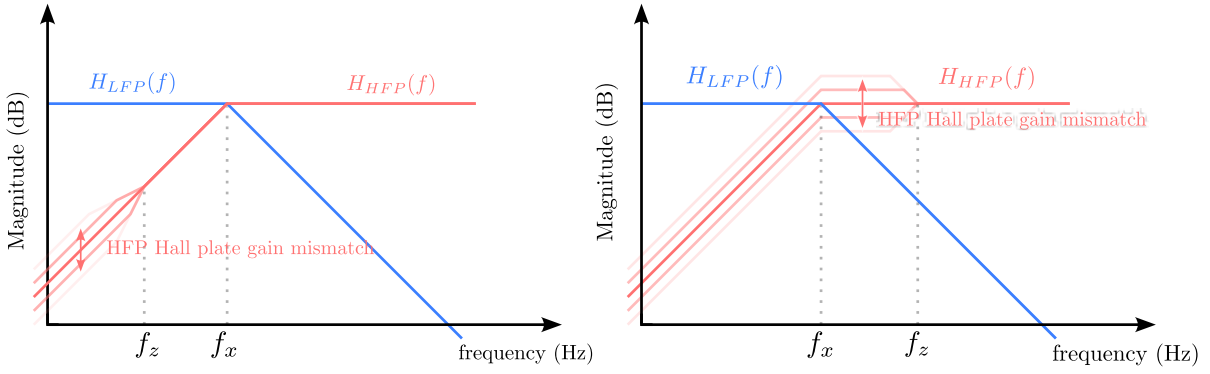


Figure 4.14: Graphical representation of the effect of mismatch between the low-frequency and high-frequency path Hall plate sensitivity. For $f_z < f_x$ (left figure), the effect is small, while for $f_z > f_x$ (right figure), the effect is large as the mismatch can still be seen when the high-frequency path takes over the overall response.

f_z can be chosen for optimum noise performance, while f_x can be chosen for sufficient ripple reduction¹⁸. Besides, this architecture can also be combined with the second-order summing filter introduced in Section 4.2.1 to obtain more ripple reduction for a smaller area penalty.

The hybrid-hybrid architecture breaks the trade-off
between **ripple** and **noise** performance at the cost of increased **power consumption**

4.3.3. GAIN FLATNESS

As will be extensively investigated in the next chapter, due to the two completely different sensing mechanisms, a downside of the coil-Hall hybrid current sensor is the requirement for (at least) two calibration steps to obtain an accurate gain over the entire frequency range of interest. As suggested in Section 2.3, this increases the manufacturing costs. While the hybrid-hybrid architecture greatly simplifies the design in many aspects, the additional Hall plate effectively adds a third signal path. Therefore, a disadvantage of this architecture would, at first sight, be increased calibration complexity and, as a result, larger manufacturing costs.

While this is true to a certain degree, there is a catch. Namely, the additional signal path is not based on some new sensing mechanism. On the contrary, just like in the low-frequency path, the additional high-frequency signal path uses Hall plates. As long as the sensitivity of the different Hall plates (and subsequent gain stages) is matched, no additional calibration is required¹⁹.

Potential gain differences in both (low- and high-frequency) Hall paths can come from mismatches in the Hall plate bias currents, mismatches between the gains A_H in both paths, or, of course, mismatches in the current-related sensitivity S_I of the Hall plates (as defined in Section 3.1). The first two sources of mismatch can be largely mitigated by accurate layout. However, the latter is unfortunately prone to effects such as on-chip stress (due to the piezo-Hall effect [29]) and surface charge traps at the oxide interface [5]. Stress-related sensitivity mismatch can be mitigated by placing the Hall plates close to each other such that they experience similar levels of stress. However, a random mismatch of roughly 1% still remains due to the charge traps at the oxide interface. If this mismatch is still too much, one could consider using buried Hall plates to move away from this "dirty" surface [5].

Lastly, it should be noted that a certain mismatch between the two Hall plates does not necessarily have to result in the same gain error in the overall frequency response. Namely, the effect on the overall frequency response highly depends on the relative location of the different cross-over frequencies f_x and f_z . If $f_z < f_x$, the coil (which is assumed to be calibrated) takes over from the Hall plate in the high-frequency path before the high-frequency path even dominates in the overall response. This means that any gain mismatch in the high-frequency Hall path does not have much of an effect on the overall response (this is graphically depicted

¹⁸Note that the noise of the Hall plates was assumed to be equal in both signal paths. However, if f_x is chosen sufficiently low such that the low-frequency path noise is insignificant, power can be saved by placing fewer Hall plates in parallel in the low-frequency path

¹⁹And even if it is needed, it can likely be done in the same step used for calibrating the low-frequency Hall path: both can be calibrated using a DC signal (as opposed to the coil). This will be further discussed in the next chapter.

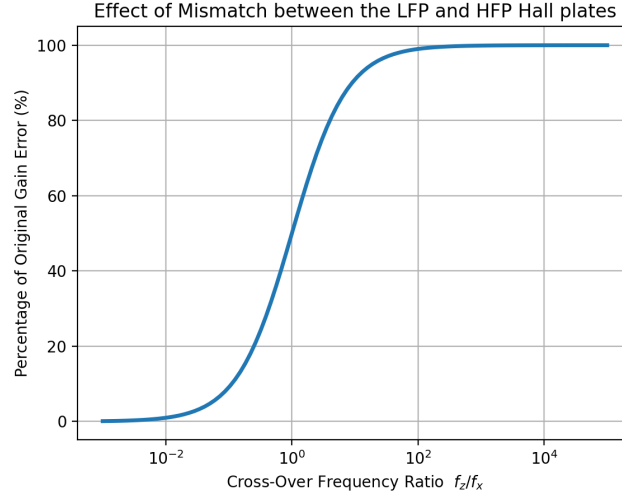


Figure 4.15: Numerical evaluation of the effect of a gain mismatch between the low- and high-frequency Hall paths as a function of f_z/f_x . The percentage indicates how much of the original mismatch is seen as a gain flatness error in the overall response. For example, 10% indicates that a mismatch error of 5% will only be seen as a maximum error of 0.5% in the gain flatness.

in the left plot of Figure 4.14). On the other hand, if $f_z > f_x$ the contrary is true, and the resulting gain error will be largely seen in the overall response (as depicted in the right plot of Figure 4.14). Figure 4.15 shows this effect quantified. Namely, this plot shows that for small values of the ratio f_z/f_x , as expected, only a certain percentage of the total mismatch error is seen in the overall response²⁰, while for large values, eventually, 100% of the error will be seen in the overall response.

**Accurate matching between the low- and high-frequency Hall paths
obviates the need for an additional calibration step**

4.3.4. OFFSET AND CLIPPING

In the integrator architecture, no offset from the high-frequency path could propagate to the overall output, yet the offset still formed an issue. Namely, because of its large DC gain, the integrator would clip on its own offset. In the hybrid-hybrid architecture, the DC gain is considerably smaller than the required integrator DC gain, and clipping should be less of an issue. However, it is still important to ensure that this is indeed the case.

The maximum amount of offset allowed at the output of the high-frequency path is ultimately limited by the signal swing and the maximum output range. As long as the output referred offset in combination with a full-scale signal does not cause clipping, no additional mitigation measures need to be taken. Therefore, the chances of clipping are reduced by:

- Allowing for a larger output range (ultimately limited by the supply voltage)
- Reducing the signal swing (by choosing a smaller gain A_H ²¹)
- By minimizing the output referred offset (by choosing a smaller gain A_H and by using large and well-matched transistors²²)

As can be seen, reducing the gain A_H is an effective way to reduce the chances of clipping and, as a result, the need for a more complex solution. Unfortunately, A_H cannot be decreased indefinitely. Namely, the gain should be sufficiently large such that offset, wide-band noise, and other interference from the common path (after summing) becomes negligible. The smaller A_H , the stricter the requirements for the common path.

²⁰For example, for $f_z = f_x$, only 50% of the error is seen as an error in the gain flatness. The 1% error due to charge traps as discussed before will then effectively only yield a 0.5% error in the overall response.

²¹ A_C is proportional to A_H and therefore also reduces when choosing a smaller A_H

²²The Hall plate offset can be minimized by placing multiple Hall plates in parallel with orthogonal biasing [5, 30].

Choosing the gain A_H is, therefore, a balancing act between reducing clipping chances and complexity in the common path.

If, for a particular design, it still turns out that additional clipping-prevention measures need to be taken, all techniques as previously mentioned for the integrator architecture can also be considered here (see Section 4.2.3). However, the same problems associated with each strategy still have to be dealt with. The only advantage of the hybrid-hybrid architecture in this regard is the smaller DC gain (compared to the integrator), which, as a result, allows for a DC servo loop to be implemented with smaller components. Besides, a simple one-time offset trim might also be sufficient now.

4.4. CONCLUSION ON THE COMBINING ARCHITECTURES

In the last two chapters, the combining of the two signal paths has been extensively discussed. In Chapter 3, several high-level calculations have been performed to be able to form general conclusions on the coil-Hall architecture without any assumptions about the implementation. Besides, the inherent coil performance was evaluated with a main emphasis on the advantages and disadvantages of either reading out the coil voltage or current. Chapter 4, on the other hand, went into some more depth on implementing the coil-Hall architecture. This chapter mainly evaluated how the high-frequency (coil) path could be implemented and combined with the low-frequency (Hall) path. This was done by evaluating the advantages and disadvantages of three different architecture "groups." This last section is meant to wrap up these two chapters by selecting a particular architecture for the prototype chip for SystematIC Design.

As mentioned in the introduction, SystematIC Design already has a fully tested, functioning design based on the multipath architecture. However, this design uses a Hall plate instead of a coil in the high-frequency path. The main goal of the new prototype chip is to replace the Hall plate in the high-frequency path with a coil and show the large benefits in terms of noise and bandwidth. Certainly, this means that the high-frequency path has to be completely redesigned; however, to minimize design time, the low-frequency path should ideally be kept unaltered. Besides, the new architecture should not cause a deterioration in all other performance aspects or cost factors (of which the requirements are given in Section 2.4).

First, the shared low-pass filter architecture can be considered. Compared to the other architectures, this architecture should be relatively area-efficient as it does not require multiple large components to implement an integrator (or low-pass filter) in addition to the summing filter (both with large time constants). However, as the high-frequency path is DC-coupled, its offset is directly added to the output. To mitigate this problem, additional circuitry and, hence, complexity (and area) are added to this seemingly simple architecture.

However, more importantly, this architecture mainly fails to be effective if no additional measures are taken to reduce the spinning ripple from the low-frequency path. Namely, if only the low-pass filter is to be used to reduce this ripple, the cross-over frequency must be set extremely low (around 235 Hz). This requires large components to implement (reducing this design's area efficiency) and results in a large noise penalty (due to a far-from-optimal cross-over frequency from a noise point of view). Reducing the ripple more efficiently would require (large) changes to the low-frequency path, which SystematIC does not desire.

One way to alleviate this tight ripple versus noise trade-off without altering the low-frequency path is to use a second-order low-pass filter. However, for this to result in an overall flat response, the high-frequency path must now contain two poles and a zero in its transfer. These must be accurately matched to the low-frequency path poles, which adds more complexity than only matching a single pole.

The second architecture, based on an integrator in the high-frequency path, allows for such a filter to be easily and accurately implemented. Indeed, it was shown that this way, area consumption and noise can be reduced while the ripple is still sufficiently attenuated. However, even with the second-order summing filter, reaching optimum noise performance is still believed to be difficult. Besides, implementing the integrator introduces issues related to clipping. Solving these issues adds significant design complexity.

The last architecture, the hybrid-hybrid architecture, solves many of these issues. Namely, using a coil-Hall hybrid sensor inside the high-frequency path allows for the noise performance to be designed completely orthogonal to the ripple performance. The ripple versus noise trade-off is effectively broken. Besides, implementing this additional coil-Hall sensor is relatively easy as offset and ripple do not play a role in the high-frequency path (i.e., the shared low-pass architecture can be implemented without any of its challenges). The main downside to this architecture is increased power consumption for biasing the additional Hall plates, and

	Shared LPF	Integrator	Hybrid-Hybrid
Area	-	-	0
Low Complexity	-	-	++
Calibration Complexity	0	0	-
Power Consumption	0	0	-
Noise	-	0	++
Gain Flatness	0	-	+
Offset	-	++	++

Table 4.1: Comparison table for the three combining architectures while considering the constraints introduced by SystematIC (i.e., no additional filtering of the ripple in the low-frequency path)

some potential added calibration complexity. However, as mentioned in Section 2.3, power consumption in this type of current sensor is believed to be of secondary importance. Increasing the power consumption in exchange for significantly reduced design complexity is certainly believed to be worth it. Given the constraints by SystematIC, the **hybrid-hybrid architecture** will therefore be chosen to be implemented.

5

GAIN CALIBRATION

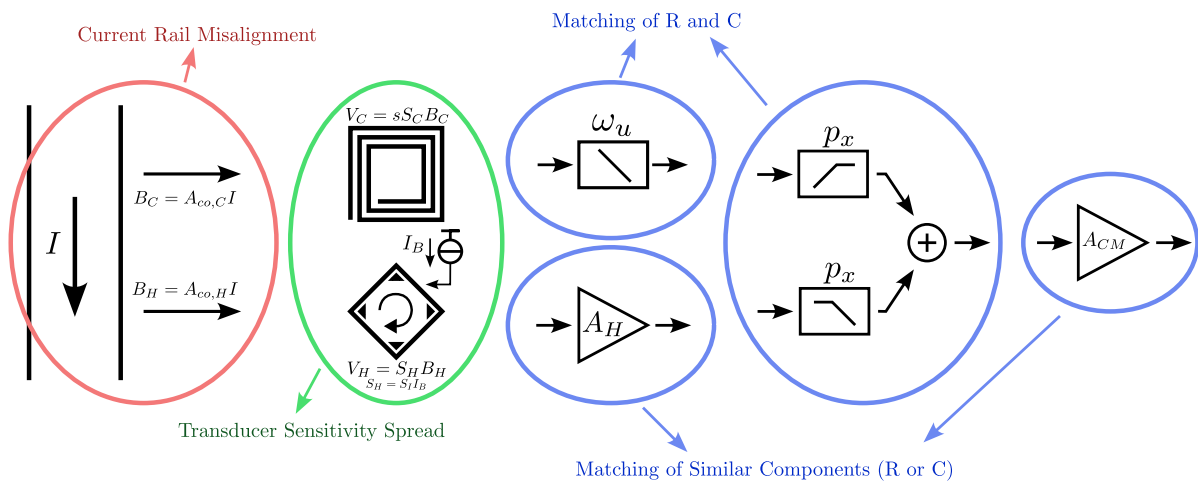


Figure 5.1: Representation of the overall system chain with an indication of potential error sources contributing to an imperfect system gain

Until now, it has been assumed that all transfer functions were perfectly accurate. While, for simplicity, this is fine to assume throughout the previous chapters, in an actual design, this does not naturally happen. The transducers themselves, components setting accurate gains, and magnetic field coupling factors, to name a few, are all prone to random sample-to-sample variations. Even worse, these variations can be highly temperature-dependent and even drift over lifetime. Therefore, measures must be taken to ensure a certain degree of gain accuracy.

These measures can be divided into two distinct groups. These will be, throughout this thesis, referred to as **gain calibration** and **gain stabilization**. Gain (foreground) calibration refers to the one-time step of removing any static absolute (sample-to-sample) gain errors after device fabrication, while gain stabilization refers to the process of ensuring that this calibrated gain does not drift under changing operating conditions. In other words, gain calibration ensures that each device outputs exactly the same absolute voltage for a certain current, while gain stabilization ensures that this still happens over a wide range of operating conditions.

As these two steps focus on quite different aspects of gain-inaccuracy, they will be discussed in two separate chapters. This first chapter will focus on gain calibration, while the next chapter will focus on gain stabilization. However, the goals of each chapter are similar. Namely, the goals are to highlight the various challenges associated with each step, present potential solutions, and indicate which solutions will be used for the prototype chip for SystematIC.

In this first chapter, firstly, an overview will be given of various potential sources of gain error in the coil-Hall hybrid architecture, and secondly, foreground calibration will be discussed in more detail.

5.1. SOURCES OF GAIN ERROR

Before it can be investigated how to fix and regulate the gain accurately, it is important first to have a picture of the different error sources contributing to a deviation from the ideal system gain. These sources can essentially be divided into the three groups graphically depicted in Figure 5.1. The three colors indicate errors related to the magnetic field coupling (red), the transducers (green), and the readout electronics (blue). This section will briefly consider these error sources in terms of their absolute sample-to-sample variability and behavior over temperature and lifetime.

MAGNETIC FIELD COUPLING FACTORS

The first group is related to the magnetic field coupling factor from the current rail's current to the magnetic field at either the coil ($A_{co,C}$) or Hall plate¹ ($A_{co,H}$). Namely, when packaging the chip, this current rail is placed at a certain position and with a certain orientation w.r.t. the die. The accuracy of this placement is limited. Misalignment of the current rail in the package used by SystematIC can, for example, cause an absolute coupling factor variation as large as 20% [31]. However, as the placement should barely change once the rail has been placed, this coupling factor is not believed to change much over temperature or lifetime.

TRANSDUCER SENSITIVITY

The second group is related to variations in the transducer sensitivities. For the coil-Hall architecture, complexity is increased as the variation of multiple transducers, based on completely different sensing mechanisms, must be accounted for in a single design.

The Hall plate sensitivity is proportional to the current- (or voltage) related sensitivity (S_I and S_V respectively) and the bias current (or voltage). The current- and voltage-related sensitivities are known to be very prone to process spread. This is because they highly depend on material parameters such as, for example, doping level and carrier mobility [5]. This also makes the sensitivity quite temperature-dependent (roughly 15% variation for S_I and 50% for S_V over the, -40 °C to +125 °C, temperature range of interest). Besides, the piezo-Hall effect introduces a cross-sensitivity to mechanical stress, which, due to stress variation in the package, can cause a lifetime drift as large as 2 - 6% [32–35].

The coil sensitivity (from magnetic field to output voltage), on the other hand, is purely related to the coil's geometry. As lithography can accurately set the dimensions, the sensitivity is not prone to large absolute, device-to-device variations. Besides, thermal or stress-related effects should not have a large influence on this geometry either. Therefore, the sensitivity is believed to remain highly stable over temperature and lifetime. However, if the coil current is used as information carrying quantity, this accurate voltage is transformed over a metal resistance, which can have an absolute variation as much as $\pm 30\%$. Besides, this resistance is also highly temperature-dependent, with a large temperature coefficient of roughly 3790 ppm/K. This results in a variation of roughly 60% over the temperature range of interest.

READOUT ELECTRONICS

Lastly, the third group is related to the readout electronics. If the amplifiers are designed properly, their temperature dependence and absolute variability should only be determined by the passive components constituting the amplifier's intended transfer. For a simple gain, flat over frequency, this is commonly set by a ratio of resistors or capacitors. Such a gain can be made very stable over temperature and lifetime. Besides, the absolute variation can be well below 1% with accurate component matching in the layout.

On the other hand, if a frequency-dependent transfer is required, this is commonly set by the product of a resistance and capacitance². As these are two completely different component types, the absolute variability is generally several tens of percent. Besides, stable operation over temperature and lifetime cannot be guaranteed as the components can behave completely differently. Such a gain error is relevant in the high-frequency path of a coil-Hall architecture due to the required integration of the coil signal.

¹As the magnetic field is (likely) not completely uniform, these coupling factors can be slightly different due to the different location and sizes of the Hall plate and coil

²In a discrete-time sampled system, it can, for example, be accurately set by a ratio of capacitors and the clock frequency. However, not all systems are suitable for such an approach

5.2. FOREGROUND CALIBRATION PROCEDURE

Because of the aforementioned error sources, the absolute sensitivity of the system is not guaranteed to be accurate when the chip has been fabricated and packaged. Therefore, the sensor has to be calibrated before it can be sold. This is referred to as (one-time) foreground calibration.

During foreground calibration, a well-known signal is applied to the system's input, the response is compared to an accurate reference and the gain is trimmed accordingly. To reach a certain degree of accuracy, a certain signal-to-noise ratio (SNR) is required. The SNR can either be increased by increasing the calibration signal or by decreasing the noise. As the system noise spectrum is already fixed by design, the only way to reduce the noise further is by limiting the detection bandwidth. As a result, the system slows down, and the foreground calibration sequence takes more time. Ideally, the calibration is performed in as little time as possible to reduce production costs. Therefore, using the largest possible calibration signal is desirable.

In Hall-only designs, fast foreground calibration is no problem. This is because a large DC calibration current of several tens of ampères can be easily and accurately generated. However, in the coil-Hall hybrid architecture, the story is different. In this case, an additional calibration step is required to calibrate the high-frequency path. This additional step in itself already increases the calibration time. However, to make matters even worse, the high-frequency calibration step cannot be performed with a DC signal because of the small coil SNR at low frequencies and a relatively high cross-over frequency between the two signal paths.

One disadvantage of the coil-Hall hybrid architecture
is the fact that an **additional** foreground calibration step is required.

POTENTIAL HIGH-FREQUENCY PATH CALIBRATION CHALLENGE

Especially in an automated calibration setup, generating a large (several tens of ampères) and accurate current becomes more difficult at higher frequencies. Namely, such a setup is designed to perform a wide variety of tests to verify the device's performance and is, therefore, not optimized for the calibration sequence only. One should not expect the performance of the automated setup to be as good as that obtained in a well-controlled and fine-tuned lab environment³ [31].

To put things into perspective, the automated calibration setup used for the calibration of the previous Hall-Hall design by SystematIC was only able to maximally generate a $10A_{pp}$, 1 kHz sinusoidal signal with sufficient accuracy⁴ (to get the gain within $\pm 1\%$ of the desired gain). However, to properly calibrate the system's high-frequency response, the calibration frequency should be well above the cross-over frequency of a few kilohertz. This requires a considerably higher calibration frequency.

Of course, more time and money can be invested in enhancing the calibration setup regarding accurate high-frequency, high-current signal generation. Certainly, no fundamental limits have been observed yet. However, for SystematIC, this is not an attractive and economically viable solution. Therefore, some other strategies that do not require high-frequency, high-current signals will be investigated in the next section.

³Think, for example, of the fact that relatively long wires are required to transport the signal from the calibration setup to the current sensor.

⁴This is certainly no fundamental limit, but it did require a custom-made, expensive piece of equipment to achieve.

5.3. FOREGROUND CALIBRATION STRATEGIES

There are believed to be four different strategies for calibrating the high-frequency path without the need for a high-frequency, high-current calibration signal. The first two strategies aim to decrease the required calibration current, while the last two aim to decrease the required calibration frequency. The latter two strategies are made possible as the system does not have to be fully operational yet during foreground calibration. This adds freedom to the calibration procedure. For each strategy, the main idea will be briefly explained. However, the exact implementation will be left open as this is highly architecture-dependent. Only some implementation examples will be given if they help clarify the main idea further. It should also be noted that the ideas are not mutually exclusive; they can be combined if that is believed to be beneficial.

REDUCED DETECTION BANDWIDTH

As mentioned before, the SNR can be increased by either decreasing the noise or by increasing the signal amplitude. In a design with a fixed noise spectrum, the noise can only be reduced further by introducing an additional bandwidth limitation, which, in turn, slows down the calibration response. This also means that, for a given SNR, the calibration signal amplitude can be reduced at the cost of calibration time.

There seems to be a trade-off between
calibration time and calibration accuracy

If this cost in calibration time is acceptable, a high-frequency but low-current calibration signal can be used. The main advantage of this idea is that the system's response can be truly calibrated at high frequencies without any changes in the signal path. This should be beneficial for calibration accuracy and obviates the need for additional on-chip circuitry. Such a calibration loop can, for example, be implemented using a highly selective bandpass filter, a lock-in amplifier (i.e., a synchronous detector followed by a low-pass filter), or by sampling and averaging over a long period of time. Also, note that this strategy can effectively be combined with any of the following strategies.

ON-CHIP COIL

Another way to decrease the required calibration current is by increasing the magnetic field strength generated for a certain amount of current. By only using the current rail, this cannot be achieved. However, for this purpose, an on-chip coil could be employed.

As the coil is considerably closer to the sensing elements (sensing coil and Hall plate), a smaller current is required for a certain magnetic field strength. If, for example, a sensing coil of $300 \times 300 \mu\text{m}$ is used, the calibration coil should also at least be $300 \times 300 \mu\text{m}$. It can be roughly calculated that⁵, with a single turn of such a coil, a magnetic coupling factor of roughly $4200 \mu\text{T A}^{-1}$ can be obtained. Of course, the smaller the calibration coil can be, the closer it is to the sensing elements and the stronger the coupling factor can be. However, even with this relatively big calibration coil, the coupling factor is already approximately an order of magnitude higher than the coupling factor from the main current rail ($\sim 400 \mu\text{T A}^{-1}$) and can even be increased further by using multiple turns.

While such a large coupling factor could help reduce the required calibration current without loss of SNR, this approach has several drawbacks. First of all, depending on the calibration coil size and trace width, the coil resistance can be relatively large. This limits the maximum possible calibration current through the calibration coil⁶, and, thus, limits the maximum calibration signal strength. Besides, by using a calibration coil, not the entire signal chain is calibrated. Namely, any errors in current rail misalignment (see Section 5.1) are not directly accounted for. To make sure this calibration strategy still gives adequate results, the (already calibrated) low-frequency path's response to a signal on the calibration coil should be used as a reference for the high-frequency path calibration. Acquiring this response requires an additional calibration step, which (slightly) increases calibration time. Besides, calibrating the high-frequency path indirectly will likely also result in a decrease in calibration accuracy⁷.

⁵The calculation procedure can be found in Appendix D

⁶e.g., a single turn of $300 \times 300 \mu\text{m}$ with a $10 \mu\text{m}$ trace width and $80 \text{ m}\Omega$ per square metal sheet resistance results in 9.6Ω . A 1 A calibration currents would require 9.6 V across the coil and causes 9.6 W of power dissipation

⁷For example, the relative relation between the four different coupling factors (from current rail or calibration coil to Hall plate or sensing coil) should not change from sample-to-sample for this solution to work

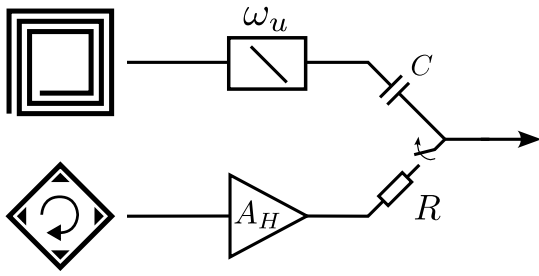


Figure 5.2: By disconnecting the resistance in the summing filter, the filter's cross-over frequency is lowered, and the high-frequency path can be calibrated with a lower frequency signal

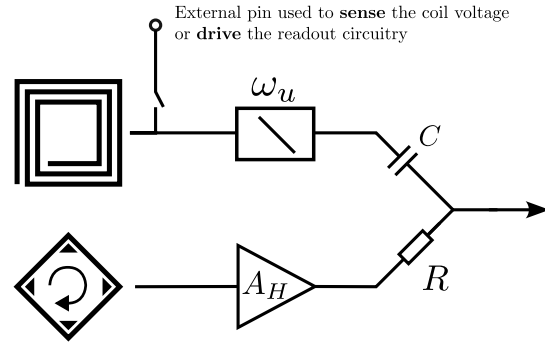


Figure 5.3: By bringing out an internal pin, the high-frequency path can be calibrated in multiple steps

TEMPORARILY CHANGE THE SIGNAL PATH

The solutions so far were aimed at allowing for a higher frequency but lower amplitude calibration signal. However, another approach would be to allow for large but low-frequency calibration signals (e.g., below 1kHz as currently used by SystematIC). During normal system operation, this would, of course, not be possible due to the system being designed for a relatively high cross-over frequency. However, during foreground calibration, the system does not necessarily have to be fully functional yet; the signal path can be temporarily changed. This can be exploited to allow for low-frequency calibration of the high-frequency path.

How this is exactly implemented heavily depends on the chosen architecture. For the integrator architecture from the previous chapter, it is relatively simple. Namely, as the integrator in the high-frequency path already flattens the coil response at low frequencies, only the summing filter cross-over frequency has to be lowered to pass low-frequency signals from the high-frequency path to the output⁸. This can, for example, be done by disconnecting the low-frequency path as shown in Figure 5.2. For the other architectures, it is probably more complex to implement as, besides lowering the cross-over frequency, an integrating transfer must be implemented to also flatten the coil response at lower frequencies.

The complexity of this solution, therefore, highly depends on the chosen system architecture. Besides, it must be carefully considered if the altered signal path still gives an accurate representation of the high-frequency behavior of the high-frequency path. Any errors in that regard decrease the obtainable accuracy of this approach.

MULTIPLE CALIBRATION STEPS

The last calibration strategy takes a slightly different approach. Instead of calibrating the entire signal chain in one go, different parts can be calibrated separately. This introduces more freedom in choosing the calibration signals. For example, the coil response can be measured at 1 kHz with a $10A_{pp}$ current on the current rail, while the subsequent readout's response can be measured by applying a more easily generated high-frequency (e.g., > 100kHz) voltage at the (high-impedance) input node of the high-frequency path. This strategy only requires some internal chip nodes to be made available to an external pin (see Figure 5.3). This can be done with minimal additions to the design.

The downside of this multiple-step approach is the potential increase in calibration time. This highly depends on how fast each intermediate step can be performed. Each step must be performed with a certain accuracy, and, as mentioned before, depending on the signal magnitude, this takes a certain amount of time. Measuring the coil response will likely take quite some time as the signals are still relatively small, even for a large current. However, measuring the readout response can probably be performed relatively quickly as relatively large voltages can be easily applied.

⁸One could also choose to directly measure the high-frequency path output before the summing filter, however, then a part of the signal chain is not considered during calibration

5.4. CONCLUSION ON FOREGROUND CALIBRATION & STRATEGY SELECTION

In this chapter, foreground calibration was discussed. It was concluded that one of the disadvantages of the coil-Hall hybrid architecture is the fact that an additional high-frequency path foreground calibration step is required. Moreover, this additional step requires a (relatively) high-frequency, high-current calibration signal. While generating such a signal is not fundamentally impossible, it would be easier to generate a lower frequency and/or lower current calibration signal. For this purpose, several high-frequency path calibration strategies were introduced.

The effectiveness of each approach is highly dependent on the specific system implementation. Therefore, no general conclusions can be drawn on which strategy is best to use. However, based on SystematIC's specific desires, a strategy can be selected for the hybrid-hybrid architecture.

Most importantly, SystematIC wants this new coil-Hall prototype to be based on an existing Hall-Hall chip with as few changes as possible. This also means that, ideally, the foreground calibration sequence is compliant with the existing calibration setup, making use of low-frequency (1kHz, 10A_{pp}) currents. This makes the first two strategies considerably less attractive.

The third strategy, based on temporarily changing the signal path, is relatively difficult to implement in the hybrid-hybrid architecture. The high-frequency Hall plate must be disconnected, and the low-pass filter (with cross-over frequency f_z , see Section 4.3) must be converted into an integrator. This also requires additional circuitry to (temporarily) make sure the signal path does not clip on its own offset. As SystematIC wants to have a relatively simple design, this is, in this case, not deemed the best calibration strategy.

As a result, mainly because of its simplicity, the **multiple calibration steps** strategy is selected. The accuracy and required calibration time are, at this point, unknown for this strategy. However, compared to the other strategies, no clear disadvantages are expected in that regard.

Lastly, it should be noted that the detailed implementation of the (external) automated calibration procedure will not be discussed further in this thesis. As agreed upon by SystematIC, this is future work for them or their potential customer(s). The only concern for this design is that the relevant internal connections are accessible at an external pin to make testing with such a calibration strategy possible.

6

GAIN STABILIZATION

After foreground calibration, the system should have a well-defined gain from input current to output voltage; all absolute static gain errors have been mitigated. However, during operation, external factors such as temperature, stress, and humidity, to name a few, could all cause this calibrated gain to drift. Therefore, to make sure the gain remains accurate over a wide range of operating conditions, additional measures have to be taken. The process of maintaining an accurate gain during system operation will, throughout this thesis, be referred to as **gain stabilization**. Its effectiveness will be expressed as a maximum gain variation over temperature and lifetime¹.

Various different parts of the signal chain contribute to gain instability. As was explained in Section 5.1, in the low-frequency path, the Hall plate sensitivity is the main culprit, while, in the high-frequency path, a resistance and capacitance (required to implement an integrating transfer) cause the most variation. Especially when reading out the coil current, the large temperature dependence of a metal resistance has to be dealt with.

Similar to the previous chapter, the goal of this chapter is to highlight various challenges associated with gain stabilization and present various solutions. Moreover, to keep a clear overview, it is believed that the numerous stabilization solutions can be categorized based on which of two general techniques they employ. These techniques will be referred to as **feedforward compensation** and **background calibration** and will be discussed in the following sections. Ultimately, a stabilization strategy will be chosen for the SystematIC prototype chip.

¹The lifetime drift specification implicitly captures the influence of all external factors unrelated to temperature

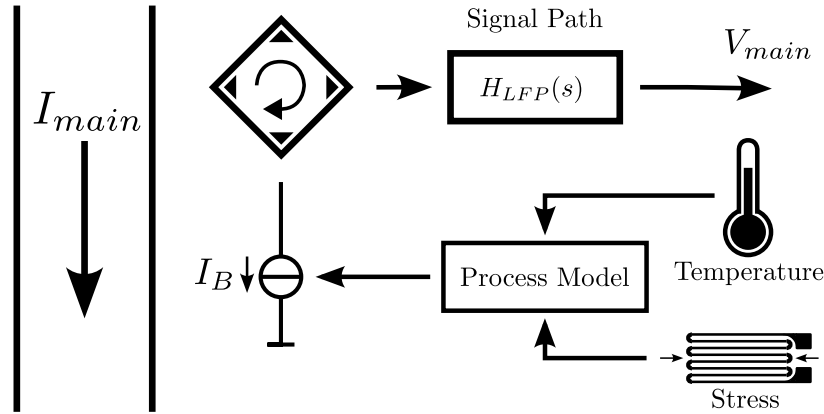


Figure 6.1: High-level representation of feedforward compensation applied to the low-frequency path

6.1. GAIN STABILIZATION: FEEDFORWARD COMPENSATION

First of all, feedforward compensation will be discussed. Feedforward compensation is a technique that uses knowledge of certain gain-disturbing processes to compensate for them. However, the gain itself is not actually measured. For the low-frequency Hall-based path, this technique is conceptually shown in Figure 6.1. As can be seen, relevant gain-disturbing external factors, such as temperature and mechanical stress, are measured and used via an underlying model to correct for any gain variation related to these disturbances. This gain correction can, in the low-frequency path, be easily implemented by controlling the Hall bias current (I_B), which sets the Hall plate sensitivity according to $S_H = S_I \cdot I_B$. Although shown for the low-frequency path, the main idea of measuring and modeling gain-disturbing factors certainly also holds for the high-frequency coil path.

The accuracy of the feedforward compensation technique mainly depends on how well the disturbing factors can be measured and how well their effect can be modeled. The model and disturbance sensor(s) should be well-behaved and should ideally not change over time or differ from sample to sample. However, this can not always be guaranteed, and, as could be seen in [35], potential (costly) foreground calibration is required to correct for such variations. Besides, with this technique, each disturbing factor must be explicitly measured and compensated for. Generally, the more disturbance factors are measured and modeled, the more accurate the gain stabilization becomes. Increased accuracy, therefore, often goes hand in hand with increased design complexity.

Implementations of the feedforward compensation technique can take on many forms. These implementations can vary widely in terms of complexity and accuracy. For the low-frequency Hall path and the high-frequency coil path, potential implementations of feedforward compensation will be discussed below.

6.1.1. FEEDFORWARD COMPENSATION IN THE LOW-FREQUENCY HALL PATH

As can be deduced from the various sources of gain error in Section 5.1, the magnetic field coupling factor (A_{co}), based on geometry, and the readout gain stages (with a frequency-independent gain), based on accurate resistor or capacitor ratios, should not contribute much to gain drift. On the other hand, the Hall plate sensitivity can drift significantly over temperature and lifetime. Therefore, gain stabilization in the low-frequency path mainly focuses on stabilizing the Hall plate.

In literature, various different implementations of the feedforward compensation technique can be found. For example, in [35], to stabilize the Hall sensitivity, both temperature- and stress sensors were employed in combination with an elaborate correction scheme. This correction scheme could be used to model the stress and temperature curve with multiple degrees of freedom. As a result, this implementation could reduce gain variation over temperature and lifetime to below 1%; however, a costly three-point temperature trim would be required to realize this.

On the other hand, in [13], feedforward temperature compensation was implemented by matching the tem-

perature coefficient (TC) of various resistors to the TC of the Hall plate current-related sensitivity. In essence, the temperature behavior of the resistors was used as a model for the temperature behavior of the Hall plate current-related sensitivity. But, as expected, the simplicity of this implementation's model resulted in a relatively large residual gain variation of roughly $\pm 2.5\%$ from -40°C to $+80^\circ\text{C}$.

SystematIC has also already implemented a feedforward compensation-based solution. This solution is very similar to the solution in [13] and is also based on matching the temperature coefficients of various resistors to the temperature behavior of the current-related sensitivity. The implementation of SystematIC even tries to compensate for a second-order component in the Hall plate temperature behavior. However, still, this solution only stabilizes the gain within $\pm 3\%$ of the ideal gain (in a temperature range from -40°C to $+125^\circ\text{C}$). Besides, lifetime drift is not accounted for as the required disturbance factors (mainly stress) are not measured and included in the model. As noted in Section 5.1, this lifetime error drift can be as large as $\pm 3\%$.

6.1.2. FEEDFORWARD COMPENSATION IN THE HIGH-FREQUENCY COIL PATH

As opposed to the low-frequency path, in the high-frequency path the transducer sensitivity is considerably more stable. Namely, the coil sensitivity (magnetic field to voltage) is purely based on geometry which is expected to remain very constant over temperature and lifetime. As a result, all residual gain drift is now caused by the various readout stages, and, thus, with proper design, any temperature and lifetime drift is caused by variations of on-chip resistances and capacitances. As the components' temperature dependencies can be quite efficiently "modeled" and compensated with other on-chip components of the same type (e.g., think of ratios of resistors), this makes feedforward compensation quite effective in the high-frequency path.

The largest temperature dependence is introduced when reading out the coil current. This is caused by a large temperature dependence of the metal resistance (with a temperature coefficient of 3790 ppmK^{-1}), which transforms the stable coil voltage to a relatively unstable current. Although it adds complexity compared to voltage readout, it can still quite accurately be compensated by matching this resistance with another metal resistance somewhere else in the signal chain. This is, for example, effectively implemented in [11] with an interesting pole-zero cancellation scheme.

There is one more catch, however. Namely, independent of voltage or current readout, to flatten the differentiating coil response, an integrating transfer would be required. This integrating transfer (either implemented by an integrator or a low-pass filter, see Chapter 4) can, in a continuous time system, only be implemented by combining a resistance and capacitance. As a result, the high-frequency path unavoidably contains a product of resistance and capacitance in its gain. These two completely different component types do not necessarily cancel each other's temperature and lifetime drift (as opposed to, for example, a ratio of two resistors or capacitors of the same type).

Fortunately, in terms of temperature drift, this effect can be minimized by using a p-poly resistor in combination with a MIM (metal-insulator-metal) capacitor. Namely, the first-order temperature coefficient of -27.23 ppmK^{-1} of the resistor and $+30.44 \text{ ppmK}^{-1}$ of the capacitor almost completely cancel. As a result, due to a remaining second-order component in the resistor's temperature dependence, only a $\pm 0.6\%$ gain variation over the temperature range from -40°C to $+150^\circ\text{C}$ remains. However, this does not necessarily mean these components also cancel each other's lifetime drift. Unfortunately, whether this drift is significant is unknown and should be tested after device fabrication. Still, the high-frequency path can be stabilized considerably more accurately and efficiently by feedforward compensation than the low-frequency Hall path.

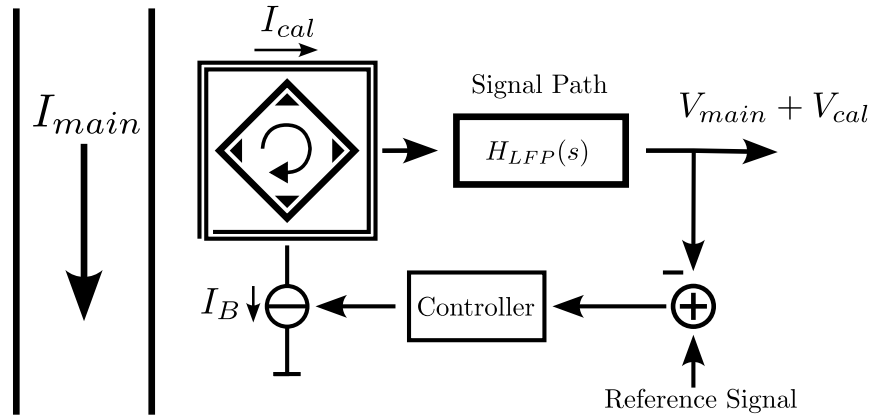


Figure 6.2: High-level representation of background calibration applied to the low-frequency path

6.2. GAIN STABILIZATION: BACKGROUND CALIBRATION

Although measuring the gain-disturbing external factors can provide much information about the gain variation, the gain should be measured directly to know the variation most accurately. As conceptually shown in Figure 6.2, the idea behind background calibration is to measure the gain directly, compare it with an accurate reference, and correct the gain accordingly. As a result, the stabilization accuracy no longer depends on how well each disturbance can be measured and how well its effects can be modeled. However, now, it is purely dependent on the stability of the calibration signals and the accuracy of the gain measurement. All possible disturbing effects are essentially taken into consideration without having to be measured explicitly. This, especially for complex underlying gain drift phenomena, could make this approach considerably more accurate than feedforward compensation.

The observing reader would have noted that this idea of measuring the gain, comparing it to a reference, and compensating it accordingly is very similar to foreground calibration. However, there is a major difference. Namely, gain stabilization by means of background calibration is performed **while the system is operational**. Therefore, the calibration must run in the background while not interfering with normal system operation (hence the name background calibration). In other words, some form of orthogonality must be introduced between the calibration loop and the main signal path. This makes implementing this technique quite challenging.

In the following sections, the main challenges associated with background calibration will be discussed, and several solutions will be provided. The main goal is to provide a clear overview of what it takes to implement this calibration technique successfully, but not necessarily what the exact implementation would look like. Besides, background calibration will mainly be considered with the low-frequency Hall path in mind. This is because it is believed that the Hall sensitivity can benefit most from a more accurate calibration scheme². Still, the ideas presented will also hold for the high-frequency coil path.

²Ideally, SystematIC wants temperature and lifetime drift to be reduced from $\pm 3\%$ to $\pm 1\%$

6.2.1. SIGNAL ORTHOGONALITY

First of all, as shown in Figure 6.2, by directly applying a calibration signal to the main signal path, the calibration signal and the input signal will interfere with each other. The calibration signal will, similar to spinning ripple, reduce the system's resolution, while the main input signals can cause calibration errors. Therefore, it is important that some form of orthogonality between the main signal path and the calibration loop is implemented somehow. This orthogonality is believed to be attainable in four different domains, with some being more effective than others. These include the frequency-, spatial-, temporal-, and common-mode differential-mode domains and will be discussed below.

FREQUENCY ORTHOGONALITY

First of all, one could consider the frequency domain to find some form of orthogonality between the input- and calibration signals. As can be seen in Figure 6.5, a calibration signal with its frequency content well above the summing filter cross-over frequency f_x could be used. This would indeed greatly help to reduce the signal deteriorating effect of the calibration signal on the main signal path.

However, this form of orthogonality does not help reduce the effect of interference from the current rail to the calibration loop. Besides, care must be taken so that the generated calibration magnetic field does not couple to the high-frequency coil path. This requires the sensing coil and Hall plate to be physically separated, which, for a (desirably) large sensing coil, can be difficult. Only if a calibration signal well outside the system bandwidth is used (and can still be accurately sensed), this solution might be more effective. However, even then, interference from the current rail could still pose an issue. This form of orthogonality on its own is, therefore, not believed to be very effective.

SPATIAL ORTHOGONALITY

The next form of orthogonality can be sought after in the spatial domain. Namely, one could place an additional Hall plate somewhere else on the chip and use this Hall plate for calibration. This Hall plate is essentially used as a "model" of the main Hall plate. It should be placed at a location where the magnetic field from the current rail is weakest and where the calibration signals do not couple to the main signal path. This way, the interference problem is greatly reduced, and the calibration loop can be significantly simplified. This is graphically depicted in Figure 6.6.

The major downside of this approach is that the Hall plate (and the subsequent readout) used in the signal path is not directly calibrated. The accuracy of this method is, therefore, highly dependent on how well the "model" Hall plate represents the temperature and lifetime drift of the main Hall plate. This concept can, therefore, also be seen as a form of feedforward compensation.

To make sure this model Hall plate behaves as similarly as possible to the main Hall plate, the two Hall plates should be located close to each other. This ensures their environment is as similar as possible such that the Hall plates experience similar temperature and stress levels³. However, this conflicts with the placement for optimum interference rejection. As no data could be found in the literature, it is difficult to tell beforehand how strict this trade-off between interference and accuracy is. However, it is believed to be worth investigating further as this solution could provide excellent orthogonality.

TEMPORAL ORTHOGONALITY

Another form of orthogonality can be found in the time domain. Namely, during one instance, the signal path can be calibrated, while in another, it is used to process information from the current rail. This, however, does obstruct the continuous time operation of the system. Therefore, this approach is best used in a discrete-time system.

For continuous-time systems, a better option would be combining temporal and spatial orthogonality, as shown in Figure 6.7. Namely, in this case, two Hall plates can be used. While one is used inside the signal path, the other one is calibrated; the Hall plates are switched periodically. In this case, the signal path always has access to a Hall plate, and thus, continuous time operation is not disrupted⁴.

Unfortunately, as opposed to the previous solution, the Hall plates in this solution are both used inside the signal path and, thus, are not shielded from any interference from the current rail. Besides, again, care must

³Varying package stress is believed to be the main cause of sensitivity drift [34, 36] and with on-chip stress differences as large as 80MPa [37] the Hall plate sensitivities could become as large as a few percent [29]

⁴In amplifier design, this technique is often referred to as the "ping-pong" technique

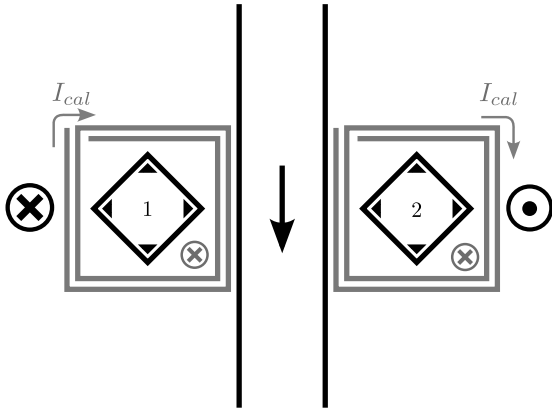


Figure 6.3: Differential signal sensing with common-mode calibration

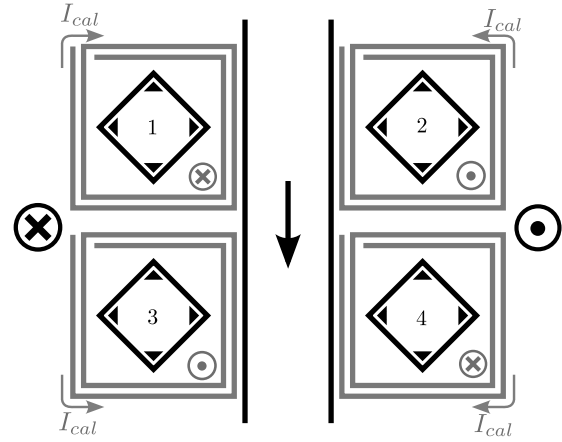


Figure 6.4: Differential signal sensing with another differential, but orthogonal calibration field.

be taken so that no calibration signal couples to the signal path. Therefore, this solution does not quite provide the desired two-way interference rejection and might only be useful when interference in the calibration loop is no problem.

COMMON-MODE DIFFERENTIAL-MODE ORTHOGONALITY

Lastly, one could consider common-mode differential-mode orthogonality. Namely, as shown in Figure 6.3, if the main signal (on the current rail) is sensed differentially (the magnetic field on either side of the current rail points in opposite directions), a common-mode field (pointing in the same direction on either side) can be applied for calibration⁵. The calibration signal can now be extracted while rejecting the main input signals and vice versa. This technique was, for example, successfully applied in [10]. However, with this approach, the main benefit of differential sensing, robustness against common-mode interferers, is lost.

To mitigate this weakness, one could aim for the approach shown in Figure 6.4. With this approach, four sensing elements are used. As opposed to the previous solution, these additional elements allow for more than two options for orthogonal signal extraction. Among these options, both the input and calibration signals can be applied differentially. As a result, robustness against common-mode interference can be maintained. This can, for example, be done with the field distribution as indicated in Figure 6.4⁶.

Although this solution seems to present the perfect two-way signal rejection, in a practical implementation, this is not exactly the case. Namely, the rejection of the different signals is highly dependent on the matching of the sensitivity of the different Hall plates (and the gain of subsequent amplifiers). For example, for a realistic Hall plate mismatch of 1%, the main input signal only gets rejected with a 40dB signal reduction from the calibration loop and vice versa. Whether this is sufficient is, of course, highly dependent on the system specifications. Still, this solution is quite promising as it could (as opposed to the other solutions⁷), to a large degree, provide the desired two-way orthogonality.

⁵In other words, the calibration field is extracted by summing the voltages ($V_{cal} = V_1 + V_2$), while the main signal is extracted by subtracting the two ($V_{main} = V_1 - V_2$).

⁶With the fields in Figure 6.4, the calibration signal can be extracted by $V_{cal} = V_1 - V_2 - V_3 + V_4$, while the main signal can be extracted by $V_{main} = V_1 - V_2 + V_3 - V_4$.

⁷Spatial orthogonality could also provide this functionality; however, in that case the signal path is not directly calibrated

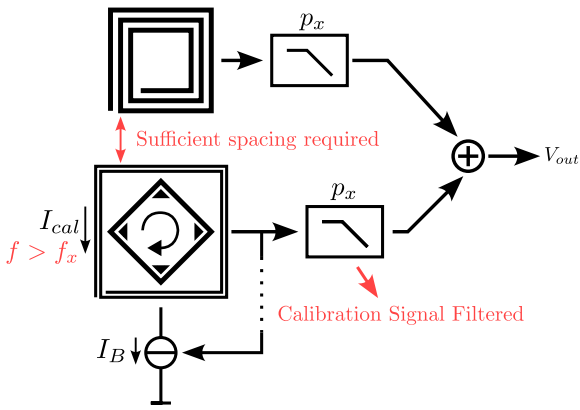


Figure 6.5: Graphical representation of using frequency orthogonality to reduce the effects of calibration signals on the main signal path

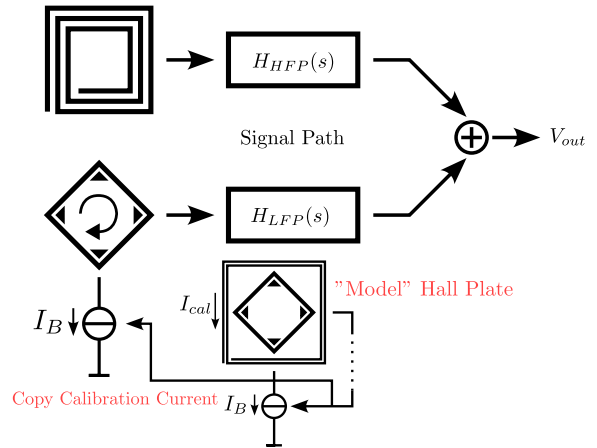


Figure 6.6: Graphical representation of spatial orthogonality making use of a "model" Hall plate placed in a different location than the main Hall plate

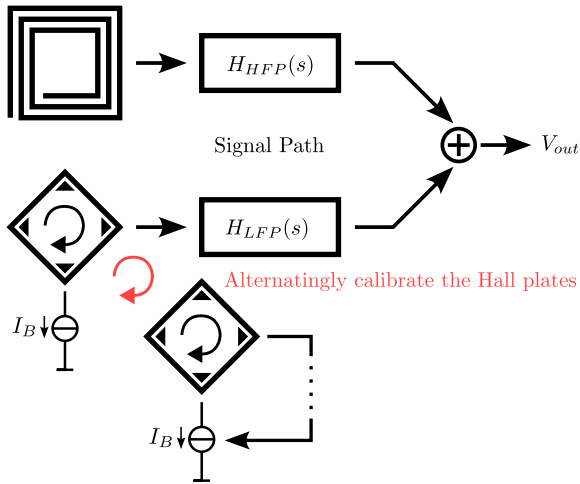


Figure 6.7: Graphical representation of using the combination of temporal and spatial orthogonality in which one instance of time the Hall plate is calibrated, and in another, it is used in the main signal path. By using two Hall plates alternatingly, continuous time operation should not be obstructed. This technique is also known as "ping-pong"

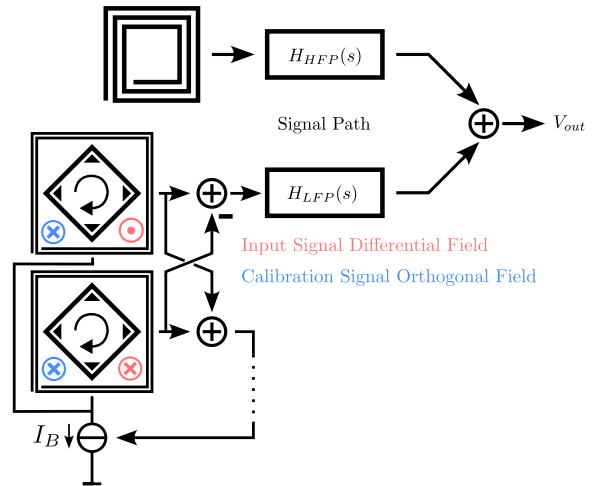


Figure 6.8: Graphical representation of orthogonality between differential- and common-mode magnetic fields (or other forms of orthogonal signal combinations)

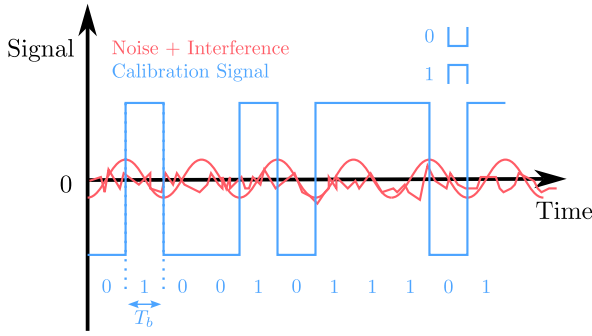


Figure 6.9: Calibration signal based on a random bit stream with bipolar nonreturn-to-zero signaling to implement a robust spread spectrum calibration sequence

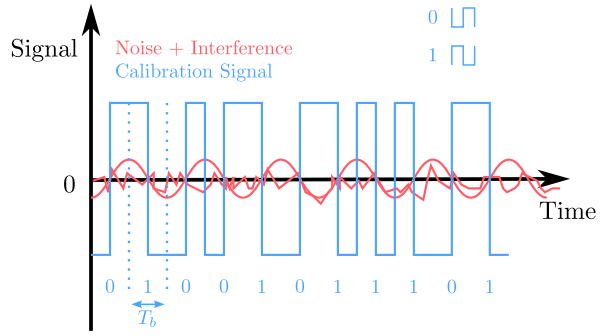


Figure 6.10: Calibration signal based on a random bit stream with manchester signaling to implement a robust spread spectrum calibration sequence. Most signal power is at

6.2.2. GAIN MEASUREMENT ACCURACY

Next, calibration accuracy should be considered. Namely, how accurately the gain can be regulated depends highly on how accurately the calibration signal can be measured at the output of the to-be-calibrated signal chain. As in any system, this accuracy is limited by noise, offset, and interference. The calibration signal should be considerably larger than these disturbances for a well-calibrated and stable gain. However, with the limited calibration signal strength of roughly 1mT ⁸ (see Appendix D), it will become clear that this is not necessarily straightforward, especially when dealing with interference. This section will investigate this further.

The most straightforward disturbances to deal with are offset and noise. Namely, offset (and low-frequency $1/f$ noise) can be mitigated by choosing a non-DC calibration signal, while noise can be reduced by limiting the detection bandwidth. For example, if the low-frequency Hall path should be calibrated within $\pm 1\%$ of the ideal gain, an SNR of at least 40dB would be required. To achieve this, with a typical Hall plate resistance R_H of $4\text{k}\Omega$, a sensitivity S_H of 100mV T^{-1} and a calibration signal B_{cal} of 1mT_{RMS} , the bandwidth should be limited to 15kHz (this value can be obtained by only considering white noise from the Hall plate and solving Equation 6.1). Such a cut-off frequency can be implemented with reasonably sized components (or can be implemented in the digital domain) and should certainly not pose any issues.

$$\frac{4kTR_H}{S_H^2} \cdot BW = B_{cal}^2 \cdot 10^{-\frac{SNR}{10}} \quad (6.1)$$

Interference, on the other hand, is more difficult to deal with. Take, for example, a calibration signal with frequency f_{cal} . Then, any signal on the current rail (or any other magnetic field) with significant signal power at f_{cal} will cause a calibration error. This error cannot be reduced by filtering, as the interference power could be heavily concentrated at the calibration frequency (as is the case with any periodic signal with frequency f_{cal}). To ensure a gain inaccuracy of $\pm 1\%$ with a calibration signal of 1mT , the signals on the current rail may not generate a signal larger than $10\mu\text{T}$ at that particular frequency. With a magnetic field coupling factor of $400\mu\text{T A}^{-1}$ (from the current rail to the magnetic field at the Hall plate) and a full-scale signal range of 80A , the interfering signals can be as large as 32mT . Even with the 40dB rejection when applying the common-mode differential-mode strategy (resulting in an equivalent interference of maximally $320\mu\text{T}$), this is still a lot, and, therefore, if no additional measures are taken, sufficient accuracy cannot be ensured for all possible signals on the current rail.

The **main challenge** with implementing background calibration
is properly dealing with **interference**

It can be concluded that, as opposed to noise and offset, the main challenge with background calibration is dealing with (unknown) interference from the current rail⁹ (i.e., the input signal). Of course, the calibration

⁸Note that this magnetic field requires a few milli-ampère to be generated, this is a small power penalty related to background calibration

⁹Or from other magnetic field sources

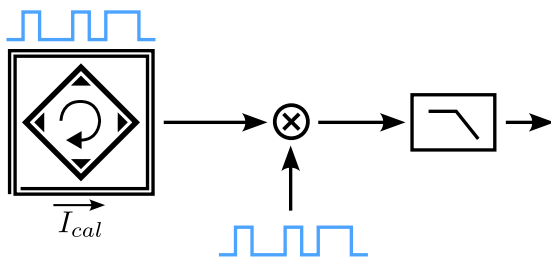


Figure 6.11: Extraction of the calibration signal from the random sequence calibration signal by means of modulation with the calibration signal sequence

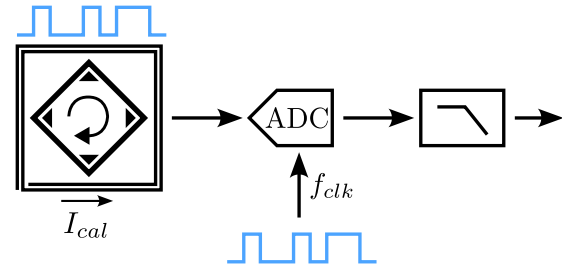


Figure 6.12: Extraction of the calibration signal from the random sequence calibration signal by means of sampling with a clock synchronised to the calibration signal

loop works perfectly fine for most input signals. However, when the input signal contains significant power at the calibration frequency, the calibration loop could easily malfunction. One could accept this risk, but ideally, the system is more robust. Some concepts from the previous section provide interference rejection to a certain degree. However, to completely mitigate the interference problem, even more measures need to be taken. Luckily, such a solution can be found by looking at more elaborate calibration signals.

6.2.3. INTERFERENCE REJECTION: SPREAD SPECTRUM SIGNALING

The only way to ensure robustness against all possible interfering signals would be by using a calibration signal that does not concentrate its power at only one (or a few) frequencies. Luckily, such signals do exist and are called spread spectrum signals. These signals are widely used in communication systems and are known for their interference rejection capabilities [38].

Spread spectrum signals are generated by multiplying the data with a so-called spreading signal¹⁰. This spreading signal is a pseudo-noise signal that causes the signal power (of a narrow-band signal) to be spread out over a wide bandwidth. To retrieve the signal, the received signal must be multiplied by exactly the same spreading signal. In that case, all received interference power is spread out over a wide bandwidth (just like the message was originally spread out) while the signal is restored to its original narrow-band form. Now, the interference, just like with white noise, can be filtered to retrieve the original signal with high fidelity.

SPREAD SPECTRUM SIGNAL GENERATION FOR BACKGROUND CALIBRATION

In communication systems, a pseudo-random code must be used so that this exact code can be reproduced at the receiver for proper despreading. However, for the background calibration, as the "transmitter" and "receiver" are close to each other (i.e., on the same chip), a fully random code can be used. This simplifies the signal generation significantly. Now, for example, the spreading signal can be a purely random bit stream that can be easily generated using a noise generator and comparator.

BIPOLAR SIGNALING

This bit-stream can, in turn, be converted into a calibration signal using bipolar signaling: a '1' represents a positive signal level, and a '0' represents a negative signal level (see Figure 6.9). By multiplying this signal with itself, it perfectly gets converted to a DC signal while all interference is spread out over a wide frequency range. Such a signal can easily be generated using a clock and a set of switches to reverse signal polarity, while it can be extracted by means of demodulation (using a simple chopper synchronized to the calibration signal) or sampling (using an ADC with a sampler synchronized to the calibration signal) as conceptually shown in Figures 6.11 and 6.12 respectively.

To understand what happens to the interference after multiplication with the random bit sequence, one should look at the power spectral density of this sequence. Namely, as can be seen from Figure 6.13 (blue curve), for bipolar signaling, the power is roughly spread out from DC to the bit-rate frequency ($f_b = \frac{1}{T_b}$, with T_b as defined in Figure 6.9). After mixing with an interfering signal, the resulting spectrum is the convolution of the bitstream spectrum and the interference spectrum. As a result, the interference power will be spread out in a bandwidth of f_b around the original interference frequency¹¹. It can thus also be concluded that the higher the bit rate, the more the interference power is spread out.

¹⁰There are actually multiple spread spectrum techniques; this particular technique is called direct sequence spread spectrum

¹¹For example, a sinusoidal interferer with a frequency equal to the bit rate would result in the spectrum shown in Figure 6.14. This spectrum, generated with a bit-rate of 1MHz, clearly shows the signal power spread out around 1MHz (blue curve)

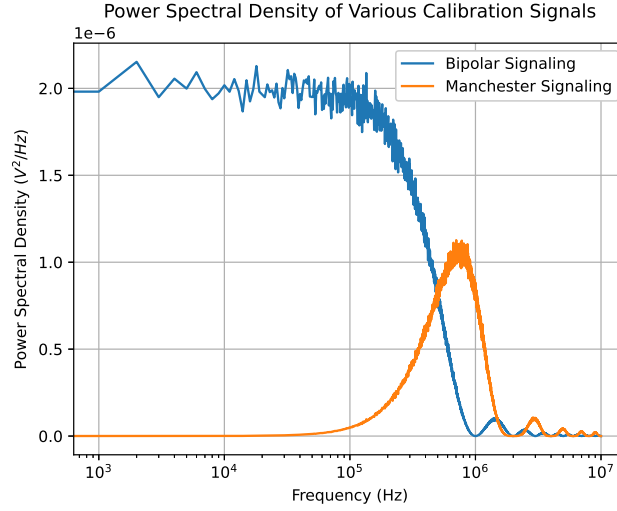


Figure 6.13: Power spectral density of a random bit stream of two different binary signaling formats. The two signals contain the same power, but, as can be seen, Bipolar signaling contains most power at low frequencies while Manchester signaling contains most power around the bit-rate frequency ($f_b = \frac{1}{T_b} = 1\text{MHz}$ in this case). The plots were obtained by averaging multiple FFTs of several randomly generated bit streams (the results are in line with the theoretical spectra from [38]).

MANCHESTER SIGNALING: ROBUSTNESS AGAINST LOW-FREQUENCY INTERFERENCE

For bipolar signaling, the interference power is spread out in a band directly around the original interference frequency. For example, low-frequency interference will be spread out at low frequencies, while high-frequency interference will be spread out at high frequencies. On the other hand, by using a different signaling technique, such as Manchester signaling [38], the interference power can, besides being spread out, also be modulated to a different frequency. Manchester signaling can be implemented by representing each bit with a switched pulse instead of a fixed signal level (see Figure 6.10). This results in the orange spectrum shown in Figure 6.13.

This signaling technique could be helpful if, for example, more robustness against low-frequency interference is desirable. However, this comes at the cost of worse robustness against high-frequency interference. This can, for example, be seen from the resulting spectrum for high-frequency interference shown in Figure 6.14. This plot shows the effect of the power spreading when the calibration signal is multiplied with a sinusoidal interferer at $f_b = 1\text{MHz}$. The interference power (orange spectrum) is clearly modulated down to low frequencies when Manchester signaling is used, but stays at high frequencies when bipolar signaling is used (blue curve).

DETECTION BANDWIDTH REQUIREMENT

Using either bipolar or Manchester signaling thus provides robustness to either high- or low-frequency interference, respectively. However, in each case, certain frequencies cause the interference power to be spread out from DC to the bit-rate frequency (f_b). For a robust system, the system must be designed to handle such worst-case signals. This essentially sets an upper limit to the detection bandwidth.

For example, the power of a full-scale 32 mT_{RMS} low-frequency sinusoidal signal would be spread out from DC to f_b if bipolar signaling would be used. For $f_b = 1\text{MHz}$, this results in a spectral density of roughly $\frac{0.032}{\sqrt{10^6}} = 32\text{ }\mu\text{T}/\sqrt{\text{Hz}}$. For an SNR of 40dB and a calibration signal of 1 mT_{RMS} , this results in an extremely small required detection bandwidth of less than 0.1Hz. However, if this type of signaling is used in combination with, for example, the common-mode differential-mode orthogonality strategy, the bandwidth requirement is relaxed to a more doable 1kHz. Therefore, these two techniques complement each other well.

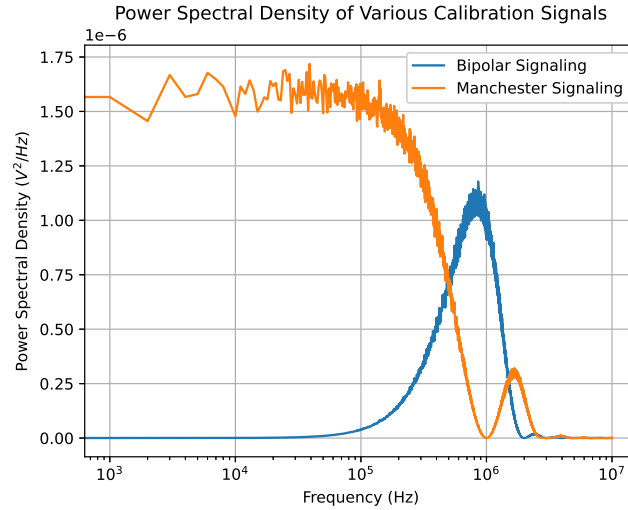


Figure 6.14: Same random bit sequences from Figure 6.14 multiplied with a 1MHz interfering signal with 1W of signal power. As can be seen, the power content from bipolar signaling shifts to a high frequency, while for manchester signaling, it shifts to DC

6.3. CONCLUSION ON GAIN STABILIZATION & STRATEGY SELECTION

In this chapter, two different gain stabilization strategies were investigated. These strategies were referred to as feedforward compensation and background calibration. With feedforward compensation, gain-disturbing processes are measured, and underlying models are used to correct for them, while, with background calibration, the gain itself is directly measured. Generally, background calibration is believed to be more accurate; however, the implementation, especially when dealing with interference, can be relatively complex.

Luckily, for the high-frequency coil path, not much effort is required to stabilize its gain. Namely, especially when voltage is read out, the only temperature dependence comes from the resistance and capacitance required to implement an integrating transfer to flatten the coil response over frequency. When this transfer is implemented with a p-poly resistor and a MIM capacitor, an overall gain variation of maximally $\pm 0.6\%$ can be assured over the temperature range from -40°C to $+125^{\circ}\text{C}$. This is more than sufficient for the to-be-designed prototype chip. Lifetime variation, on the other hand, is difficult to predict and should be measured afterward.

It is a different story for the low-frequency Hall path. SystematIC had previously implemented a feedforward-based compensation scheme that tries to match the temperature behavior of several resistor types to the temperature behavior of the Hall plate current-related sensitivity. However, these resistors cannot model the Hall plate behavior quite as accurately, and thus, roughly a $\pm 3\%$ variation over temperature remains [31]. Besides, lifetime drift was also not compensated for.

Ideally, SystematIC would like to keep this gain drift within $\pm 1\%$. For this purpose, background calibration was extensively investigated. As can be concluded from this investigation, implementing a robust calibration scheme can be quite complex, especially when dealing with interference. However, combining so-called common-mode differential-mode orthogonality with a spread spectrum calibration signal is believed to be a promising solution. Besides, the accuracy of the spatial orthogonality strategy is interesting to investigate further, as it could provide a less complex calibration alternative.

Unfortunately, for this first prototype chip, SystematIC had decided that neither of these two calibration schemes would be implemented. Both add too much complexity and, thus, design time to the project. To reduce research and development costs, the main focus of this first chip slightly shifted halfway through the project to only focus on the high-frequency coil path implementation. Therefore, the **existing feedforward compensation scheme** will be reused for this design. However, for a future design, SystematIC has certainly expressed an interest in implementing background calibration. Although this chapter did not discuss any in-depth implementation, the important challenges and solutions related to this technique are believed to be exposed. This should greatly help SystematIC for future development.

SYSTEM IMPLEMENTATION & VERIFICATION

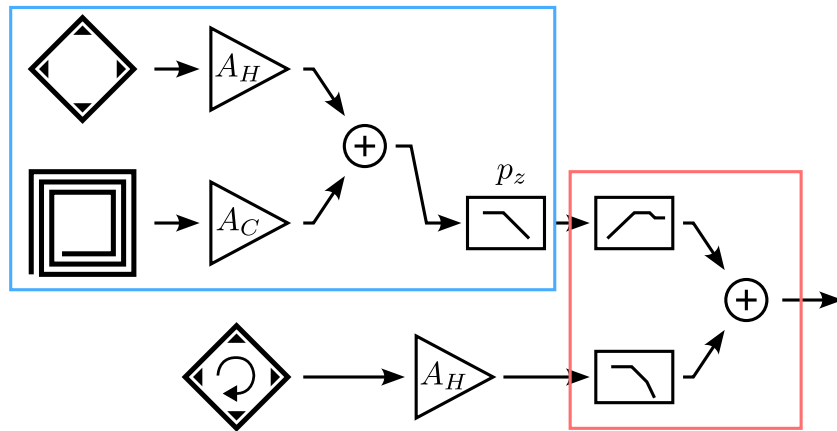


Figure 7.1: "Hybrid-Hybrid" architecture with an indication of the various implemented blocks. The blue block indicates the high-frequency path, and the red block is the summing filter. This thesis is not concerned with the already implemented low-frequency path. However, for the sake of completeness, a full system overview can be found at the end of this chapter.

In the previous chapters, extensive research has been done on many different aspects of the coil-Hall hybrid sensor architecture. In Chapter 4, various high-frequency path architectures have been investigated, while in Chapters 5 and 6, the calibration and stabilization of the low- and high-frequency paths have been discussed. In the end, it was decided that the so-called "hybrid-hybrid" architecture would be implemented. This architecture breaks the tight trade-off between ripple and noise performance without requiring alterations to the low-frequency path and, maybe even more importantly, allows for a relatively simple implementation of the high-frequency path. As to the calibration and stabilization, it was decided that to reduce design time, no additional circuitry would be implemented; for stabilizing the Hall plates, a previous solution by SystematIC will be re-used, while for the coil, reading out the coil voltage in combination with suitable resistor and capacitor types should provide a sufficiently stable response.

As the low-frequency path will be reused from a previous design¹, this chapter will only discuss the circuit implementation of the hybrid-hybrid architecture high-frequency path and the summing filter (i.e., the outlined blocks of Figure 7.1). This implementation will be part of a prototype chip for SystematIC, which aims to show the large noise and bandwidth benefits associated with the coil-Hall architecture. Ideally, the noise power should be reduced by a factor of 5x (from 180mA_{RMS} to less than 80mA_{RMS}) and the bandwidth should be increased by a factor of 6x (from 320kHz to 2MHz) compared to SystematIC's previous Hall-Hall design. All other specifications should remain the same and are specified in Section 2.4.

¹This low-frequency path (designed by SystematIC) consists of some transimpedance and transconductance amplifiers, a simplified overview is given in Figure 7.16

	Value
Hall Sensitivity (S_H)	80 mV T ⁻¹
Hall Plate Resistance (R_H)	700 Ω
Coil Sensitivity (S_C)	2.5 μ V T ⁻¹ rad ⁻¹ s
Coil Resistance (R_C)	1.4 k Ω
Coil Capacitance (C_C)	3.3 pF
Coil Inductance (L_C)	0.18 μ H

Table 7.1: Transducer model parameters (typical values, they are prone to process spread)

Moreover, to implement this design, a team of multiple people at SystematIC have worked on changing various circuit blocks from the previous Hall-Hall design (not only the high-frequency path). For the implementation (circuit design and layout), a total time of approximately eight weeks was allocated. Therefore, to save design time and costs, SystematIC highly valued simplicity in the design. Already existing circuit blocks should be re-used when possible, and ideally, as few possible new circuits should be designed. As long as, of course, the requirements can be met. Besides, implementing the design in a simple way is the perfect way to highlight the low complexity associated with the hybrid-hybrid architecture, one of its main strengths.

This chapter commences with a brief overview of the chosen transducer parameters. Next, an overview of the chosen high-frequency path readout- and summing filter architecture and the associated design choices. Next, specifications for the various amplifiers will be given, and the transistor-level design of these amplifiers will be considered. Lastly, this chapter concludes with the verification of the specifications by means of simulations, and a conclusion will be drawn, including some general recommendations for future designs. The testing itself will not be part of this thesis as the chip will not be manufactured in time. However, a test plan will be provided in Appendix A. Besides presenting and verifying the prototype chip, this chapter also aims to highlight various challenges related to implementation that are difficult to show in the previous chapters. Moreover, the true strength of the simplicity of the hybrid-hybrid architecture will become even more evident.

7.1. TRANSDUCER PARAMETERS

Before discussing the high-frequency path implementation, it is good first to present various design decisions related to the transducers:

- The sensing coil is designed according to the procedure presented in Appendix C. As the voltage is read out, the sensitivity is chosen to be limited to 2.5 μ V T⁻¹ rad⁻¹ s. This ensures that a full-scale (± 80 A), maximum frequency (2MHz) signal in combination with the largest possible magnetic field coupling factor of 400 μ T A⁻¹, results in a maximum coil signal V_{max} of ± 1 V. By differentially reading out the coil, with a worst-case supply voltage of 3.3V, processing this signal should not pose any issues. This coil is split into two halves to implement differential sensing (see Section 6.2.1). These coils are placed on both sides of the current rail and have an outer dimension of roughly 300x600 μ m, 8 turns, and a wire width of 2 μ m.
- As for the Hall plates, it is chosen to place eight plates in parallel: four on each side of the current rail for differential sensing. This is done for two reasons. First of all, by orthogonally biasing the parallel Hall plates (i.e., letting the Hall bias current flow in different directions through each Hall plate), the offset can be reduced as several geometry-related sources of offset cancel out [5] (the resulting offset will be well below 1mV). Secondly, by placing them in parallel, the resistance can be reduced. This is beneficial for the overall noise performance². Biasing the Hall plates requires roughly 2-3mA. This power consumption is the main cost related to the hybrid-hybrid architecture³.

The resulting model parameters for the coil and Hall plates can be found in Table 7.1.

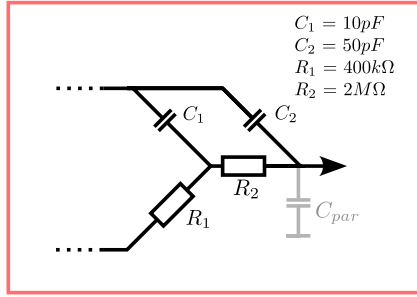


Figure 7.2: Implementation of the summing filter as indicated by the red outline in Figure 7.1. The high-frequency path will be connected to the capacitance, while the low-frequency path will be connected to the resistance.

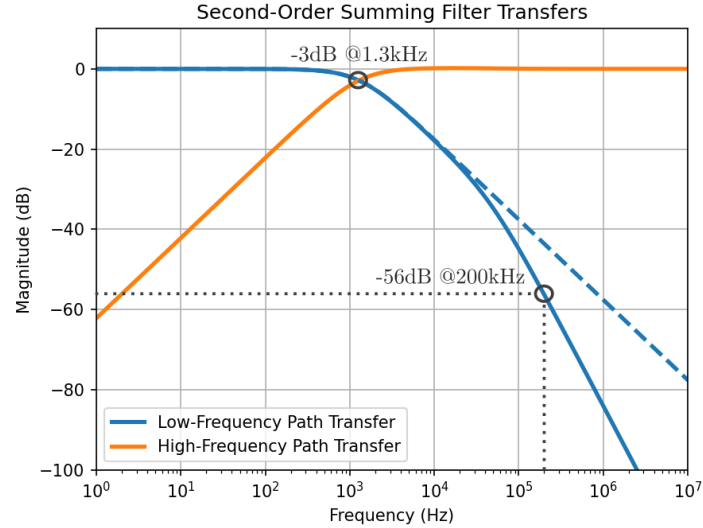


Figure 7.3: Simulated low-pass and high-pass transfer of the summing filter shown in Figure 7.2

7.2. SUMMING FILTER IMPLEMENTATION

Next, the implementation of the summing filter, the red block in Figure 7.1, will be briefly discussed. This filter, used to combine the low- and high-frequency path, is shown in Figure 7.2. As discussed in Section 4.2.1, to obtain a certain level of attenuation at the spinning frequency, a second-order filter generally requires considerably smaller components than a first-order filter. This is why this implementation is chosen for this design. Besides, this network is quite simple to implement as it does not require any active circuitry.

According to previous design data by SystematIC, a typical equivalent spinning ripple of approximately $19A_{RMS}$ is generated at the Hall plates of the low-frequency path⁴. In this design, to keep this ripple well below the integrated noise, a reduction of 56dB at the spinning frequency of 200kHz is aimed for (resulting in an equivalent ripple of $30mA_{RMS}$ ⁵). This can be achieved with the component sizes as indicated in Figure 7.2. These particular values are chosen based on area and gain flatness considerations. Namely, to reduce any gain flatness errors, C_2 should be considerably larger than any parasitic capacitance (C_{par}) at the output of the summing filter⁶. The resulting frequency magnitude response of both paths is shown in Figure 7.3 and clearly shows the filtering benefit of the second-order filter compared to a first-order filter with roughly the same component sizes (dotted line).

7.3. HIGH-FREQUENCY PATH READOUT ARCHITECTURE

After discussing the summing filter, in this section, the implementation of the high-frequency path, the blue block in Figure 7.1, will be discussed. The chosen implementation is shown in Figure 7.4.

First of all, as can be seen, this implementation realizes the summation of the Hall plate and coil signal in the voltage domain by placing the devices in series. This is mainly done for the sake of simplicity. Namely, in this case, no separate amplifiers must be designed to read out the coil- or Hall voltage in the high-frequency path. Besides reducing design time, this should also help reduce area and power consumption. Also, note that two coils are placed on either side of the Hall plate to create a balanced signal; the electrical model is shown in Figure 7.5.

²For the same reasons, also eight Hall plates are used in the low-frequency path

³However, compared to the previous Hall-Hall design by SystematIC, effectively power consumption does not increase as in that design also eight Hall plates were used in the high-frequency path

⁴This is essentially the only required information from the low-frequency path to be able to design this new prototype chip properly

⁵Some margin is taken into consideration w.r.t. $80mA_{RMS}$ of noise to account for filter cross-over variation due to process spread and variation in the ripple magnitude (due to different Hall plate offsets)

⁶According to simulations, this parasitic capacitance is around 1pF, $C_2 = 50\text{pF}$ should keep a gain flatness error well below 1%

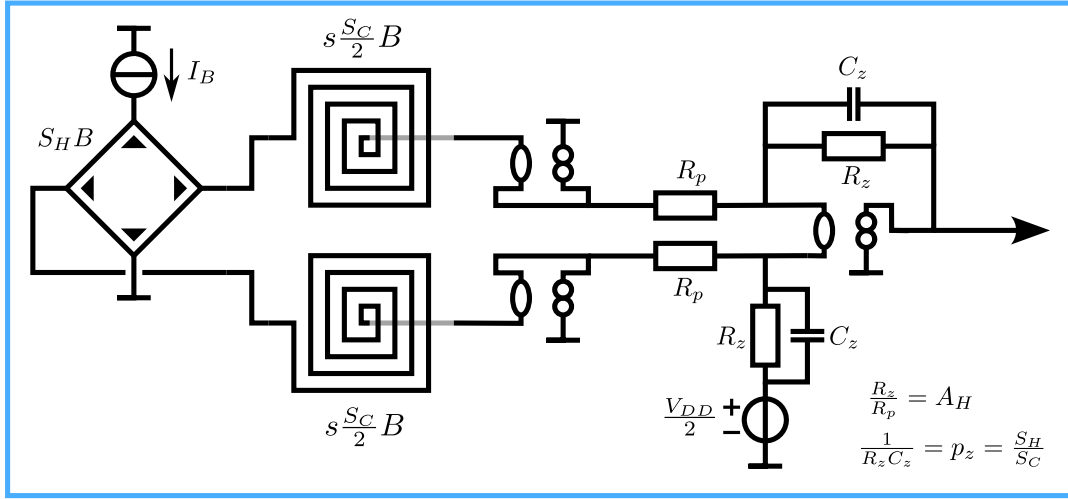


Figure 7.4: Implementation of the high-frequency path as indicated by the blue outline in Figure 7.1

As the coil and Hall plate experience the same DC gain, the downside of this approach is that the cross-over frequency $f_z = \frac{p_z}{2\pi}$ (see Figure 7.1) cannot be freely chosen anymore and must be equal to $f_z = \frac{1}{2\pi} \frac{S_H}{S_C} \approx 5.1\text{kHz}$. Unless the equivalent noise resistances of the coil and Hall plates are equal (see Section 3.4), this cross-over frequency is not optimal for noise performance⁷. However, as will be shown later, the noise specifications can still be met with this approach, and therefore, this non-optimal cross-over frequency is no large problem in this design.

Next, the summed Hall and coil voltages should be properly read out. This is done by means of buffers (to obtain a large input impedance) and a transimpedance-like amplifier. This amplifier, henceforth referred to as a transimpedance amplifier or TIA, has two functions. Namely, it implements the required low-pass filtering to obtain a flat response over frequency and realizes a differential to single-ended conversion.

Again, simplicity plays a big role in the selection of this particular implementation. Namely, the controllers [17] of the buffer and the transimpedance amplifier have very similar specifications regarding noise performance, drive capability, output swing, offset, etc. Therefore, both controllers can be implemented in almost the same way. Although other considered approaches could potentially be more power-, noise- and area-efficient, they were believed to be more time-consuming to implement. Besides, as will be shown, the specifications can still be met with this solution. Moreover, this implementation especially highlights the strength of the hybrid-hybrid architecture in the sense that good performance can still be obtained with a relatively simple design.

7.4. CHOOSING THE COMPONENT SIZES

To properly combine the low- and high-frequency paths. Both paths should have the same gain right before the summing filter. In other words, at any frequency, the high-frequency path must have a gain equal to the low-frequency path gain $S_H A_H$ from the magnetic field to the path's output voltage. At low frequencies, the high-frequency path response is dominated by the Hall plate transfer S_H , and the amplifier's low-frequency gain $\frac{R_z}{R_p}$. As a result, if the low-frequency and high-frequency path Hall plates have the same sensitivity, the following equality must be adhered to:

$$S_H \cdot \frac{R_z}{R_p} = S_H A_H \quad (7.1)$$

On the other hand, at high frequencies, the high-frequency path's gain is determined by the coil transfer $s S_C$, and the amplifier's high-frequency transfer $\frac{1}{s R_p C_z}$ (asymptote of the low-pass filter⁸). As a result, also the

⁷ Another disadvantage is that $f_z > f_x$ and thus, as explained in Section 4.3.3, any mismatch between the low- and high-frequency path will be seen as a gain flatness error. For $f_z = 5.1\text{kHz}$ and $f_x = 1.3\text{kHz}$ a $\pm 1\%$ error will be seen as a $\pm 0.8\%$ error in the gain flatness

⁸ As mentioned before, this resistance-capacitance product is the main source of high-frequency gain drift

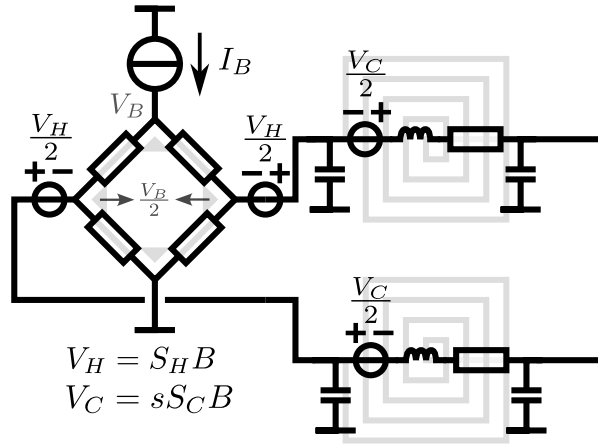


Figure 7.5: The transducer symbols in Figure 7.4 replaced with their electrical model

following equality must be adhered to:

$$\frac{S_C}{R_p C_z} = S_H A_H \quad (7.2)$$

Especially this latter equation highlights an interesting trade-off. Namely, as can be seen from Figure 7.4, the resistance R_p directly adds noise as if it were in series with the coil (and Hall plate). Therefore, ideally, R_p should be considerably smaller than the coil resistance R_c to obtain a noise figure well below 3dB. However, after rewriting Equation 7.2 as:

$$C_z = \frac{S_C}{R_p} \cdot \frac{1}{A_H S_H} \quad (7.3)$$

It can be seen that the required capacitance C_z is inversely proportional to R_p . Reducing the readout noise⁹, therefore, comes at the cost of increased capacitor area, which, for typical component values, could easily be several tens of pF. Besides, the current drive capability of the amplifiers must be increased. Namely, both buffers and the transimpedance amplifier must be able to sink and source a current as large as $\frac{V_{max}}{R_p}$. This could put strict requirements on the output stage. It should be noted that this trade-off does not only hold for the hybrid-hybrid architecture but holds for any other implementation in which the coil response is integrated/flattened in one step¹⁰. This trade-off was quite important during the design of the prototype chip.

There seems to be a trade-off between **noise performance**, **area consumption** and **current drive capability** when implementing the integrating transfer to flatten the coil response

SELECTING THE GAIN A_H

As can be deduced from the equations above, before selecting the passive component sizes, R_p , R_z , and C_z , the gain A_H should first be selected. Selecting this gain is also about balancing certain trade-offs. Namely:

- Choosing a large gain A_H results in a smaller required capacitance C_z and puts less strict noise and offset requirements on the amplifiers in the common path (i.e., the signal path after combining the low- and high-frequency paths) due to a larger preceding gain.
- On the other hand, choosing a small gain A_H reduces the chances of clipping in the hybrid-hybrid architecture (due to amplifier offset) and reduces the required component size for R_z . Whether R_z or

⁹Note that another advantage of reading out the coil current is the fact that R_c itself is used in the transfer, and, thus, no additional resistance R_p is required which adds noise

¹⁰In [11–13] the area trade-off was broken by splitting up the integrating transfer in multiple steps by matching various poles and zeros. However, as with many of the trade-offs so far, this increases design complexity

	Value
Cross-Over Frequency $f_z = \frac{p_z}{2\pi}$	5.1kHz
Cross-Over Frequency $f_x = \frac{p_x}{2\pi}$	1.3kHz
DC-gain A_H	160x (44dB)
Resistance R_p	7.5k Ω
Resistance R_z	1.2M Ω
Capacitance C_z	26pF
Coupling Factor A_{co}	310 μ T A ⁻¹

Table 7.2: Several typical values associated with the high-frequency path implementation

C_z dominates in the total area is highly dependent on the coil sensitivity and resistance. However, for a small coil, C_z generally requires a considerably larger area than R_z .

With these trade-offs in mind, the gain for the prototype chip was chosen to be 160x (44dB).

SELECTING R_p , C_z AND R_z

Having selected the gain A_H , the various components, R_p , C_z , and R_z , constituting the feedback network, can be selected. Moreover, according to Equations 7.1 and 7.2, only one component value can be selected independently. Eventually, as the available area was limited, a resistance R_p of 7.5k Ω was chosen, which, in turn, resulted in a resistance R_z of 1.2M Ω and a typical capacitance C_z of 26pF¹¹. Besides, the current drive capability requirement of the amplifiers is only $\frac{V_{max}}{R_p} \approx \pm 130\mu$ A, which is quite doable.

As a downside, with a 15k Ω resistance effectively placed in series with the 1.4k Ω coil- and 700 Ω Hall plate resistance, the large resistance R_p causes a relatively large noise penalty. This causes a noise increase from roughly 20mA_{RMS} to approximately 55mA_{RMS} (using the Equations presented in Section 4.13). As mentioned, reduced capacitor area and a reduced drive capability requirement come at the cost of reduced noise performance. In this design, this large penalty is not a big issue as it still leaves sufficient margin for other noise contributions before reaching the 80mA_{RMS} specification. Namely, it was determined that in total, approximately an equivalent noise resistance of 35k Ω could be placed in series with the coil/Hall plate before 80mA_{RMS} is reached¹². A margin of 20k Ω therefore remains.

Lastly, it should be noted that the resistances must be implemented with p-poly resistors and the capacitances with MIM capacitors for a low variation over temperature. Besides, C_z is implemented as an 8-bit capacitor bank. This is done to be able to trim the gain of the high-frequency path during foreground calibration¹³. Moreover, to be able to account for all variations from the magnetic field coupling factor A_{co} (from 250 to 400 μ VT⁻¹), the resistance spread of R_p ($\pm 20\%$) and the capacitance spread of C_z itself ($\pm 20\%$). The capacitor must, instead of 26pF be as large as 42pF. For the sake of clarity, the following sections will only present the nominal case (with the various parameters as shown in Table 7.1 and 7.2). However, it was verified that the chip still works properly for the full range of component values.

¹¹Note that for the chosen implementation, each component has to be implemented two times

¹²Note that almost all noise comes from the high-frequency path, as the low-frequency path noise is significantly filtered by the summing filter with its low-frequency cross-over frequency

¹³The low-frequency gain (in both the low- and high-frequency path) can be trimmed with the Hall plate bias current

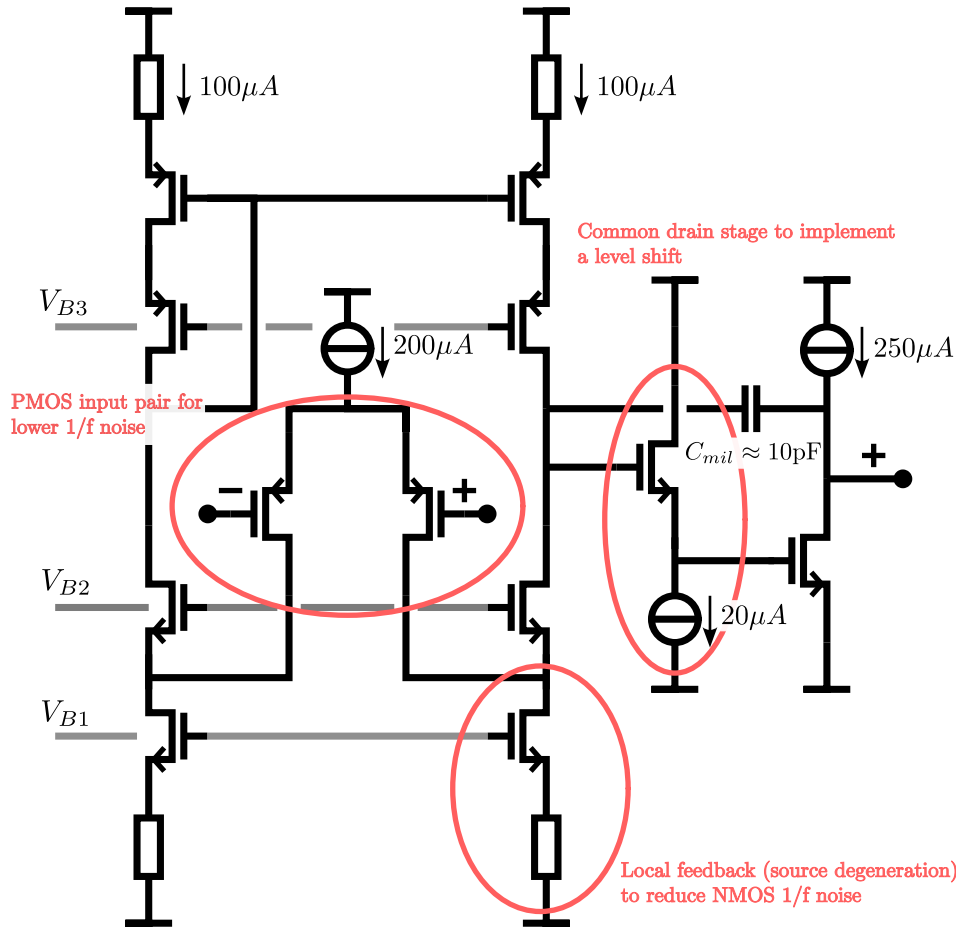


Figure 7.6: Architecture for the controller of the transimpedance amplifier and buffers. V_{B1} , V_{B2} and V_{B3} indicate a connection to a certain bias voltage.

7.5. CONTROLLER IMPLEMENTATION

Having obtained the various values for the feedback network, the next step is to implement the controllers for the buffers and the TIA. As mentioned before, the same controller implementation will be used for both amplifiers, as they require very similar specifications. First, these specifications will be derived, and secondly, the chosen controller architecture will be elaborated on.

7.5.1. CONTROLLER REQUIREMENTS

The controller requirements are given in Table 7.3. Most requirements are quite self-explanatory. However, some requirements should be briefly clarified.

1. As mentioned before, to keep the integrated noise level below 80mA_{RMS} , in total, maximally $35\text{k}\Omega$ can effectively be placed in series with the coil. With $2R_p = 15\text{k}\Omega$ already used in the feedback network, this means that each amplifier's controller is maximally allowed to add roughly the noise of a $6.7\text{k}\Omega$ equivalent resistance. However, in this case, $1/f$ -noise has not yet been considered. It was calculated that, with $1/f$ -noise taken into consideration, each controller may add the noise of a $1\text{k}\Omega$ resistor ($4\text{ nV}/\sqrt{\text{Hz}}$) if the $1/f$ corner frequency is below 20kHz . Of course, it is also possible to have a slightly higher noise floor in combination with a lower $1/f$ corner frequency.
2. The amplifier output swing should be considered. Namely, with a $\pm 80\text{A}$ input signal, a $310\mu\text{T}^{-1}$ coupling factor, a $80\text{mV}\text{T}^{-1}$ Hall plate sensitivity and a gain of $160\times$, the high-frequency paths output swing is roughly $\pm 320\text{mV}$. However, that is not all. Namely, the system must still work for various levels of offset. Therefore, the larger the possible output swing, the larger the offset can be. In this design, an output swing of $\pm 1.55\text{V}$ is aimed for. For the minimum supply voltage of 3.3V , with the output biased at $\frac{V_{DD}}{2}$, this leaves 100mV of headroom to both ground and V_{DD} .

Requirement	Value
Input Referred Noise (v_n)	$< 4 \text{ nV Hz}^{-0.5}$
1/f Noise Corner (f_c)	$< 20\text{kHz}$
Output Swing	$\pm 1.55\text{V}$
Offset (σ)	$< 1.48\text{mV}$
Minimum Input Voltage	$< 100\text{mV}$
Maximum Input Voltage	$> \frac{V_{DD}}{2} + 0.5\text{V}$
Current Drive Capability	$\pm 130\mu\text{A}$
Low-Frequency Gain Inaccuracy	$< 1\%$
Gain Temperature Variation	$< \pm 3\%$
Gain-bandwidth product	$> 5\text{MHz}$
Quiescent Bias Current	$< 800\mu\text{A}$
Supply Voltage	$3.3 - 5.5\text{V}$

Table 7.3: Requirements for each of the controllers used in the buffers and transimpedance amplifier

- Using this output swing, the maximum allowed offset can be determined. Namely, with the given signal swing and output range, the output referred offset is maximally allowed to be 1.23V. To obtain a chip yield of more than 99.7% (3σ), the standard deviation (σ) of the total offset should be below roughly 360mV. Input referred with a gain of 160x and divided over three amplifiers (with uncorrelated offset sources), this results in a maximum standard deviation of roughly 1.48mV for each controller.
- The input voltage range should be considered. As the input common-mode voltage of the amplifiers is essentially set at half the Hall plate bias voltage (see Figure 7.5), the input voltage can be relatively low. More specifically, depending on the bias current (which is used to trim the Hall plate sensitivity), this common-mode voltage can be as low as 600mV or as high as $\frac{V_{DD}}{2}$. With a maximum coil signal of $\pm 500\text{mV}$ (single-ended, or $\pm 1\text{V}$ differential), this means that the amplifier input must be able to go from roughly 100mV to $\frac{V_{DD}}{2} + 0.5\text{V}$.
- The amplifiers must be able to at least source and sink $130\mu\text{A}$. This is calculated using the maximum coil voltage of $\pm 1\text{V}$ and the resistance R_p of $7.5\text{k}\Omega$.
- The low-frequency gain inaccuracy due to insufficient loop gain should be well below 1%. This is done such that it does not introduce an additional significant mismatch between the low- and high-frequency Hall plates. This helps reduce the chances that an additional third foreground calibration step would be required.
- The gain temperature variation is explicitly restated to indicate that the controller should not cause significant temperature variation due to, for example, too little loop gain at high frequencies.

All other requirements should be quite self-explanatory.

7.5.2. CONTROLLER ARCHITECTURE

To meet these requirements, the three-stage controller architecture shown in Figure 7.6 is believed to be quite effective. As can be seen, the first stage consists of a PMOS differential pair, an NMOS (folded) cascode, and a PMOS current mirror, while the second and third stage consists of a common-drain stage followed by a common-source class-A output stage. A brief justification behind these signal path design choices is presented below:

- The PMOS differential pair is mainly chosen for its superior 1/f noise performance compared to the NMOS variant. Besides, with a PMOS input pair, the amplifier can, as required, still function well with a low input voltage. A single CS stage was also considered for better noise performance. However, properly biasing such a stage adds complexity and, eventually, has its own noise penalty.
- The NMOS cascode is added to increase the low-frequency loop gain such that the low-frequency gain inaccuracy is well below $\pm 1\%$. Compared to a PMOS cascode¹⁴, this folded cascode architecture is less power efficient due to the additional required biasing branch. However, in this way, a large input swing can be maintained.
- The PMOS current mirror is used for a proper differential to single-ended conversion inside the controller. This has the advantage that the circuit is more robust to common-mode interference (e.g., bi-

¹⁴Resulting in a so-called telescopic architecture

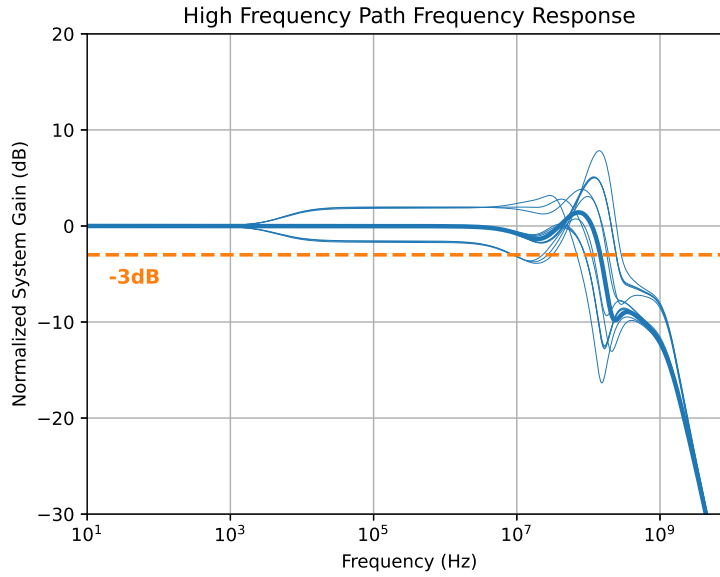


Figure 7.7: Frequency response of the high-frequency path (including coil parasitics, the buffer, and TIA) over all corners and various temperatures. The shift in gain at high frequencies is due to changing capacitance over corners; this can be trimmed out. As can be seen, a bandwidth above 10MHz is easily obtained.

asing interference or noise). Besides, not half of the current is "wasted," and, as a result, the transconductance of the first stage is effectively doubled.

- A single stage is not sufficient as it mainly conflicts with the large output swing and the low-frequency loop gain requirement (because of the relatively low 7.5k Ω load resistance). Therefore, another stage is required.
- For this purpose, a simple class-A output stage is chosen. This is mainly done for the sake of simplicity. Besides, with the drive capability requirement of $\pm 130\mu\text{A}$, the quiescent power consumption can be kept within budget.
- Lastly, a common-drain stage is placed between the input- and output stages to provide a level shift. Namely, for certain temperatures and corners, the output stage's gate voltage would be so low that it would force the input stage's NMOS cascode out of saturation.

Lastly, it should be noted that frequency compensation is implemented by means of Miller compensation. Besides, the implementation of the bias current sources of the input stage is explicitly shown in Figure 7.6 to indicate that local feedback (source degeneration) is used to reduce the effect of $1/f$ -noise of the biasing transistors. This is, for the same reason, also done for the PMOS current mirror. In the next section, the performance of this controller inside the high-frequency path will be verified through simulations. In that section, some more details about this controller's implementation will also be provided.

7.6. HIGH-FREQUENCY PATH PERFORMANCE VERIFICATION

Having discussed the implementation of the high-frequency path, this section will verify whether this implementation adheres to the requirements imposed by SystematIC. Both the system requirements (as can be found in Section 2.4) and some more specific high-frequency path implementation requirements (as introduced in this chapter) will be discussed. The general system requirements are mainly related to bandwidth, noise, and gain stability, while the more specific high-frequency path requirements are related to the output swing and offset. Besides, the various cost factors will be briefly discussed, and a simulation related to coil signal clipping will be presented. Lastly, it should be noted that the system functionality was extensively verified over various process corners, temperatures, and supply voltages. However, not all simulations are shown for the sake of conciseness.

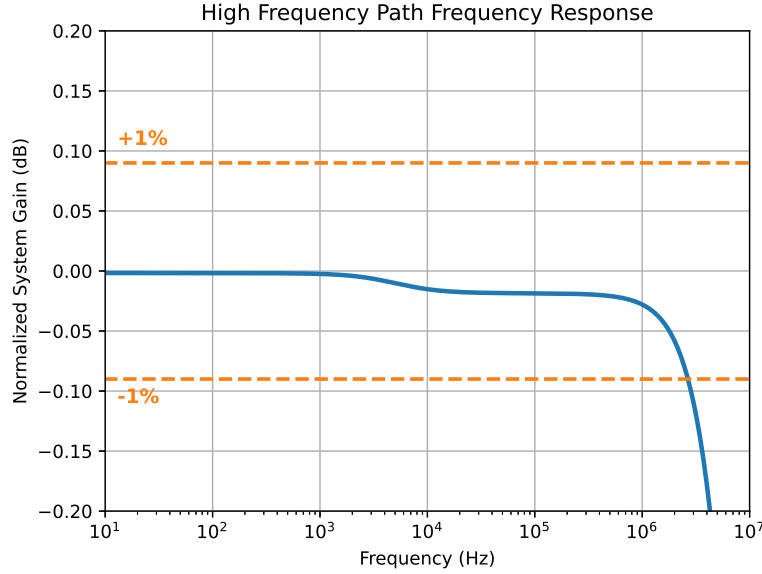


Figure 7.8: Frequency response of a trimmed high-frequency path. The gain stays well within $\pm 1\%$ gain flatness up until roughly 4MHz

7.6.1. BANDWIDTH & GAIN FLATNESS

First of all, the bandwidth should be discussed. In the common path (after combining the low- and high-frequency path), SystematIC already had an amplifier limited to 2MHz. Therefore, to ensure the bandwidth is not further limited, for the high-frequency path, it was decided that the bandwidth should be above 5MHz. With the buffers and transimpedance amplifier, this could quite easily be achieved as the feedback network does not present much attenuation at high frequencies¹⁵. As can be seen in Figure 7.7, over corners, the overall high-frequency path bandwidth easily stays above 10MHz¹⁶. Moreover, as can be seen in Figure 7.8, the (trimmed) gain shows a small gain change around the cross-over frequency (due to finite trimming accuracy) but stays well within the $\pm 1\%$ gain flatness bound.

7.6.2. NOISE

Next, noise should be considered. For this purpose, it was calculated that each amplifier should maximally have an input referred noise of an equivalent $1\text{k}\Omega$ noise resistance and a $1/f$ corner frequency of maximally 20kHz to keep the integrated noise level below 80mA_{RMS} . As shown in Figure 7.9, these specifications have been met. Namely, the obtained spectrum (orange curve) is below the specified maximum spectrum (blue curve). However, to obtain this performance, relatively large input devices with a $200\mu\text{A}$ bias current were required (1.2mm width and 500nm length). Moreover, local feedback had to be employed to reduce the $(1/f)$ noise contribution from (especially the NMOS) bias sources. All in all, obtaining this noise spectrum was not necessarily straightforward.

After fully implementing the high-frequency path, the system noise spectrum, as shown in Figure 7.10, was obtained. In this figure, the purple curve indicates the simulated noise spectrum, while the red curve shows a calculated noise spectrum¹⁷. Using this (accurate) calculation script to gain some more insight, the various noise contributions from the low-frequency Hall path, the high-frequency Hall path, and the high-frequency coil path could also be estimated (blue, orange, and green curves, respectively). Namely, it can, for example, be seen that the high-frequency coil path contributes significantly more (75mA_{RMS}) noise than the high-frequency Hall path (11mA_{RMS}). By shifting the cross-over frequency f_z to a higher frequency, this noise could have been balanced better, and for $f_z \approx 30\text{kHz}$, it was determined that an overall integrated noise level

¹⁵The buffer has unity gain feedback, while the transimpedance amplifier effectively has a low-frequency (phantom) zero in the loop gain at f_z which extends the bandwidth

¹⁶Note that the variation of the high-frequency gain is caused by the varying feedback resistance and capacitance over corners, this can be trimmed out using the capacitor bank

¹⁷This was done using a Python script, which calculates the noise spectrum according to various given transfer functions for the low- and high-frequency path. In this calculation, also a $1/f$ noise spectrum is incorporated in the high-frequency path

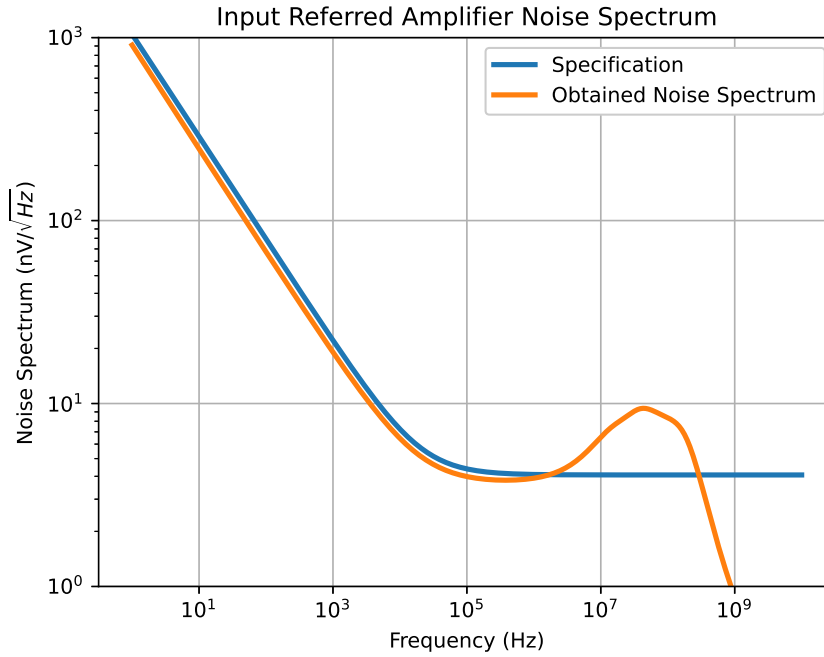


Figure 7.9: Specification input referred amplifier voltage noise spectrum (blue) versus the obtained noise spectrum (orange)

of 40mA_{RMS} could have been obtained. However, as explained before, this would have come at the cost of complexity, which, for this first prototype, was undesirable. Besides, with a total integrated noise level of 77mA_{RMS} , the noise specification of 80mA_{RMS} has been met anyway.

Lastly, another strength of the hybrid-hybrid architecture could be highlighted. Namely, using the same script used to calculate the noise spectra shown in Figure 7.10, it can be estimated how a normal hybrid architecture (using, for example, an integrator as explained in Section 4.2) would have performed with the given readout architecture (i.e., the architecture shown in Figure 7.4 without the Hall plate and with R_z considerably larger to effectively implement an integrator). The resulting noise spectrum, as shown in Figure 7.11, clearly indicates a huge noise increase to roughly 560mA_{RMS} . This is mainly caused by the amplifier $1/f$ noise and the fact that the main summing filter cross-over frequency f_x is far from optimal¹⁸. Besides having to deal with the already discussed difficulties associated with a non-hybrid-hybrid architecture (see Chapter 4), reducing $1/f$ noise would have resulted in an additional challenge. This again beautifully highlights the strength of the hybrid-hybrid architecture in which noise performance can be orthogonally designed from ripple performance.

7.6.3. GAIN STABILITY

Another important design aspect discussed throughout this thesis is gain stability over temperature (and lifetime). For this prototype chip, it was specified that the gain should not drift more than $\pm 3\%$. However, in a future design, ideally, SystematIC would like to see a maximum error of $\pm 1\%$. For the low-frequency path, no changes to SystematIC's previous Hall-Hall design have been made and will, therefore, not be considered further. For the high-frequency gain, it was determined that using a p-poly resistor and MIM capacitor to implement the integrating transfer should keep the gain temperature drift relatively low. Namely, it was calculated that over the specified temperature range from -40°C to $+125^\circ\text{C}$, a maximum temperature variation of less than 1% could be obtained. The orange curve in Figure 7.12 shows this calculated temperature dependence.

The simulation results, on the other hand, show a larger temperature dependence. Namely, the simulated gain variation, as shown by the blue curve in Figure 7.12, varies by roughly 2.5%. It was determined that this larger-than-expected temperature variation came from the temperature variation of the transimpedance am-

¹⁸Note that if this architecture would have been used, f_x could be increased to balance ripple and noise performance better. But still, it is believed that a noise level below 150mA_{RMS} , let alone 80mA_{RMS} , would have been difficult to obtain

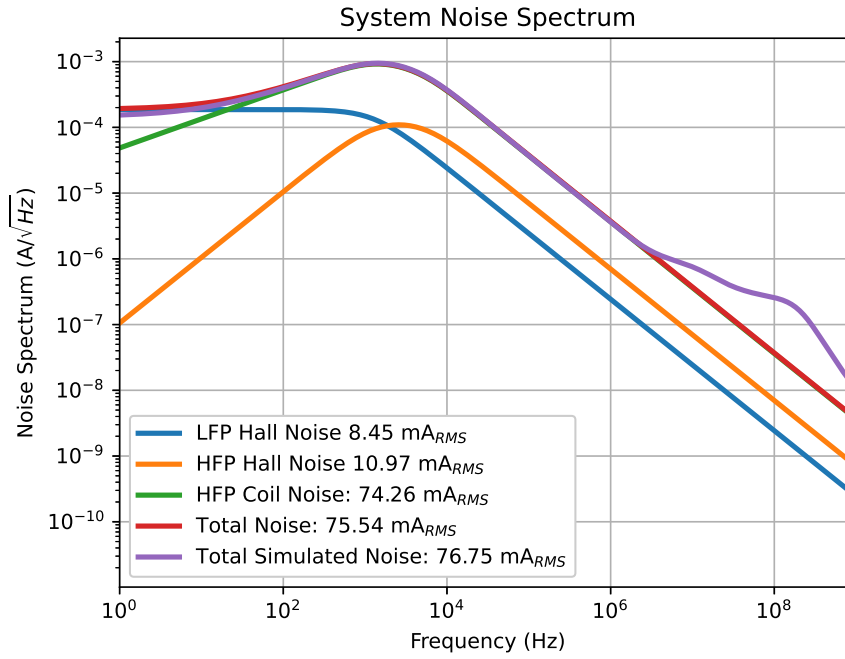


Figure 7.10: Simulated noise spectrum (purple) versus calculated noise spectrum (red) for the hybrid-hybrid architecture. The blue, orange, and green curves show the calculated low-frequency Hall path, high-frequency Hall path, and high-frequency coil path noise contributions, respectively (this includes noise from the readout amplifiers). The calculation matches the simulation quite accurately. The total integrated noise is below $80mA_{RMS}$.

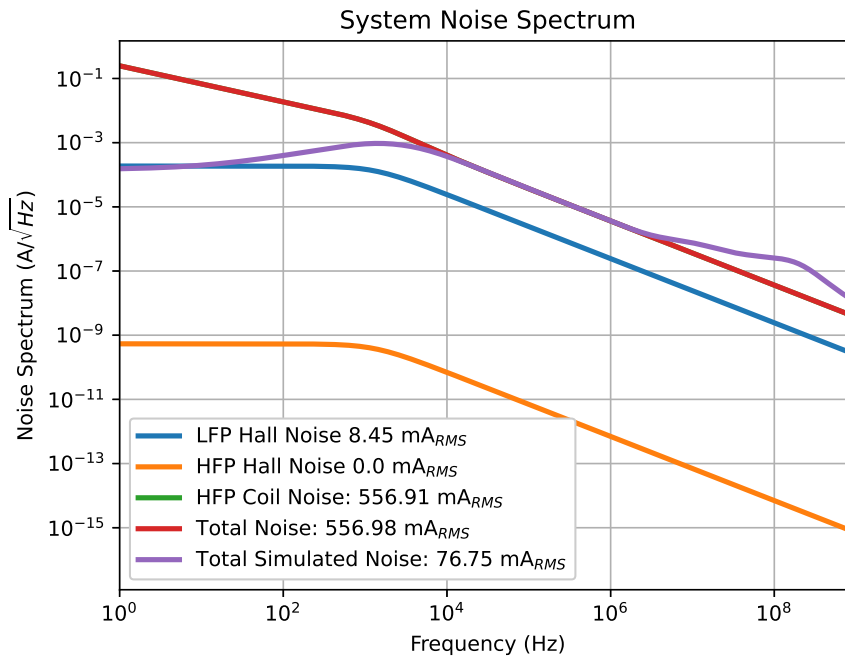


Figure 7.11: Using the same calculation script for Figure 7.10, the total noise in case a hybrid-hybrid architecture were not to be used is calculated to be roughly $560mA_{RMS}$. This huge noise increase is mainly caused by $1/f$ noise and a far from optimal cross-over frequency $f_x = 1.3kHz$.

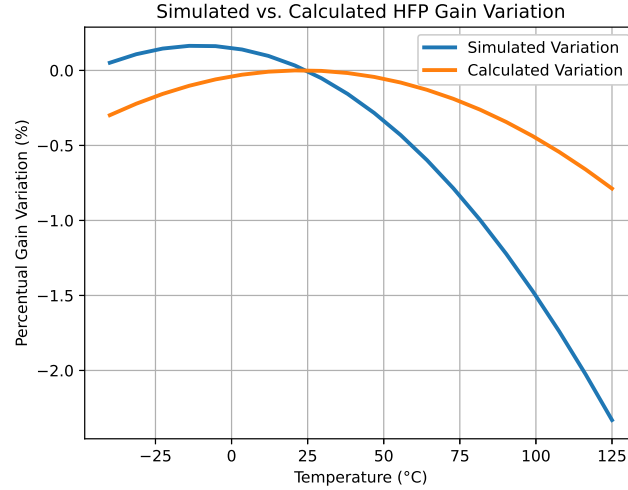


Figure 7.12: Simulated versus calculated high-frequency gain variation over temperature. The calculated variation only takes into consideration resistance and capacitance temperature variation. The simulated variation also taken various second-order effects into consideration, the temperature variation was simulated at 100kHz

plifier loop gain at high frequencies. Namely, at high frequencies, the loop gain would only be around 30dB. Moreover, a variation of this loop gain from 30dB to 27dB (as observed) roughly introduces this additional 1.5% of gain variation. Although for this design, the gain variation over temperature is within specification, for a future design, this additional temperature variation should be looked into. Moreover, lifetime drift cannot easily be simulated and should be characterized after device fabrication.

7.6.4. OFFSET & OUTPUT SWING

Offset is another important design aspect when considering the coil-Hall architecture. With the chosen hybrid-hybrid architecture, the high-frequency path does not contribute any offset to the overall system output. However, as mentioned in the controller requirements, offset in combination with the amplifier output swing is still important to consider in light of signal clipping. Namely, it was specified that each amplifier should maximally have an input referred offset standard deviation of 1.48mV and that the output swing should be at least 1.55V (leaving 100mV headroom to the ground and the supply).

The offset standard deviation was obtained using a Monte Carlo simulation with 1000 runs. The resulting offset distribution can be found in Figure 7.13. Moreover, the calculated standard deviation is roughly 1.21mV. This is well within the specifications¹⁹. The simulated output swing of the high-frequency path for various offset levels is shown in Figure 7.14. In these simulations, a full-scale (80A), 2MHz input signal was applied. As can be seen, this signal can still be processed within the specified output range (as indicated by the dashed red lines). With these results, it can be calculated that at least 99.98% of the chips should not clip on their own offset²⁰. With a larger supply voltage, this chance is even reduced further²¹.

7.6.5. COIL CLIPPING & SLOW SETTling

Next, an effect related to coil voltage readout should be discussed. Namely, the coil signal can become quite large when a fast input signal is present. Especially when reading out the coil voltage, this signal should be limited by a clamping circuit so that it does not cause transistor breakdown. However, this, in turn, causes significant distortion. As mentioned in Section 3.6, especially for short circuit detection, this can be quite an issue.

Namely, if the signal rise time is too fast, all of a sudden, the system reacts considerably slower. Instead of detecting the short circuit within less than a microsecond (associated with a 2MHz bandwidth), it can

¹⁹Moreover, this offset level could be reached with the input pair dimensions already required for noise performance

²⁰SystematIC also required multiple gain settings, for some of these gain settings this chance is reduced to 99.7%. However, this is still within specification.

²¹Of course, the contrary is also true for a smaller supply voltage

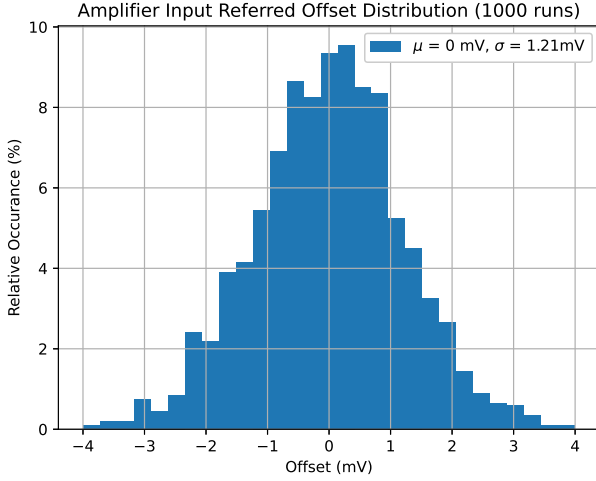


Figure 7.13: Distribution of the input-referred offset of the controller shown in Figure 7.6. This data was obtained using a Monte Carlo simulation with 1000 runs.

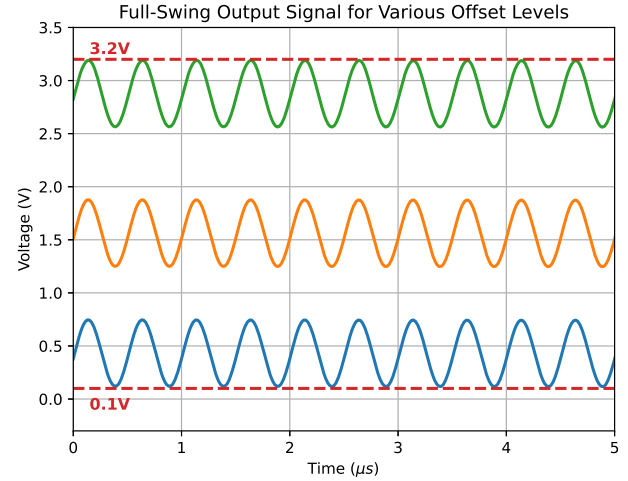


Figure 7.14: Simulated transimpedance amplifier output with a full-scale (320mV), 2MHz sinusoidal signal with various levels of offset

take tens or even hundreds of microseconds associated with the low-frequency path cross-over frequency (1.3kHz). This can be seen in the simulation result shown in Figure 7.15.

In this figure, both the coil signal (top plot) and the output signal after summing the low- and high-frequency path (middle and bottom plots) can be seen for three different rise times of a full-scale signal step. For a 400ns and 100ns rise time (blue and green curves, respectively), the coil signal is not clipped, and it can be seen that the output signal is, as expected, a full-scale step with the specified rise times. However, for a 25ns rise time, the coil signal becomes too large and is clipped against the 3.3V supply. As a result, the high-frequency path cannot fully respond, and the signal only rises to a fraction of the intended value. The final value is only reached after the low-frequency path has settled; this takes orders of magnitude longer than the time associated with the 2MHz bandwidth. This can be seen in the blue curve from Figure 7.15.

As mentioned in Section 3.6, for voltage readout, the chances of this happening can only be reduced by using a smaller coil, by using a larger supply voltage, and by using transistors rated for a higher voltage. When reading out the current, this clipping can certainly also occur. However, in this case, the maximum signal is limited by the amplifier's current drive capability. This can more effectively be increased, by, for example using a class-AB output stage with large transistors.

7.6.6. POWER- & AREA CONSUMPTION

Lastly, the power- and area consumption should be briefly discussed. SystematIC indicated that this new coil-Hall prototype should fit within the same area as the previous Hall-Hall design and consume roughly the same amount of power. As the new design uses the same amount of Hall plates in the high-frequency path (requiring the same bias current), the new architecture does not introduce a power penalty in that regard. Furthermore, as the previous high-frequency path amplifiers required roughly 2.4mA of bias current, it was specified that each of the three amplifiers in the new design should, ideally, require less than 800 μA each. As indicated in Figure 7.6, each amplifier eventually required approximately 750 μA . This is within specification. Regarding area consumption, the large passive components (especially the capacitors) required to implement the various transfer functions introduced some difficulties. However, eventually, the new design was made to (just) fit within the old Hall-Hall layout. This was made possible by SystematIC's great layout team.

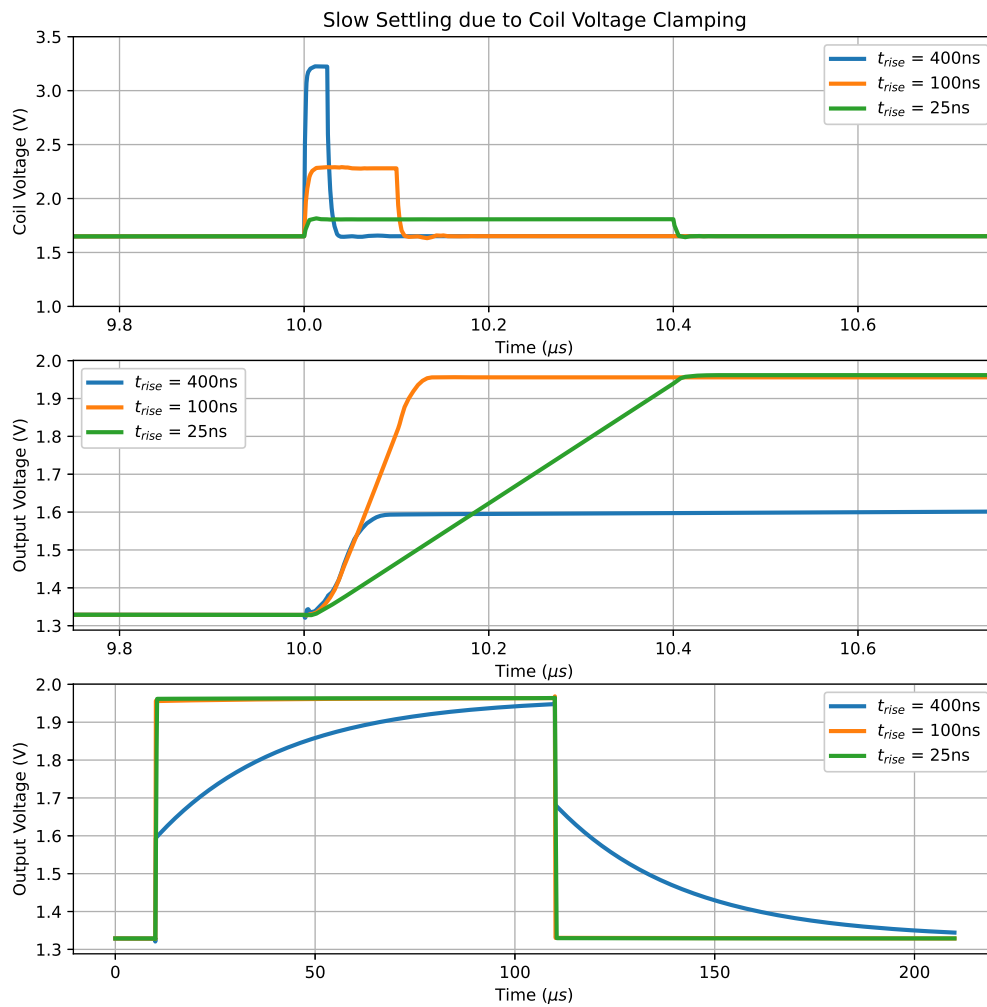


Figure 7.15: If the coil signal clips in the high-frequency path, a large slow-settling component occurs at the system output as the high-frequency path cannot fully react to the fast input signal. The top plot shows the coil signal, while the middle and bottom plots show the resulting system output (with the middle plot being a zoomed-in version of the bottom plot at 10μs). The input signal is a full-scale 80A step with three different rise times.

7.7. CONCLUSION & FUTURE DESIGN RECOMMENDATIONS

This chapter presented the design of the coil-Hall hybrid current sensor for SystematIC. In this design, the so-called hybrid-hybrid architecture was chosen for two main reasons. Firstly, it alleviates the tight trade-off between ripple and noise performance without requiring any changes to the low-frequency path. Secondly, it allows for a relatively simple implementation without having to deal with large offset-related problems and $1/f$ -noise.

These advantages were clearly demonstrated in the implementation discussed in this chapter. For the high-frequency path, the only significant design effort was developing a single, relatively simple controller used in both the buffer and the TIA. No additional measures were needed against offset and $1/f$ noise. For other combining architectures, such a straightforward solution likely would not have been sufficient. For example, in Section 7.6.2, it was shown that the lack of orthogonality in the design of noise and ripple performance would have resulted in an integrated noise level well above 150 mA_{RMS} . Besides, a larger design effort would have been required to deal with issues related to offset. All in all, this chosen architecture was the perfect choice to satisfy the needs of SystematIC, and, as shown in Table 7.4, the new design is a great step forward compared to their previous design²².

FUTURE RECOMMENDATIONS

However, there is certainly still some room for improvement. Namely, the simplicity of the design mainly comes at the cost of gain stability and noise performance²³. If SystematIC would like to improve the design in the future, some recommendations can be given. These are ordered based on the level of required modification and added complexity. In terms of gain stability, two points of improvement can be given:

1. First of all, the higher-than-expected high-frequency gain variation over temperature should be looked into. This can probably be solved by designing an amplifier controller (for the TIA) which allows for more loop gain at higher frequencies.
2. Secondly, a better gain stabilization technique for the Hall plates could be considered. For this, various ideas have been presented in Chapter 6.

In terms of noise performance, some more recommendations can be given. These are as follows:

1. Firstly, the easiest way to optimize noise performance would be to optimize the cross-over frequency f_z . To achieve this, instead of placing the Hall plate directly in series with the coil, the Hall plate must be read out with a separate amplifier. This allows for a difference in DC gain between the coil- and Hall path, and, as a result, f_z can be freely chosen. For f_z of approximately 30kHz instead of 5.1kHz, the noise can likely be reduced to roughly 45 mA_{RMS} .
2. Secondly, to truly reach excellent noise performance, reading out the coil current instead of voltage should be investigated. This way, a considerably larger sensitivity can be obtained, which mainly relaxes the noise requirements of the various amplifiers. Besides, the coil resistance itself can be used to implement the integrating transfer, and no additional (large) resistance, contributing a significant amount of noise, has to be added in the signal path (e.g., R_p in Figure 7.4). Of course, reading out the coil current comes at the cost of complexity; however, with this approach, state-of-the-art performance is believed to be within reach.
3. Thirdly, it should be noted that the high-frequency path is not the limiting factor for the overall bandwidth. If SystematIC wants to improve bandwidth performance, they should look into increasing the bandwidth of the amplifiers in the common path.
4. Fourthly, if SystematIC does manage to reduce the noise, the spinning ripple should also be reduced further. Of course, this can be done by lowering the summing filter cross-over frequency f_x . However, area-wise, it might be beneficial to look into changing the low-frequency path and add, for example, ripple reduction loops [13].
5. Lastly, the most drastic change would be to investigate a fully digital implementation of the signal path. To limit the scope of this thesis, this was not considered until now. This solution most likely has its own challenges that should be carefully investigated. However, mainly in terms of area (due to the large required passive components to implement large time constants), a fully digital implementation could be beneficial.

²²Of course, the coil-Hall performance still needs to be verified by measurements, for this purpose, a test plan is given in Appendix A

²³ 77 mA_{RMS} was achieved, while the optimum noise level associated with the transducers would have been 20 mA_{RMS}

Specification	Hall-Hall	coil-Hall
Integrated Noise	180mA _{RMS}	77mA _{RMS}
Bandwidth	318kHz	2MHz ²⁴
Input Range	±80A	±80A
Gain Stability (Temperature, LFP)	-1.5% to 2.6%	-1.5% to 2.6%
Gain Stability (Temperature, HFP)	-1.5% to 2.6%	-2.3% to 0.2%
Gain Stability Lifetime (LFP)	±3%	±3%
Gain Stability Lifetime (HFP)	±3%	-
Gain Flatness over Frequency	±3% ²⁵	<±1%
Offset	± 70mA (3σ)	± 70mA (3σ)
Spinning Ripple	80mA _{RMS}	30mA _{RMS}
Supply Current	20mA	20mA
Chip Area	4mm ²	4mm ²

Table 7.4: Comparison between SystematIC's previous Hall-Hall design and the new coil-Hall design obtained specifications.

Regardless of SystematIC's future decisions, the hybrid-hybrid architecture is believed to remain highly effective. Namely, not having to deal with the various offset-related challenges associated with the other architectures is considered worth the slight power penalty. Furthermore, this architecture does not require navigating significant trade-offs, making it highly scalable for achieving even better performance in the future.

²⁴Limited by amplifiers in the common path (not the high-frequency path). The high-frequency path bandwidth is significantly larger than 10MHz

²⁵In the initial design, SystematIC had an additional pole in the low-frequency path

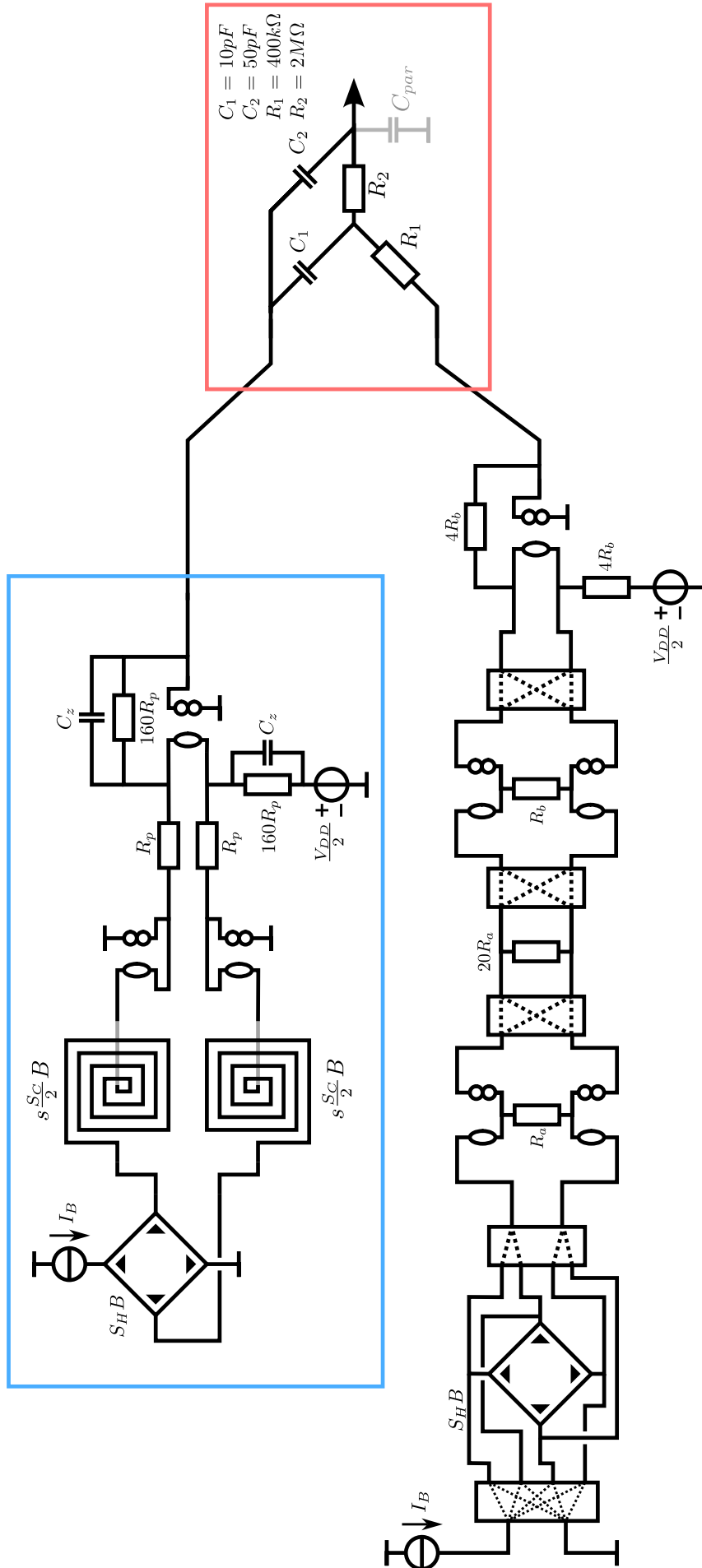


Figure 7.16: System overview including the low- and high-frequency path. The low-frequency path was not designed in this thesis and was completely reused from a previous design of SystematIC. It is, therefore, not further explained in this thesis.

8

CONCLUSION

Throughout this thesis, the design of a coil-Hall hybrid current sensor has been thoroughly investigated. This CMOS-compatible multi-path sensor utilizes a Hall plate for low-frequency sensing and a coil for high-frequency sensing. This enables excellent noise and bandwidth performance. However, as examined in Chapters 3 to 6, implementing this current sensor involves various challenges and trade-offs related to combining, calibrating, and stabilizing the two distinct signal paths. For each of these design aspects, multiple solutions could be employed, each with its clear advantages and disadvantages.

One of the primary goals of this thesis was to provide a comprehensive overview of these challenges and solutions. This is believed to have been successfully accomplished. After reading this thesis, the reader should have obtained sufficient insight to determine their own approach to designing a coil-Hall hybrid current sensor and navigate the various trade-offs according to their own requirements.

The second goal of this thesis was to apply this insight to design a coil-Hall hybrid current sensor tailored to the needs of SystematIC Design B.V. This design was intended as a prototype to demonstrate a significant performance improvement compared to SystematIC's previous Hall-Hall design. For this prototype, not necessarily state-of-the-art performance, but simplicity was highly valued. However, as extensively discussed in Chapter 3, combining the two signal paths involves managing tight trade-offs related to noise, ripple, offset, and gain flatness, which, until now, has led to various complex system solutions.

To address this, the so-called "hybrid-hybrid" architecture was developed. Instead of only a coil, this architecture uses both a coil and a Hall plate in the high-frequency path, adding an additional degree of freedom to the design at the cost of a slight power penalty. This added design flexibility helps alleviate several tight trade-offs and considerably simplifies the overall implementation. In Chapter 7, this was verified by presenting an implementation that could, despite its simplicity, meet the specifications imposed by SystematIC. Moreover, due to its simplicity, this innovative hybrid-hybrid architecture is expected to pave the way for easier attainment of improved performance in the future.



PROTOTYPE CHIP TEST PLAN

As the prototype chip for SystematIC has not been manufactured in time, no measurement results have been presented in this thesis. However, it is still good to present a test plan to make sure all relevant design aspects get investigated and verified once the chip returns. This test plan mainly focuses on the various aspects related to the new coil-Hall architecture and not necessarily on the already verified aspects of SystematIC's previous Hall-Hall design. These relevant aspects mainly include noise, ripple, bandwidth, gain flatness over frequency, temperature drift, and lifetime drift.

A.1. CHIP STARTUP & ENTERING TEST MODE

First of all, it should be verified if the chip starts up properly and is not broken somehow. This can be done as follows:

1. Solder the chip package to a test PCB / breakout board previously used for verification of SystematIC's Hall-Hall product (the same package is used for this new design).
2. Connect the power supply and set its voltage to 5V (nominal voltage) and limit the current to 24mA.
3. Turn on the power supply and check if the supply hits its current limit. If everything functions properly, the chip should consume around 18 - 22 mA.
4. Check the voltages of the various (output) pins (see Table A.1). Both the output pin (VOUT) and the analog reference pin (VOCM) should be at $2.5V (\frac{V_{DD}}{2})$. The fault detection pin (FAULT) should be at 5V (V_{DD}).

If the chip functions properly, the next step is to enter the so-called chip test mode. In this mode, the chip can be programmed via the FAULT pin (e.g., the gain settings can be trimmed). Besides, various on-chip test buses can be connected to the VOCM pin to monitor certain internal analog and digital voltages. The exact procedure on how to enter this test mode is extensively described in some (confidential) documents of SystematIC.

pin	Function
FAULT	Serial Communication IO
GND	Main Ground Connection
VDD	Main Supply Connection
VOUT	Analog Output
VOCM	Analog Reference Output / Testbus Connection
IP+, IP-	Current Rail Connections

Table A.1: The seven package pins and their functionality in system test mode

A.2. HALL PLATE CALIBRATION & MATCHING

Next, before being able to verify the chip's various performance aspects, the low- and high-frequency gain should be calibrated. In this first step, the low-frequency Hall plate and the high-frequency Hall plate will be trimmed. The Hall plate sensitivity can be trimmed by changing the bias current. In this test, it will also be verified if the low- and high-frequency Hall plate can be trimmed in a single step. The procedure can be as follows:

1. Connect the VOCM pin to a multimeter
2. Program the chip such that the output of the high-frequency path is connected to VOCM
3. Measure the voltage at VOCM and check if this voltage is within 0.5V and 4.6V ($\frac{V_{DD}}{2} - 0.4V$). If this is not the case, the output might clip when an input signal is applied, and another chip should be tested¹.
4. Connect the VOUT pin to a multimeter
5. Measure the voltage at VOUT. This signal will be referred to as the quiescent output voltage.
6. Apply a well-known DC current to the current rail. This can, for example, be generated by applying a well-known voltage (using a power supply) across a well-known resistance (power resistor). Ideally, this signal should be as large as possible, but most importantly, the signal should be accurate. To be sure, it could be measured with another multimeter.
7. Measure the chip output voltage at VOUT with the multimeter. If need be, average the signal several times such that an SNR of at least 40dB is obtained.
8. Subtract the voltage obtained in step 6 from the voltage in step 7 and divide it by the input current. This is the measured sensitivity.
9. Trim both the low- and high-frequency Hall plate bias current simultaneously by programming the gain settings via the FAULT pin. If the sensitivity is below 25mV/A, increase the bias current. If the sensitivity is above 25mV/A, decrease the bias current.
10. Measure the output voltage again and repeat steps 6 - 9 until the sensitivity is trimmed within 1% of 25mV/A.
11. Connect the VOCM pin to a multimeter
12. Program the chip such that the output of the low-frequency path is connected to VOCM
13. Measure the voltage at VOCM
14. Apply the same DC current as used in step 6 to the current rail
15. Measure the voltage at VOCM, subtract the voltage measured at step 13, and divide by the input current. This is the measured sensitivity at the output of the low-frequency path.
16. Program the chip such that the output of the high-frequency path is connected to VOCM
17. Apply the same DC current as used in step 6 to the current rail
18. Measure the voltage at VOCM, subtract the voltage measured at step 13, and divide by the input current. This is the measured sensitivity at the output of the high-frequency path.
19. Compare the sensitivities obtained in steps 15 and 18 and check if they are within 1% of each other. If this is the case, it is an indication that the low- and high-frequency Hall plates can be calibrated in one step². If this is not the case, repeat the calibration steps (7 - 11) purely for the high-frequency Hall plate.

A.3. HIGH-FREQUENCY GAIN TRIMMING

Next, the high-frequency (coil) gain should be trimmed. This can be done using a variable capacitor bank in the high-frequency path (i.e., by changing C_z in Figure 7.4). Calibrating the high-frequency gain should be performed at a frequency well above the cross-over frequency f_x . Generating such a signal and accurately measuring it is less straightforward compared to a DC signal. Again, to obtain sufficient resolution, ideally, a large input signal is applied. However, if this is not possible, the signal can be averaged over a long period of time.

Currently, SystematIC does not have any specialized equipment to generate high-frequency, high-current sinusoidal signals. However, a PCB has been designed that can apply short and accurate 40A pulses (with a rise time of less than 100ns) to the current rail. The pulse duration should be considerably shorter than the time constant related to the cross-over frequency (roughly 150μs) but significantly longer than the time constant related to the system bandwidth (roughly 80ns). This way, the low-frequency path does not react, but the high-frequency path can fully settle and its gain can be measured. This can be achieved with the PCB

¹The chances of this happening should be well below 1% as indicated in Section 7.6.4

²Of course this should be validated by measuring more chips

used by SystematIC. Using this PCB the procedure can be as follows:

1. Connect the VOUT pin to an oscilloscope
2. Connect the current rail to the custom PCB
3. Apply the 40A current pulse
4. Measure the pulse height with the oscilloscope (this probably only has to be done once, as a 40A signal should result in a considerably larger SNR than 40dB)
5. Divide this voltage by the 40A input current. This is the high-frequency gain.
6. Trim the capacitor bank by programming the gain setting via the FAULT pin. If the sensitivity is below 25mV/A, decrease the capacitance. If the sensitivity is above 25mV/A, increase the capacitance.
7. Measure the output voltage again and repeat steps 3 - 6 until the sensitivity is trimmed within 1% of 25mV/A.

A.4. BANDWIDTH MEASUREMENT

Having trimmed the low- and high-frequency gain, various performance aspects can be validated. First of all, the system bandwidth can be measured. For this purpose, the aforementioned PCB was initially designed. Namely, the settling time of such a pulse is an accurate measure of the bandwidth. The test procedure can be as follows:

1. Connect the VOUT pin to an oscilloscope
2. Connect the current rail to the custom PCB
3. Apply the 40A current pulse
4. From the oscilloscope plot, measure how long it takes for the output signal to settle within 5% of the final value. This is equal to three time constants τ . The bandwidth can be calculated according to $BW = \frac{1}{2\pi\tau}$.

A.5. NOISE PERFORMANCE & SPINNING RIPPLE

Besides a large bandwidth, excellent noise performance is also one of the main benefits of the coil-Hall architecture. In this design, 80mA_{RMS} was strived for. This should be verified. The total RMS noise can be measured with a wide range of instruments. However, to also see the noise spectrum, a spectrum analyzer would be a good choice. In this way, the tones associated with the Hall plate spinning ripple can also be measured. This can be done as follows:

1. Connect the spectrum analyzer to VOUT
2. Disable Hall plate spinning in the low-frequency path
3. Measure the output spectrum and the (AC) RMS value of the output signal
4. Divide this RMS signal by the system gain to obtain the RMS noise signal
5. Enable Hall plate spinning again
6. Measure the spectrum and verify that the tones associated with spinning are well below the integrated noise level
7. Also measure the RMS value of the signal and see if it has not significantly increased compared to the noise measurement
8. Repeat this measurement for various other samples to get an estimate on how much the noise and ripple differs from sample to sample (also gain calibration should be done for these samples)

A.6. FREQUENCY RESPONSE MEASUREMENT & GAIN FLATNESS

Besides bandwidth and noise performance, SystematIC has also specified certain limits to gain inaccuracy; these will be validated in the last few tests. First of all, SystematIC has specified that the gain is maximally allowed to deviate with $\pm 1\%$ from the ideal gain over frequency. To verify this, the frequency response must be measured with more than 99% accuracy. Similar to high-frequency calibration and bandwidth verification, this requires a high-frequency and, ideally, high-current signal. However, instead of pulses, in this case, sinusoidal signals at several frequencies must be applied. The previously mentioned PCB cannot be used for this purpose. If SystematIC wants to do such a measurement, they should invest in a better measurement setup.

It is believed that the most cost-effective solution would be to design a custom PCB containing a buffer with

high current drive capability³ (ideally a few ampère). Moreover, the output of this buffer can be connected to the current rail in series with an accurate (power) resistor. By also applying an accurate voltage to the buffer input, an accurate current is generated⁴. A measurement procedure could, for example, be as follows:

1. Connect the VOUT pin to an oscilloscope
2. Connect the current rail in series with an accurate (power) resistor to the buffer output
3. Connect a function generator at the input of the buffer
4. Connect a voltage probe across the power resistor to another oscilloscope channel
5. Using the function generator, apply a sinusoidal voltage signal with a certain frequency f at the buffer input
6. Measure the signal magnitude at VOUT and across the power resistor
7. The signal magnitude measured at VOUT, divided by resistor voltage magnitude, multiplied by the resistance value of the power resistor, is the gain at the frequency f . If the measurement is not sufficiently accurate, several measurements should be averaged.
8. Repeat steps 5 - 7 for multiple frequencies. The signals should range from DC to 2MHz
9. The various gain measurements can be plotted against frequency to obtain the frequency (magnitude) response⁵.

Lastly, it should be noted that if SystematIC decides to invest in such a measurement setup with high-current sinusoidal signals, it could also be used for calibration and bandwidth measurements, as described in the previous two sections.

A.7. GAIN DRIFT OVER TEMPERATURE

Next, gain drift over temperature should be investigated. Namely, as specified by SystematIC, both at low- and high frequencies, the gain should not vary with more than $\pm 3\%$. A similar setup can be used as for the frequency response measurement, however, in this case, it is probably sufficient to measure at only two frequencies (DC and well above the cross-over frequency, for example, at 100kHz). To automate the setup, it is probably best to use a multimeter with AC RMS measurement capabilities (such as, for example, the Agilent 34410A Multimeter available at SystematIC) instead of an oscilloscope. A multimeter could not be used in the frequency response measurement as, at least the multimeters at SystematIC, do not have sufficient bandwidth. A potential test procedure could be as follows:

1. Connect VOUT to the multimeter
2. Connect the current rail in series with an accurate (power) resistor to the buffer output
3. Connect a function generator at the input of the buffer
4. Connect another multimeter across the power resistor
5. Place the test PCB with the chip in a temperature/climate chamber
6. Set the temperature at -40°C
7. Measure the gain at DC and at 100kHz according to a similar procedure presented in the previous section (steps 5 - 7)
8. Repeat these measurements for multiple temperatures up until $+125^{\circ}\text{C}$

A.8. GAIN DRIFT OVER LIFETIME

Lastly, gain drift over lifetime should be measured. For gain drift over lifetime, a very similar procedure as used for gain drift over temperature could be used. However, it takes a significant amount of time and this test should probably be performed last. Especially in this situation, it is important that the measurement setup is automatized as such measurement take quite a long time. The lifetime drift of the high-frequency gain is especially interesting to characterize. This is because no information is available on the drift of the resistor and capacitor required in this gain. Of course, the severity of the lifetime drift of the Hall sensitivity is also interesting to investigate. A potential measurement plan could be as follows:

³Instead of designing such a PCB, SystematIC could also choose to buy an off-the-shelf waveform amplifier (e.g., see the TS250 of Accel Instruments)

⁴Note that a buffer is not necessarily required, and a function generator can be directly connected to an accurate resistance. However, the current, due to the 50 Ω function generator output impedance, is likely quite small, and a lot of measurements need to be averaged

⁵Of course, with this relatively simplistic approach, only some discrete measurement points can be obtained. To obtain a more continuous spectrum, SystematIC could, for example, invest in a network analyzer

1. Connect VOUT to the multimeter
2. Connect the current rail in series with an accurate (power) resistor to the buffer output
3. Connect a function generator at the input of the buffer
4. Connect another multimeter across the power resistor
5. Place the test PCB with the chip in a temperature/climate chamber
6. Set the temperature at roughly 80°C and a humidity at 85%. This was, for example, also done in [36] as this way, moisture absorption of the package can be sped up. This absorption is believed to be one of the main ways in which package stress changes, which, in turn, is believed to be one of the main causes of sensitivity lifetime drift.
7. Automatically measure the gain at DC and 100kHz at a fixed time interval
8. Leave this test running for multiple days (or even weeks)

B

COIL MODEL

To determine the different component values for this model, the rectangular planar coil shown in Figure 3.3 can be divided into a set of closed rectangles as shown in Figure 3.4. The different model parameters can be related to the total coil length (L_{tot}) and the total coil area (A_{tot}). With the coil winding inwards, the length and enclosed area of each rectangle decrease slightly for each additional turn. Therefore, the total length is determined by calculating the perimeter of each rectangle and then summing these for all N rectangles, while the total area is determined by first calculating the area enclosed by each rectangle and then summing these values for all N rectangles. As long as the trace width (w) is not too large compared to the outer dimensions (x and y), this approximation is believed to be accurate.

Following these definitions, the length of the n -th turn can be written as:

$$l_n = 2 \cdot (x - w - 2(n-1)(w+s)) + 2 \cdot (y - w - 2(n-1)(w+s)) \quad (\text{B.1})$$

while the area enclosed by the n -th turn can be written as:

$$A_n = (x - w - 2(n-1)(w+s)) \cdot (y - w - 2(n-1)(w+s)) \quad (\text{B.2})$$

By summing these lengths and areas for all N turns, the total length and area can be written as:

$$L_{tot} = N \cdot (2x + 2y + 4s) - N^2 \cdot (4w + 4s) \quad (\text{B.3})$$

$$A_{tot} = N \cdot [(x + w + 2s)(y + w + 2s) - 2(N+1)(w+s) \cdot [\frac{1}{2}(x+y) + w + 2s - \frac{1}{3}(2N+1)(w+s)]] \quad (\text{B.4})$$

A more in-depth derivation can be found in [39]; the equations listed here are a slightly modified version to allow for rectangular coils instead of only square ones.

After obtaining the total length and area, the model parameters S_C , R_C and C_C can be written as shown by Equations B.5 - B.7.

$$S_C = A_{tot} \quad (\text{B.5})$$

$$R_C = \frac{L_{tot}}{w} \cdot R_{sheet} \quad (\text{B.6})$$

$$C_C = \epsilon_0 \epsilon_{ox} \frac{w L_{tot}}{d} \quad (\text{B.7})$$

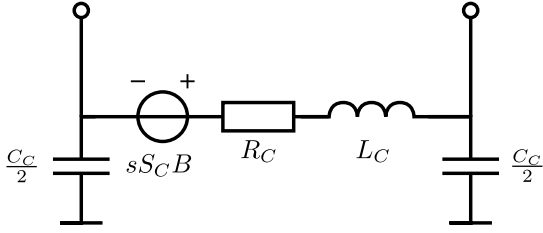


Figure B.1: Electrical model of the coil (the same figure from Section 3.1.2 reshown for the sake of clarity)

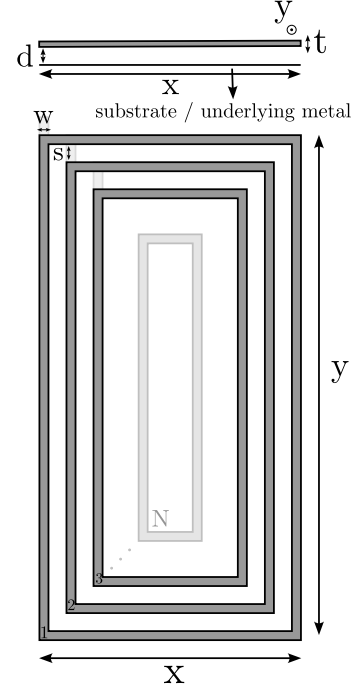


Figure B.2: Coil representation with geometry-related parameter definitions for calculation purposes (the same figure from Section 3.1.2 reshown for the sake of clarity)

In these equations, w is the trace width (in m), R_{sheet} the metal sheet resistance (in Ω per square), ϵ_0 the permittivity of free space ($8.854 \cdot 10^{-12}$ F m $^{-1}$), ϵ_{ox} the relative permittivity of the silicon oxide (roughly 4.2), and d the distance to the underlying oxide or nearby metal layer (in m). R_C is determined using the metal sheet resistance and the total number of squares derived from L_{tot} and w while the capacitance C_C is estimated as a parallel plate capacitance to ground (either the substrate or an overlapping metal layer).

The inductance L_C can be approximated using a modified version of Wheeler's formula [22]. This results in the following equation:

$$L_C = K_1 \mu_0 \frac{N^2 d_{avg}}{1 + K_2 \rho} \quad (\text{B.8})$$

In this equation, K_1 and K_2 are geometry-dependent constants (of 2.34 and 2.75), μ_0 is the permeability of free space ($4\pi \cdot 10^{-7}$ H m $^{-1}$), while d_{avg} and ρ are the so-called average diameter and fill-factor respectively [22]. The last two parameters are based on the outer and inner diameter of the coil according to $d_{avg} = \frac{1}{2}(d_{out} - d_{in})$ and $\rho = (d_{out} - d_{in}) / (d_{out} + d_{in})$ with $d_{out} = x$ and $d_{in} = x - 2(N - 1)(w + s) - 2w$. [22] states that these equations have a typical error of around 1 - 2% compared to a field solver. However, the equation was derived for a square coil; for a rectangular coil, this equation might produce less accurate results.

B.1. PLOTS

Using these equations, various plots can be generated to help dimension the coil more accurately and obtain the model parameters; these plots include:

- Coil Sensitivity S_C
- Coil Resistance R_C
- Coil Capacitance C_C
- Coil Inductance L_C
- Coil Bandwidth Estimate¹ $\frac{1}{2\pi R_C C_C}$
- Coil Bandwidth Estimate $\frac{R_C}{2\pi L_C}$
- Coil Noise Factor $\frac{S_C}{\sqrt{R_C}}$
- Coil Current Sensitivity $\frac{S_C}{R_C}$

as a function of the number of coil turns N for different outer coil dimensions x and y ($x = y$) and wire widths w .

To generate these plots, the following parameter values have been used:

- Spacing s of 0.28 μm
- Distance to the substrate / nearby metal layer d of 800 nm
- Sheet Resistance R_{sheet} of 80 m Ω per square

All other parameter values have been given in the previous section. Also, note that each curve ends for a different amount of turns N ; this is limited by the number of turns with width w that can fit in the area limited by x and y ; the larger w , or the smaller x or y , the fewer turns will fit.

¹Based on simulation data, the actual bandwidth for a distributed network is likely higher

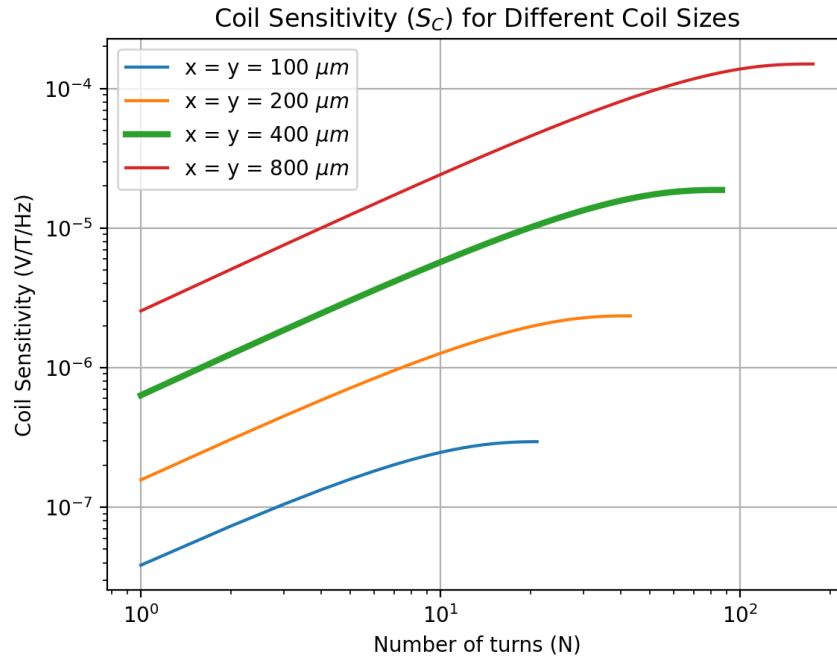


Figure B.3: Coil sensitivity (S_C) for different coil sizes and number of turns, $w = 2 \mu m$ (calculated with the equations from Section 3.1.2)

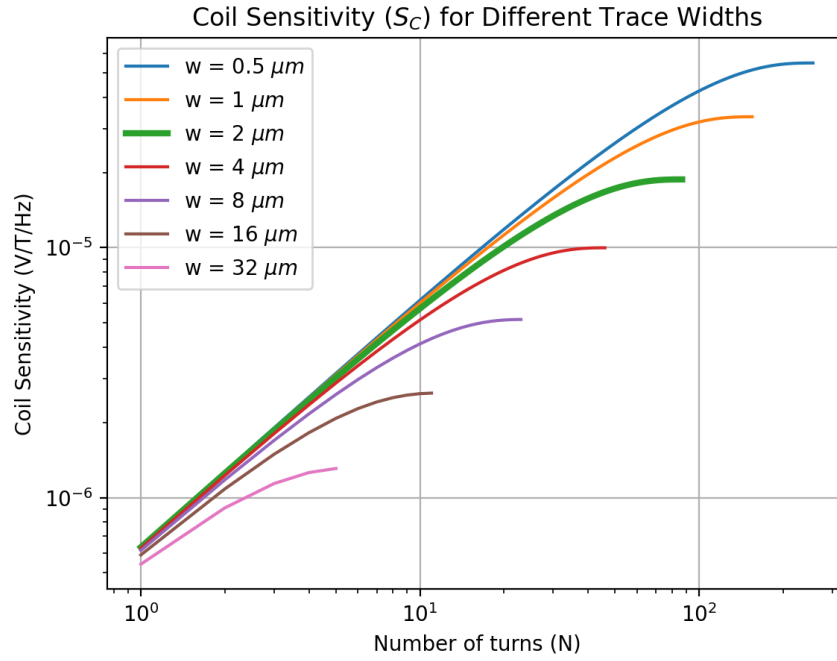


Figure B.4: Coil sensitivity (S_C) for different coil widths and number of turns, $x = y = 400 \mu m$ (calculated with the equations from Section 3.1.2)

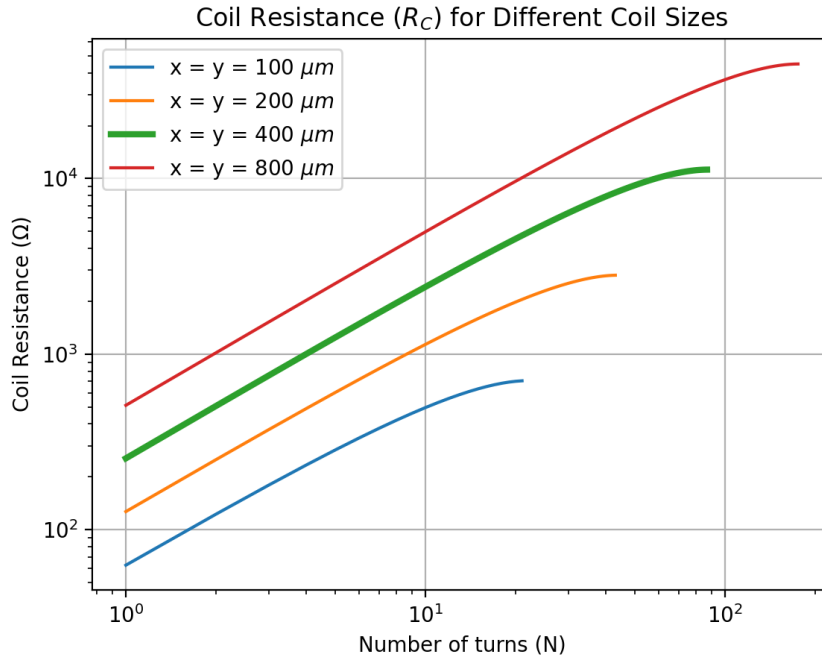


Figure B.5: Coil resistance (R_C) for different coil sizes and number of turns, $w = 2 \mu m$ (calculated with the equations from Section 3.1.2)

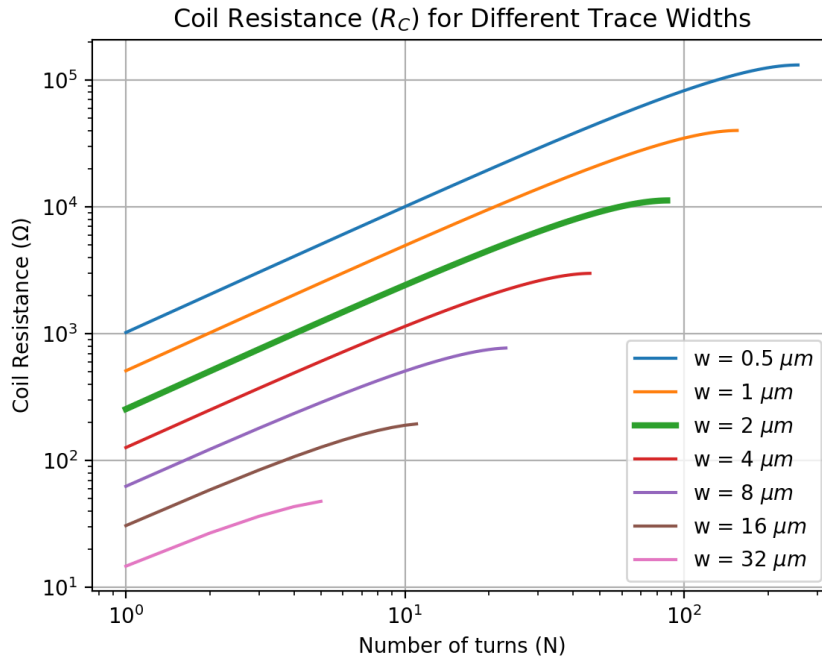


Figure B.6: Coil resistance (R_C) for different coil widths and number of turns, $x = y = 400 \mu m$ (calculated with the equations from Section 3.1.2)

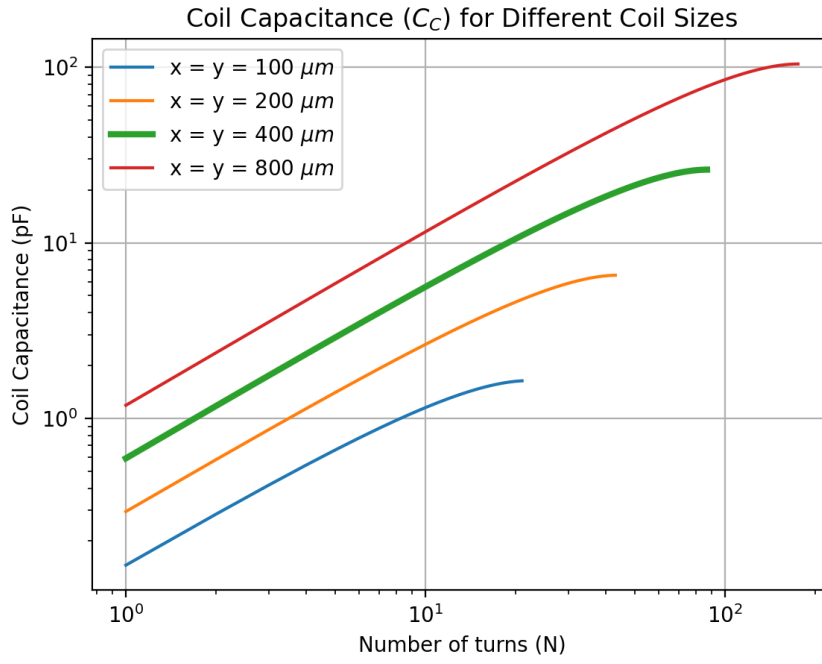


Figure B.7: Coil capacitance (C_C) for different coil sizes and number of turns, $w = 2 \mu m$ (calculated with the equations from Section 3.1.2)

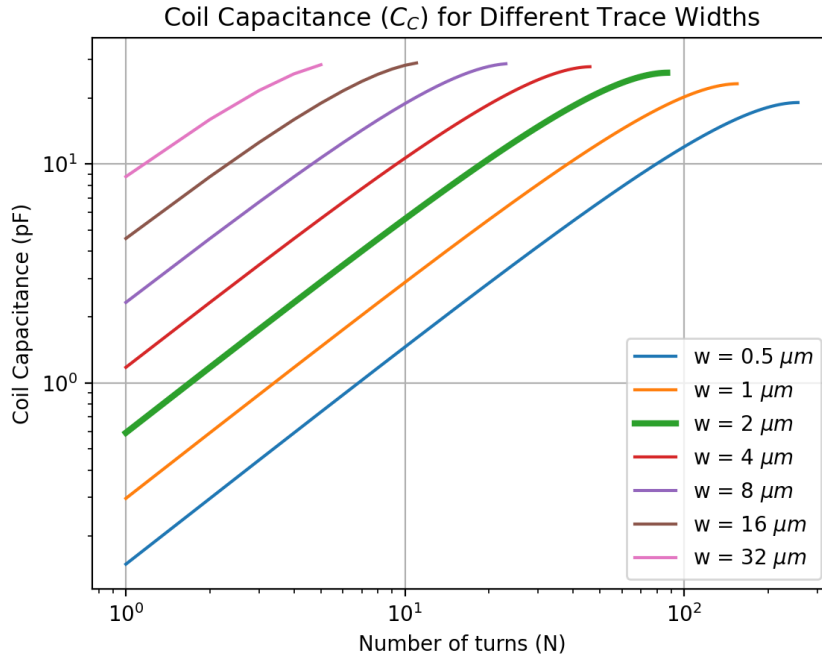


Figure B.8: Coil capacitance (C_C) for different coil widths and number of turns, $x = y = 400 \mu m$ (calculated with the equations from Section 3.1.2)

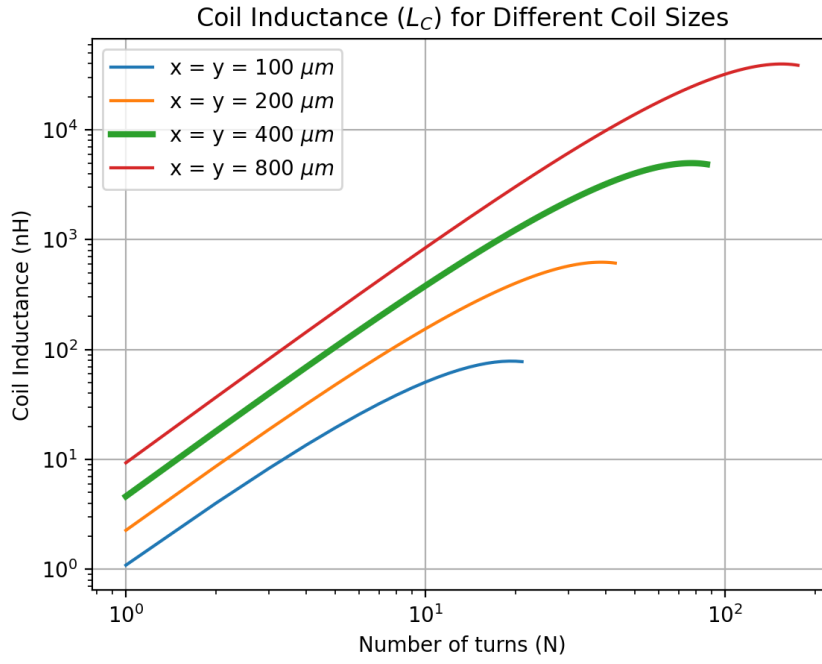


Figure B.9: Coil inductance (L_C) for different coil sizes and number of turns, $w = 2 \mu m$ (calculated with the equations from Section 3.1.2)

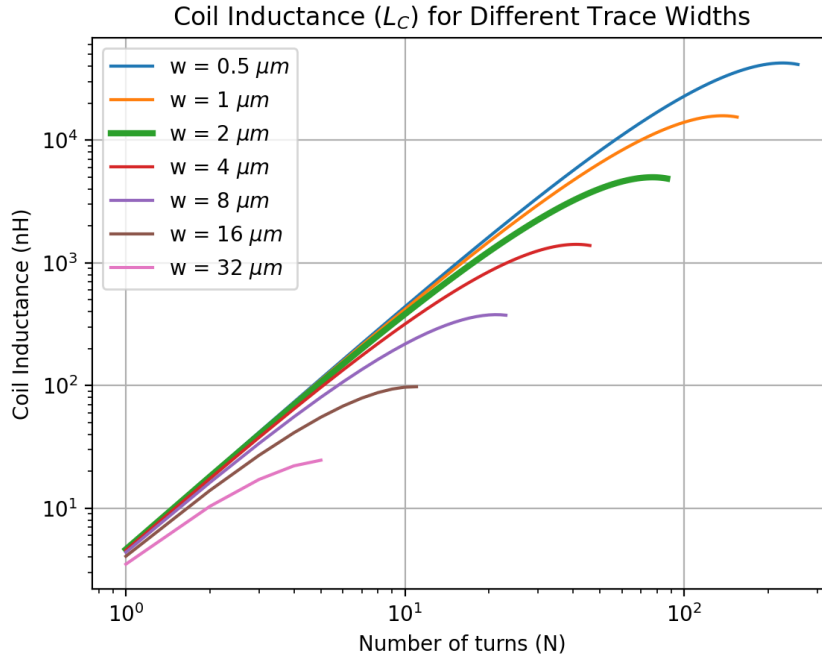


Figure B.10: Coil inductance (L_C) for different coil widths and number of turns, $x = y = 400 \mu m$ (calculated with the equations from Section 3.1.2)

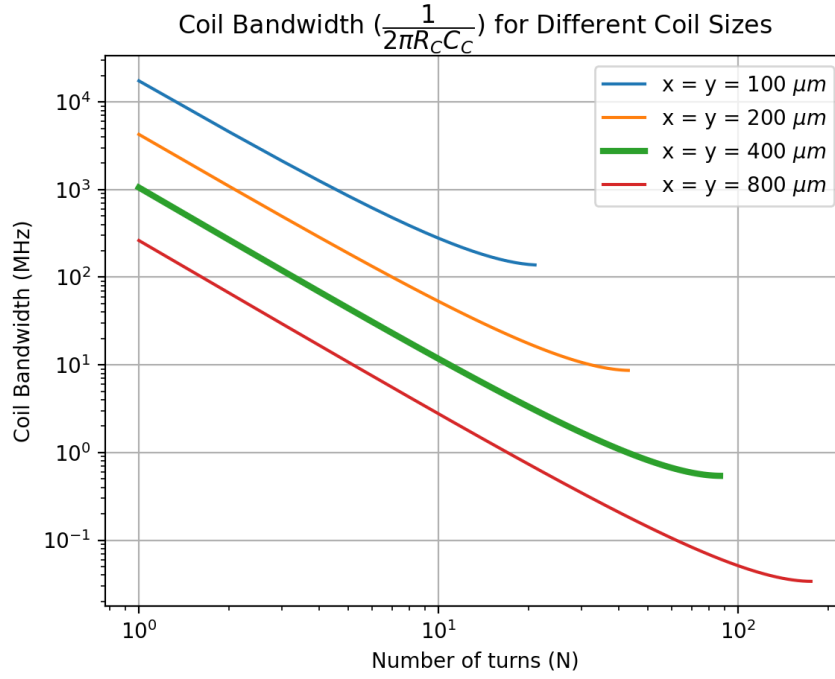


Figure B.11: Coil bandwidth estimate ($\frac{1}{2\pi R_C C_C}$) for different coil sizes and number of turns, $w = 2 \mu m$ (calculated with the equations from Section 3.1.2)

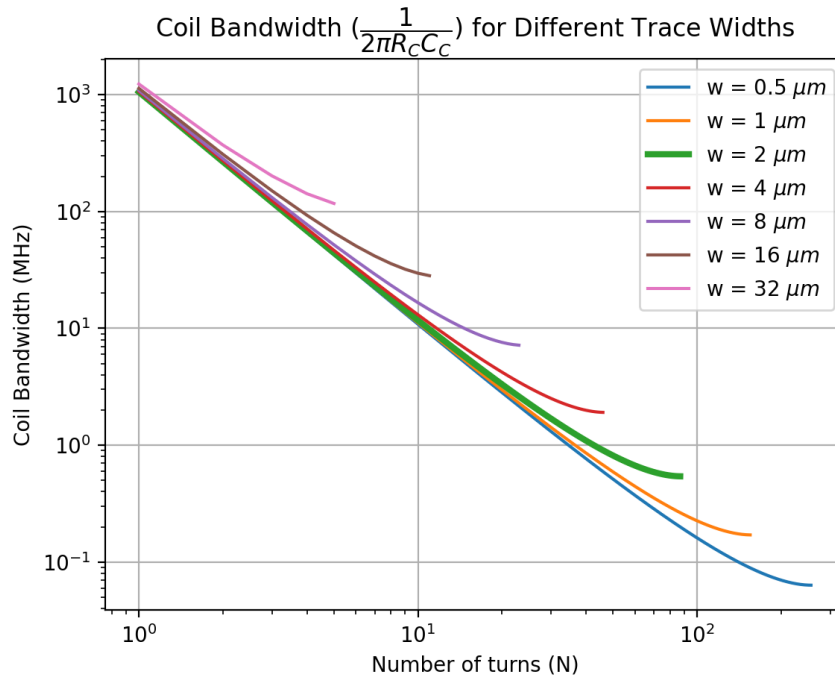


Figure B.12: Coil bandwidth estimate ($\frac{1}{2\pi R_C C_C}$) for different coil widths and number of turns, $x = y = 400 \mu m$ (calculated with the equations from Section 3.1.2)

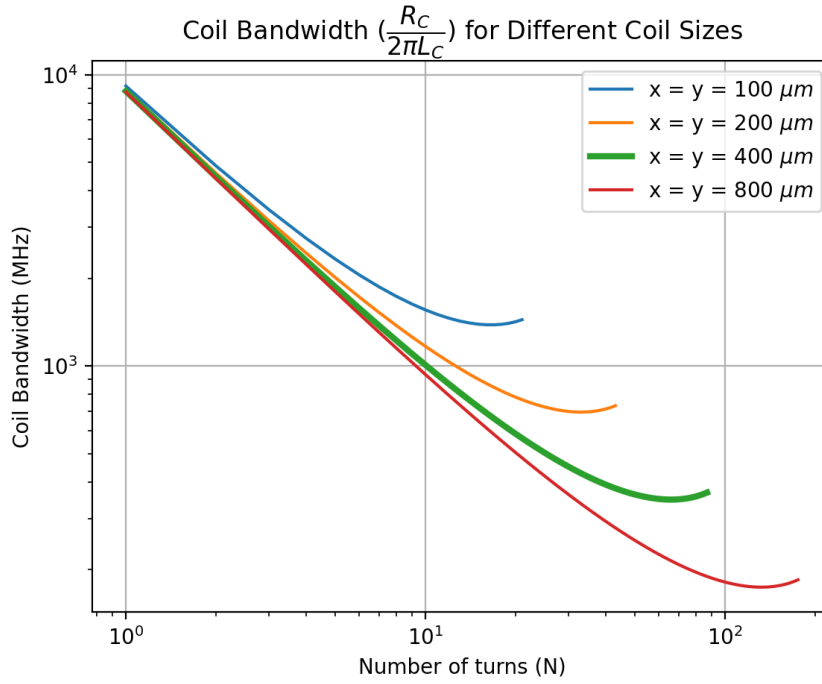


Figure B.13: Coil bandwidth estimate ($\frac{R_C}{2\pi L_C}$) for different coil sizes and number of turns, $w = 2 \mu m$ (calculated with the equations from Section 3.1.2)

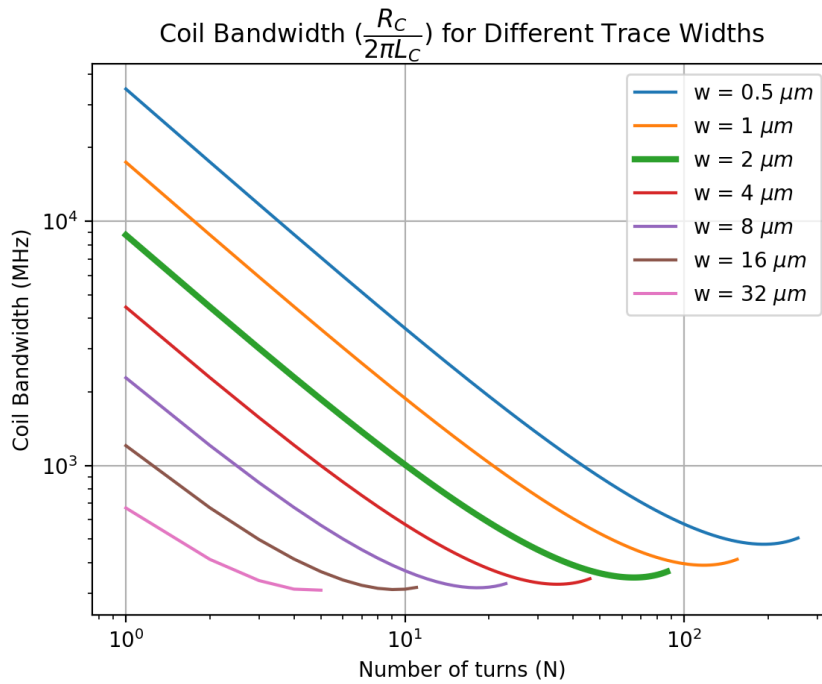


Figure B.14: Coil bandwidth estimate ($\frac{R_C}{2\pi L_C}$) for different coil widths and number of turns, $x = y = 400 \mu m$ (calculated with the equations from Section 3.1.2)

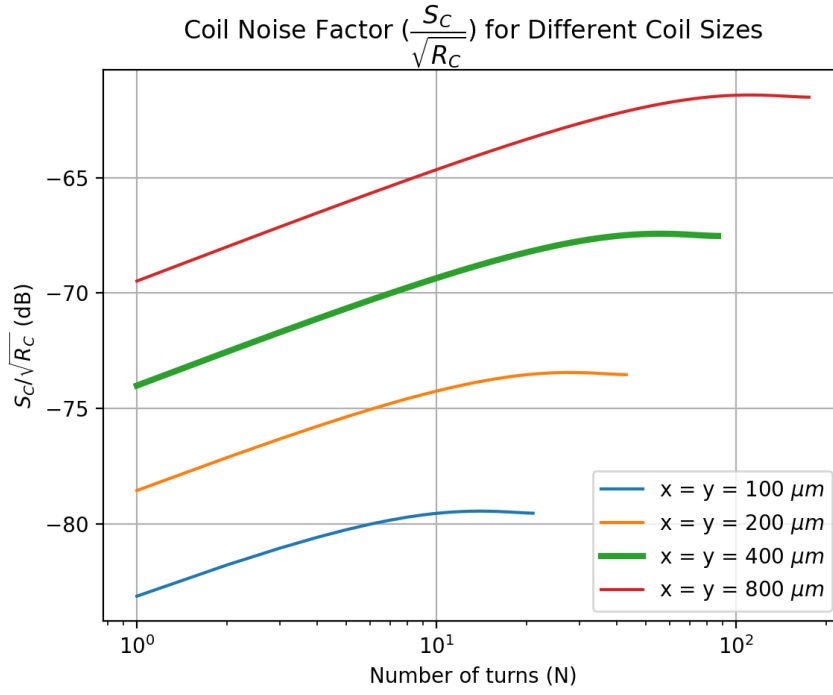


Figure B.15: Coil noise factor ($\frac{S_C}{\sqrt{R_C}}$) for different coil sizes and number of turns, $w = 2 \mu m$ (calculated with the equations from Section 3.1.2)

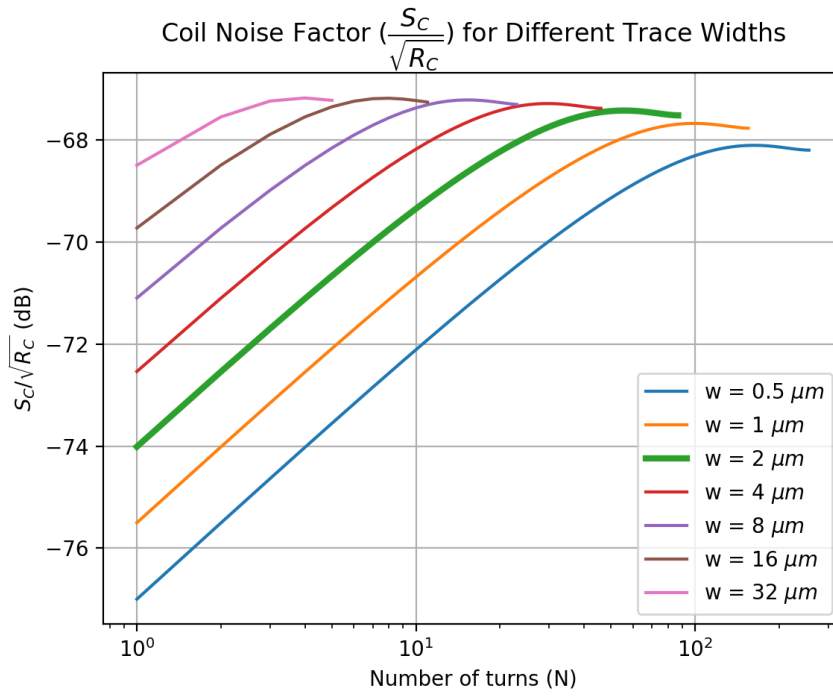


Figure B.16: Coil noise factor ($\frac{S_C}{\sqrt{R_C}}$) for different coil widths and number of turns, $x = y = 400 \mu m$ (calculated with the equations from Section 3.1.2)

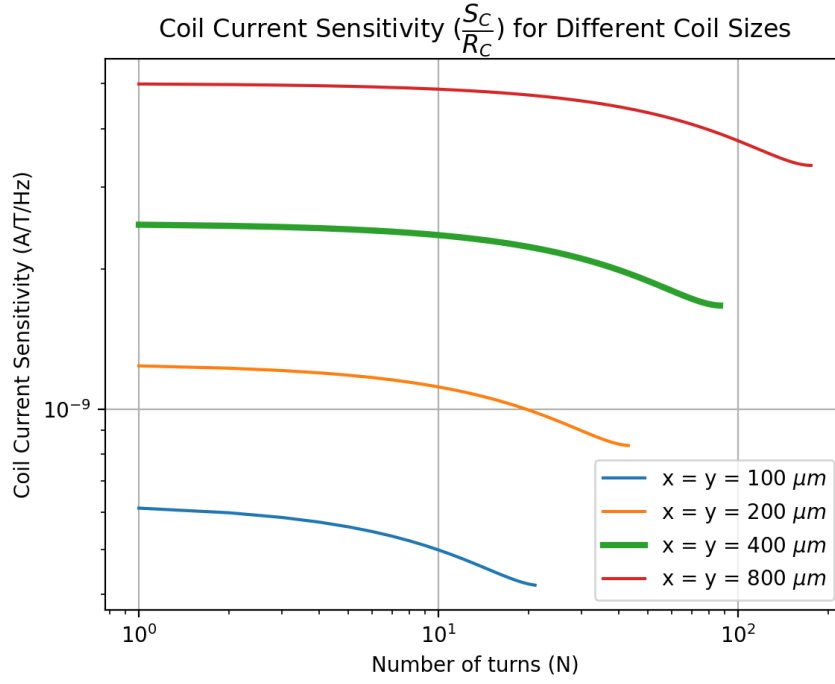


Figure B.17: Coil current sensitivity ($\frac{S_C}{R_C}$) for different coil sizes and number of turns, $w = 2\mu m$ (calculated with the equations from Section 3.1.2)

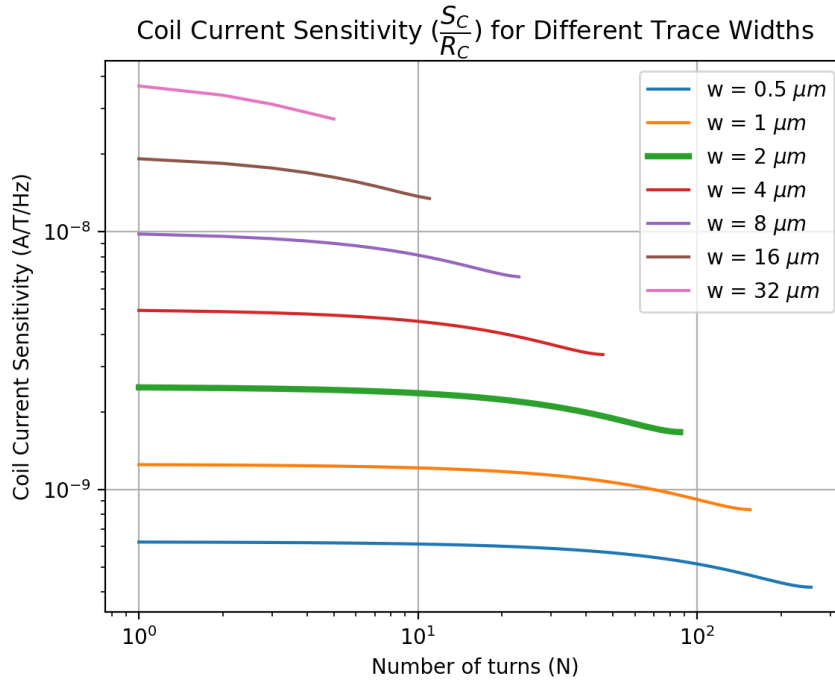


Figure B.18: Coil current sensitivity ($\frac{S_C}{R_C}$) for different coil widths and number of turns, $x = y = 400\mu m$ (calculated with the equations from Section 3.1.2)

C

DIMENSIONING OF THE COIL

As explained in Section 3.1.2, the coil can be modeled by means of the four parameters S_C , R_C , C_C , and L_C . The resistance R_C limits the coil noise performance, while the reactive elements C_C and L_C (in combination with R_C) limit the inherent coil bandwidth. On the other hand, a larger sensitivity results in better noise performance and also introduces benefits in various trade-offs related to the signal path combining. Therefore, ideally, the sensitivity S_C is made as large as possible for small values of the coil resistance, capacitance, and inductance.

In this appendix, a strategy will be devised for dimensioning the coil for optimal performance—that is, optimizing the sensitivity for the best inherent noise and bandwidth performance. Besides, as explained in Section 3.6, bandwidth limitations, drive capability, and signal handling constraints are quite different for reading out either voltage or current. As a result, they require different dimensioning strategies. These will be discussed below.

C.1. VOLTAGE READOUT

As mentioned in Section 3.6.1, the coil voltage can become extremely large for high-speed, high-amplitude signals. The maximum voltage level the readout circuitry can handle is determined by the input range of the subsequent amplifier (commonly limited by the supply) or, more strictly, by the breakdown voltage of the transistor.

Suppose the system will be designed to handle coil voltages of maximally $\pm V_{max}$ and is rated for an input range of $\pm I_{max}$ and a maximum frequency of f_{max} . For this system to handle a full-scale sinusoidal signal $I(t) = I_{max} \cdot \sin(2\pi f_{max} t)$, the maximum coil sensitivity should be limited to:

$$S_{C,max} = \frac{V_{max}}{A_{co} \cdot 2\pi f_{max} \cdot I_{max}} \quad (C.1)$$

For some realistic values of $f_{max} = 2\text{MHz}$, $I_{max} = 80\text{A}$, $V_{max} = \pm 1\text{V}$ and $A_{co} = 400\text{ }\mu\text{T A}^{-1}$, S_C should be limited to roughly $2.5 \cdot 10^{-6} \text{ VT}^{-1} \text{ Hz}^{-1}$. This is a relatively small value and can generally be obtained for a small coil with more than sufficient bandwidth remaining¹. In other words, unless one aims for an extremely high bandwidth, this sensitivity limit is likely lower than the limit set by the inherent coil bandwidth. This low sensitivity limits system performance regarding noise, but also in various other aspects, as indicated in Chapter 4.

One could consider neglecting the large coil signals and specify the full-power bandwidth differently from the small-signal bandwidth. However, as explained in Section 3.6.1, one should be aware that this also increases the chances of the system failing to respond quickly to a short circuit event².

¹for example, a $200 \times 200\text{ }\mu\text{m}$ coil with 15 turns can reach this sensitivity with a bandwidth of at least 50 MHz remaining, see, for example, Figure B.3 and B.11 in Appendix B

²The minimum rise-time t_r the system can handle (for a $-I_{max}$ to $+I_{max}$ step) is related to the chosen f_{max} by $t_r = \frac{1}{\pi f_{max}}$. $f_{max} = 2\text{ MHz}$ corresponds to $t_r \approx 160\text{ ns}$

Still, whether or not the sensitivity is chosen to be limited by the aforementioned constraint, it is educational to investigate how to properly dimension the coil. For this purpose, writing the different model parameters as a set of proportionalities related to coil dimensions is deemed more insightful than blindly following the relatively complex equations derived in Appendix B. Still, these more complex equations are valuable for verifying the intuition and obtaining accurate model parameters. For the interested reader, plots of various coil-related parameters have been given for some typical coil dimensions in Appendix B.

By assuming that y is scaled proportional to x (i.e., $y \propto x$), the sensitivity (being related to the total coil area), the resistance (being related to the total coil length), and the capacitance (also related to the total length) can be written as shown below:

$$S_C \propto x^2 \cdot N \quad (\text{C.2})$$

$$R_C \propto \frac{x \cdot N}{w} \quad (\text{C.3})$$

$$C_C \propto x \cdot N \cdot w \quad (\text{C.4})$$

Furthermore, as the inherent coil bandwidth is commonly limited by the coil resistance and capacitance³. The bandwidth limitation can be written as:

$$BW_{voltage} \propto \frac{1}{R_C C_C} \propto \frac{1}{x^2 \cdot N^2} \quad (\text{C.5})$$

Lastly, the inherent noise performance of the coil can be characterized by the factor $\frac{S_C}{\sqrt{R_C}}$, which, according to equation 3.11, should be maximized. This factor essentially indicates the minimum noise the coil path contributes if the readout noise figure would be 0 dB. Using Equations C.2 and C.3, this factor can be written as:

$$\frac{S_C}{\sqrt{R_C}} \propto x^{\frac{3}{2}} \cdot N^{\frac{1}{2}} \cdot w^{\frac{1}{2}} \quad (\text{C.6})$$

These proportionalities give valuable insight. Most importantly, it can be seen that the sensitivity can only be increased by increasing the outer dimensions x or the number of turns N . However, the bandwidth is inversely proportional to the square of these parameters. This results in the trade-off between sensitivity and bandwidth. Still, Equation C.2 indicates that the sensitivity scales fastest with larger values of x while the bandwidth in Equation C.5 decreases at the same rate for both x and N . Therefore, to maximize the sensitivity for a given bandwidth, the outer dimensions (x and y) should be as large as possible. These outer dimensions are either limited by the available chip area or by the size of the area to which the magnetic field is concentrated. This strategy also causes the inherent noise factor from Equation C.6 to increase the fastest.

After maximizing x , the only way to increase the sensitivity further is by increasing N . However, as this causes the bandwidth to decrease drastically, this can only be done up to a certain point. Moreover, sufficient margin must be taken into consideration to account for process variation and additional bandwidth limitation caused by the subsequent amplifier's input capacitance. If the coil is not fully filled with turns at this point (either limited by bandwidth or by the maximum coil signal constraint), the trace width can be increased to reduce the coil's resistance. This does, to a first-order approximation, not affect the sensitivity and bandwidth, but does increase the fundamental noise factor $\frac{S_C}{\sqrt{R_C}}$. Reducing the resistance without the decrease in sensitivity always helps achieve better noise performance. However, the benefit is only marginal if the readout noise is not also decreased simultaneously; only a larger sensitivity can soften the readout noise requirements

³This can, for example, be deduced by comparing Figures B.11 and B.13 in Appendix B; the bandwidth estimate based on R_C and C_C becomes significantly smaller for larger coil dimensions compared to the bandwidth estimate based on R_C and L_C

When dimensioning the coil, there is a trade-off between
sensitivity, bandwidth and area

C.2. CURRENT READOUT

The story is slightly different when reading out the coil current. Most importantly, the coil signal amplitude constraint is now determined by the maximum coil current instead of voltage⁴. The maximum current the system can handle is set by the subsequent amplifier's current drive capability ($\pm I_{drive}$). This can be increased by proper design and is not limited by the strict transistor breakdown constraint. Similar to Equation C.1, the following limit can be devised:

$$\left. \frac{S_C}{R_C} \right|_{max} = \frac{I_{drive}}{A_{co} \cdot 2\pi f_{max} \cdot I_{max}} \quad (C.7)$$

Besides having more flexibility in choosing I_{drive} compared to V_{max} , another advantage is that the factor $\frac{S_C}{R_C}$ should be limited instead of S_C . Because the coil resistance and sensitivity both increase for larger coil dimensions, this factor can remain relatively small for large values of S_C . The following proportionality captures this:

$$\frac{S_C}{R_C} \propto x \cdot w \quad (C.8)$$

While S_C increases with x^2 and N , $\frac{S_C}{R_C}$ only increases with x and w . Therefore, with a small trace width, a large number of turns, and large outer dimensions, a large sensitivity can be obtained for a small value of $\frac{S_C}{R_C}$. For example, the same values of f_{max} , I_{max} and A_{co} used in the previous section in combination with a moderate drive capability of $I_{drive} = \pm 1$ mA result in a maximum $\frac{S_C}{R_C}$ of $2.5 \cdot 10^{-9}$ A T⁻¹ rad⁻¹ s. For a fully filled coil of 400x400 μ m with a trace width of 1 μ m (roughly 100 turns), a sensitivity of $4 \cdot 10^{-5}$ V T⁻¹ rad⁻¹ s can be obtained while $\frac{S_C}{R_C}$ stays well below $2.5 \cdot 10^{-9}$ A T⁻¹ rad⁻¹ s (see the orange curves in Figures B.18 and B.4 in Appendix B). This sensitivity is already an order of magnitude larger than the limit for voltage readout ($2.5 \cdot 10^{-6}$ V T⁻¹ rad⁻¹ s). Besides, the current drive capability can be relatively easily increased. This means that the chance of clipping can also be even further reduced.

Lastly, the coil bandwidth should be considered. For voltage readout, it became evident that the bandwidth quickly drops for a large sensitivity. In that case, the coil resistance and capacitance limit the bandwidth. However, as explained in Section 3.6.2, when reading out the coil current, the capacitance is (to a first-order approximation) shorted to ground, and the coil resistance and inductance now determine the bandwidth. The relevant proportionalities could be written for this inductance-based bandwidth limitation. However, as it turns out, for a typical sheet resistance of 80 m Ω per square, the bandwidth, even for relatively large coils with a high turn count, generally stays above 100 MHz (see, for example, Figure B.14). Bandwidth, when reading out the coil current instead of voltage is therefore less of a limitation

When dimensioning the coil for a current readout based system, the trade-off between
sensitivity and bandwidth is alleviated

As a result of this mitigation of the bandwidth constraint, the dimensioning strategy becomes relatively simple. Ideally, one wants to maximize x for a large sensitivity. However, x is either limited by the available chip area or by the drive capability constraint. To avoid being limited by the latter, w should be chosen small enough such that when x occupies the maximum allowed amount of chip area, the required drive capability is not too large. This likely results in w being set at the minimum trace width of a certain process.

The final step is to fill this coil with as many turns as possible; this results in the largest sensitivity and inherent noise factor for the given outer dimensions. Unless the outer dimensions are extremely large (e.g., larger than

⁴If the current is read out using a proper transimpedance or current amplifier

1 mm), the resistance and inductance probably do not limit the bandwidth too much (i.e., the bandwidth does not drop below 100 MHz).

Lastly, it should also be noted that the resistance and capacitance of such a large coil are likely relatively large (tens of kilo-ohms and pico-farads). To a first-order approximation, they should have no effect on the bandwidth. However, in a distributed model, they do. To verify that the bandwidth is not limited too much, it is important that simulations are also performed using a distributed model using a cascade of multiple sections⁵ of the model shown in Figure B.1.

⁵With adequately scaled component sizes

D

ON-CHIP CALIBRATION COIL

To implement background calibration (see Section 6.2) or foreground calibration without using the current rail (see Section 5.3), a magnetic field must be generated on chip. For this purpose, an on-chip coil should be used. To be able to perform some rough calculations, it is good to know what field strengths can approximately be generated.

To generate a magnetic field, a current is required. Ideally, this current is sufficiently small so that the calibration coil does not consume too much power. By estimating the square calibration coil (see Figure D.3) as a circular current loop with radius r , the magnetic field strength at its center can be estimated as:

$$B = \frac{\mu_0 I}{2r} \quad (\text{D.1})$$

Besides, if a trace width w is chosen, the resistance of one turn can be estimated as:

$$R_{coil} = \frac{8r}{w} \cdot R_{sheet} \quad (\text{D.2})$$

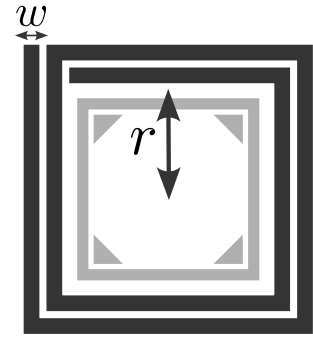


Figure D.3: On-chip calibration coil around a single Hall plate with an indication of the dimensions.

With R_{sheet} being the sheet resistance, which, for a metal trace, is nominally $80\text{m}\Omega$ per square.

Using these equations and considering that the radius increases for each additional coil turn, the plots in Figures D.1 and D.2 can be generated. These plots give an indication of the magnetic field coupling factor (in mT A^{-1}) and coil resistance (in Ω) as a function of the number of coil turns N for a coil with inner radius $r = 25\mu\text{m}$ and trace width $w = 1\mu\text{m}$. The radius is such that it fits just around a single $50 \times 50\mu\text{m}$ Hall plate¹ (which is typically used by SystematIC).

As can be seen in the plots, for roughly 20 turns, the coil has a coupling factor of approximately 350 mT A^{-1} and a resistance of roughly 400Ω . This resistance, combined with the maximum voltage headroom, limits the maximum current through the coil and, as a result, puts a limit to the maximum magnetic field. For this 20-turn coil, a voltage of 1.2V result in a current of roughly 3 mA and a magnetic field of approximately 1 mT. These values are by no means completely accurate (due to the simple assumptions made before); however, for rough calculations, they can still be used.

¹ Note that if larger sensing elements (such as a large sensing coil or multiple Hall plate in parallel) are to be calibrated, the radius r has to be larger. This is not beneficial for the magnetic field coupling factor and thus likely results in a larger current requirement

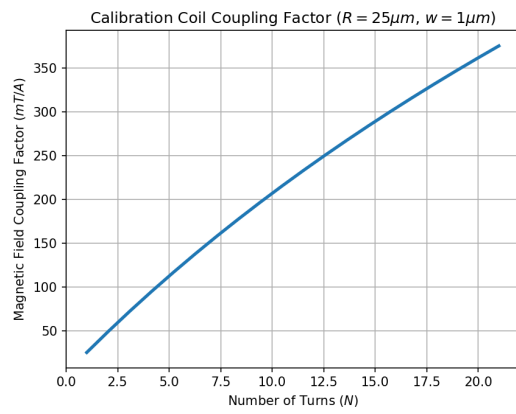


Figure D.1: Rough estimation of the magnetic field coupling factor (mTA^{-1}) of an on-chip calibration coil with $R = 25\mu\text{m}$ and $w = 2\mu\text{m}$ (with R and w as specified in Figure D.3)

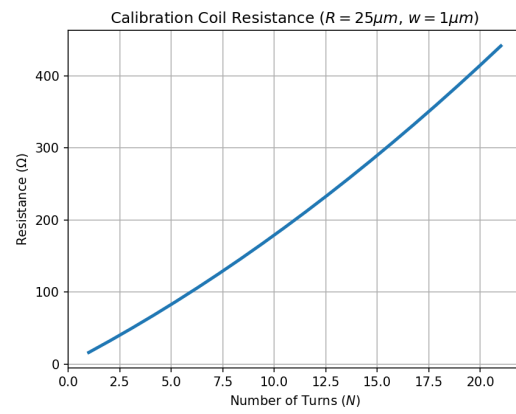


Figure D.2: Rough estimation of the resistance of an on-chip calibration coil with $R = 25\mu\text{m}$ and $w = 2\mu\text{m}$ (with R and w as specified in Figure D.3)

BIBLIOGRAPHY

- [1] S. H. Shalmany, D. Draxelmayr, and K. A. Makinwa, *A $\pm$ 36-a integrated current-sensing system with a 0.3% gain error and a 400-ua offset from -55 to 85*, IEEE Journal of Solid-State Circuits **52**, 1034 (2017).
- [2] A. Patel and M. Ferdowsi, *Current sensing for automotive electronics—a survey*, IEEE Transactions on Vehicular Technology **58**, 4108 (2009).
- [3] P. Garcha, V. Schaffer, B. Haroun, S. Ramaswamy, J. Wieser, J. Lang, and A. Chandraksan, *A duty-cycled integrated-fluxgate magnetometer for current sensing*, IEEE Journal of Solid-State Circuits **57**, 2741 (2022).
- [4] L. Jogschies, D. Klaas, R. Kruppe, J. Rittinger, P. Taptimthong, A. Wienecke, L. Rissing, and M. C. Wurz, *Recent developments of magnetoresistive sensors for industrial applications*, Sensors **15**, 28665 (2015).
- [5] R. S. Popovic, *Hall effect devices* (CRC Press, 2003).
- [6] P. Munter, *A low-offset spinning-current hall plate*, Sensors and Actuators A: Physical **22**, 743 (1990).
- [7] Z. Xin, H. Li, Q. Liu, and P. C. Loh, *A review of megahertz current sensors for megahertz power converters*, IEEE Transactions on Power Electronics **37**, 6720 (2021).
- [8] M. Biglarbegan, S. J. Nibir, H. Jafarian, and B. Parkhideh, *Development of current measurement techniques for high frequency power converters*, in 2016 IEEE International Telecommunications Energy Conference (INTELEC) (IEEE, 2016) pp. 1–7.
- [9] J. Jiang and K. A. Makinwa, *Multipath wide-bandwidth cmos magnetic sensors*, IEEE Journal of Solid-State Circuits **52**, 198 (2016).
- [10] J. Jiang and K. A. Makinwa, *A hybrid multi-path cmos magnetic sensor with 76 ppm/° c sensitivity drift and discrete-time ripple reduction loops*, IEEE Journal of Solid-State Circuits **52**, 1876 (2017).
- [11] A. Jouyaeian, Q. Fan, M. Motz, U. Ausserlechner, and K. A. Makinwa, *5.6 a 25a hybrid magnetic current sensor with 64ma resolution, 1.8 mhz bandwidth, and a gain drift compensation scheme*, in 2021 IEEE International Solid-State Circuits Conference (ISSCC), Vol. 64 (IEEE, 2021) pp. 82–84.
- [12] A. Jouyaeian, Q. Fan, R. Zamparetti, U. Ausserlechner, M. Motz, and K. A. Makinwa, *A hybrid magnetic current sensor with a multiplexed ripple-reduction loop*, IEEE Journal of Solid-State Circuits (2023).
- [13] A. Jouyaeian, Q. Fan, U. Ausserlechner, M. Motz, and K. A. Makinwa, *A hybrid magnetic current sensor with a dual differential dc servo loop*, IEEE Journal of Solid-State Circuits (2023).
- [14] A. Jouyaeian, Q. Fan, M. Motz, U. Ausserlechner, and K. A. Makinwa, *23.3 a 51a hybrid magnetic current sensor with a dual differential dc servo loop and 43ma rms resolution in a 5mhz bandwidth*, in 2023 IEEE International Solid-State Circuits Conference (ISSCC) (IEEE, 2023) pp. 22–24.
- [15] T. Funk, J. Groeger, and B. Wicht, *An integrated and galvanically isolated dc-to-15.3 mhz hybrid current sensor*, in 2019 IEEE Applied Power Electronics Conference and Exposition (APEC) (IEEE, 2019) pp. 1010–1013.
- [16] Y. Li, M. Motz, and L. Raghavan, *A fast t&h overcurrent detector for a spinning hall current sensor with ping-pong and chopping techniques*, IEEE Journal of Solid-State Circuits **54**, 1852 (2019).
- [17] A. J. Montagne, *Structured Electronics Design - A conceptual approach to amplifier design*, Vol. 3 (TUDelft Open, 2023).

- [18] M. Crescentini, R. Ramilli, G. P. Gibiino, M. Marchesi, R. Canegallo, A. Romani, M. Tartagni, and P. A. Traverso, *The x-hall sensor: Toward integrated broadband current sensing*, IEEE Transactions on Instrumentation and Measurement **70**, 1 (2020).
- [19] J. M. G. Rey, J. A. Palechor, L. F. Ariza, A. G. Roza, and F. Segura-Quijano, *Coil-on-chip: Design of integrated coils for inductive telemetry system*, in 2012 IEEE 3rd Latin American Symposium on Circuits and Systems (LASCAS) (IEEE, 2012) pp. 1–4.
- [20] C.-H. Wu, C.-C. Tang, and S.-I. Liu, *Analysis of on-chip spiral inductors using the distributed capacitance model*, IEEE journal of solid-state circuits **38**, 1040 (2003).
- [21] C. P. Yue, C. Ryu, J. Lau, T. H. Lee, and S. S. Wong, *A physical model for planar spiral inductors on silicon*, in International Electron Devices Meeting. Technical Digest (IEEE, 1996) pp. 155–158.
- [22] S. S. Mohan, M. del Mar Hershenson, S. P. Boyd, and T. H. Lee, *Simple accurate expressions for planar spiral inductances*, IEEE Journal of solid-state circuits **34**, 1419 (1999).
- [23] P. Dimitropoulos, P. Drljaca, and R. Popovic, *A 0.35 μm -cmos, wide-band, low-noise hall magnetometer for current sensing applications*, in SENSORS, 2007 IEEE (IEEE, 2007) pp. 884–887.
- [24] J. Jiang, W. J. Kindt, and K. A. Makinwa, *A continuous-time ripple reduction technique for spinning-current hall sensors*, IEEE Journal of solid-state circuits **49**, 1525 (2014).
- [25] F. Witte, K. Makinwa, and J. Huijsing, *Dynamic offset compensated CMOS amplifiers* (Springer Science & Business Media, 2009).
- [26] R. Wu, J. H. Huijsing, K. A. Makinwa, R. Wu, J. H. Huijsing, and K. A. Makinwa, *Dynamic offset cancellation techniques for operational amplifiers*, Precision Instrumentation Amplifiers and Read-Out Integrated Circuits, 21 (2013).
- [27] C. C. Enz and G. C. Temes, *Circuit techniques for reducing the effects of op-amp imperfections: autozeroing, correlated double sampling, and chopper stabilization*, Proceedings of the IEEE **84**, 1584 (1996).
- [28] C.-G. Yu and R. L. Geiger, *An automatic offset compensation scheme with ping-pong control for cmos operational amplifiers*, IEEE Journal of Solid-State Circuits **29**, 601 (1994).
- [29] B. Halg, *Piezo-hall coefficients of n-type silicon*, Journal of applied physics **64**, 276 (1988).
- [30] J. T. Maupin and M. L. Geske, *The hall effect in silicon circuits*, in [The Hall Effect and Its Applications](#), edited by C. L. Chien and C. R. Westgate (Springer US, Boston, MA, 1980) pp. 421–445.
- [31] SystematIC Design B.V., *Design data*, (2024).
- [32] U. Ausserlechner, M. Motz, and M. Holliber, *Compensation of the piezo-hall effect in integrated hall sensors on (100)-si*, IEEE Sensors journal **7**, 1475 (2007).
- [33] S. Huber, C. Schott, and O. Paul, *Package stress monitor to compensate for the piezo-hall effect in cmos hall sensors*, IEEE Sensors Journal **13**, 2890 (2013).
- [34] U. Ausserlechner, M. Motz, and M. Holliber, *Drift of magnetic sensitivity of smart hall sensors due to moisture absorbed by the ic-package [automotive applications]*, in SENSORS, 2004 IEEE (IEEE, 2004) pp. 455–458.
- [35] S. Huber, W. Leten, M. Ackermann, C. Schott, and O. Paul, *A fully integrated analog compensation for the piezo-hall effect in a cmos single-chip hall sensor microsystem*, IEEE sensors journal **15**, 2924 (2014).
- [36] D. Manic, J. Petr, and R. Popovic, *Short and long-term stability problems of hall plates in plastic packages*, in 2000 IEEE International Reliability Physics Symposium Proceedings. 38th Annual (Cat. No. 00CH37059) (IEEE, 2000) pp. 225–230.
- [37] H. Husstedt, U. Ausserlechner, and M. Kaltenbacher, *In-situ analysis of deformation and mechanical stress of packaged silicon dies with an array of hall plates*, IEEE Sensors Journal **11**, 2993 (2011).
- [38] L. W. Couch, M. Kulkarni, and U. S. Acharya, *Digital and analog communication systems*, Vol. 8 (Cite-seer, 2013).
- [39] J. Boele, *Cmos magnetic field sensor using a pick-up loop: for measuring electric currents*, (2022).