



universität
wien

DISSERTATION

Titel der Dissertation

Construction of frames by discretization of phase space

Verfasser

Christoph Wiesmeyr

angestrebter akademischer Grad

Doktor der der Naturwissenschaften (Dr.rer.nat.)

Wien, 2013

Studienkennzahl lt. Studienblatt:

A 791 405

Dissertationsgebiet lt. Studienblatt:

Mathematik

Betreuerin / Betreuer:

Prof. Dr. Hans Georg Feichtinger

Contents

1	Introduction	1
1.1	Motivation	2
1.2	The continuous Fourier Transform	2
1.2.1	The Fourier Transform and some basic operations	4
1.2.2	Hermite functions	6
1.2.3	Uncertainty	9
1.3	The Discrete Fourier Transform	10
1.3.1	The DFT and some basic operations	11
1.3.2	Discrete Hermite Functions	12
1.4	Frames	14
1.4.1	Important operators	15
1.5	The Short-Time Fourier Transform	15
1.5.1	Localization operators	18
1.5.2	STFT of distributions	20
1.6	Gabor frames	20
1.6.1	From continuous to discrete Gabor frames	22
1.6.2	Structure of Gabor systems	23
1.6.3	Multi-window Gabor frames	24
2	Gabor transforms on general lattices	27
2.1	Lattices in the TF-plane	27
2.1.1	Counting subgroups of $\mathbb{Z}_L^2$	30
2.2	Metaplectic operators	32
2.2.1	Continuous	32
2.2.2	Discrete	35
2.3	Computation on non-separable lattices	38
2.3.1	Correspondence via multiwindow Gabor	39
2.3.2	Correspondence via Smith normal form	41
2.3.3	Correspondence via shearing	41
2.3.4	Further optimization	44
2.3.5	Extension to higher dimensions	47
2.4	Implementation and timing	47
2.4.1	Implementation of the shear algorithm	48
2.4.2	Dual and tight windows	50
2.4.3	Analysis of the computational complexity	51
2.4.4	Numerical experiments	52

3	Optimal Gabor frames	55
3.1	The Janssen test and tilings of the TF-plane	55
3.2	Optimal sampling for Gaussian windows	59
3.3	Adaption of the window - optimal ambiguity functions	60
3.3.1	Concentration measures	60
3.3.2	Continuous optimization principle	62
3.3.3	Numerical experiments	66
4	Beyond Gabor - Wavelets and ERB-lets	71
4.1	Filter banks via warping	72
4.2	Warped frames	75
4.3	Construction of tight warped frames	80
4.4	A continuous warped transform	84
4.5	Discretization and numerical examples	86

Abstract

The purpose of this thesis is the construction of frames that are useful in applications. Most of the presented theory and experiments are tailored to one-dimensional signals, therefore the presented work has potential impact on one-dimensional signal processing, e.g. audio signal analysis. The first part of the thesis investigates *Gabor systems*, especially on *non-separable lattices*. We present some efficient algorithms for the realization of the underlying transform and give theoretical and experimental arguments why non-separable sampling schemes have potentially better properties than the traditional rectangular sampling of the underlying phase space. The second part of the thesis uses *warping* in the frequency domain to construct more flexible function systems. We will use the framework of *non-stationary Gabor systems* to construct systems which are adapted to a given scale. The proposed method allows for an easy construction of tight frames based on *painless* sampling of the *phase space*, which in this case depends on the warping function. *Wavelet* and Gabor transforms are special cases of this more flexible transform. As a possible application we also construct filters that follow the *ERB-scale*, which is adapted to the human hearing. This framework should be particularly interesting for audio signal processing.

Zusammenfassung

Das zentrale Thema dieser Doktoratsarbeit ist die anwendungsorientierte Konstruktion von *Frames*. Die besprochene Theorie und die vorgestellten Experimente sind zu einem großen Teil auf eindimensionale Signale zugeschnitten. Deshalb findet das vorgestellte Material auch hauptsächlich in diesem Bereich Anwendungen, wie z.B. in der Audiosignalanalyse. Der erste Teil der Arbeit beschäftigt sich mit *Gaborsystemen*, im Speziellen auf nicht separablen Gittern. Neben effizienten Algorithmen für die Realisierung der Transformation werden wir auch theoretische und numerische Experimente besprechen, die zeigen, dass nichtseparable Gitter im Phasenraum oft bessere Eigenschaften haben als die traditionellen Rechtecksgitter. Der zweite Teil der Arbeit verwendet eine Verzerrung der Frequenzachse, um flexiblere Transformationen realisieren zu können. Die *nicht-stationäre Gabor Transformation* ist ein Konzept, das flexibel genug ist, um Systeme zu konstruieren, die an eine gegebene Skala angepasst sind. Die vorgestellte Methode erlaubt es außerdem, ohne großen Aufwand *tight Frames* zu konstruieren basierend auf einer Abtastung des zugehörigen Phasenraumes, die bekannt ist als *painless sampling*. Diese flexible Transformation beinhaltet *Gabor-* und *Wavelettransformationen* als Spezialfälle, ist also eine echte Verallgemeinerung von bestens bekannten Transformationen. Als eine mögliche Anwendung konstruieren wir so eine Filterbank, die an die sogenannte *ERB-scale* angepasst ist. Diese spezielle Frequenzskala ist an das menschliche Gehör angepasst und die zugehörige Transformation hat potentielle Anwendungen im Bereich der Audiosignalanalyse.

Acknowledgements

I want to thank a lot of people for their support during the last three years. To start with the scientific side, I want to thank my advisor, Hans G. Feichtinger, for providing the possibility to work in an international EU funded project (UNLocX), which sparked many interesting collaborations and lead to interesting insights. Within that project there is a number of people I want to thank for their support, the following list is neither exhaustive, nor is it in a specific order: Bruno Torr  ssani, Pierre Vandergheynst, David Shuman, Benjamin Ricaud, Guillaume Stempfel and H  l  ne Lachambre. At the University of Vienna I am glad that I could be part of the Numerical Harmonic Analysis Group (NuHAG), where I found a very hospitable environment making the last years a nice experience. I want to particularly mention Nicki Holighaus, with whom I had very fruitful collaborations and whose ability to work through terrible equations I admire. I am glad to call him a friend beyond the scientific exchange. Also Gino Velasco should be mentioned, who I want to thank for a beautiful tour through the Philippines. I also want to thank: Peter, Sebastian, Angelika, Veit-Severin, Markus, Andreas, Martin and Darian.

It is important for me to thank my mother Maria, who passed away in 2006, and my father G  nter. Their patience and support gave me the opportunity to pursue my mathematical interest. I also want to thank my sisters Ursula and Magdalena, as well as my brother Johannes, who I enjoy being with. Also my grandparents, Alois and Maria, need to be mentioned as they have always been supportive and the priority they give to the family is comforting. Also a number of friends from my home town have been supporting me and I want to thank especially Stefan, Peter and Philipp for staying in touch even though we are not living in the same city anymore.

Chapter 1

Introduction

In this chapter we discuss the fundamentals of Fourier analysis, frames and Gabor transforms, as well as the Short Time Fourier Transform. The Fourier Transform and frame theory are the cornerstones for the further theory of Gabor frames. We will show there how we can decompose square integrable functions into basic building blocks. These atoms have a certain interpretation and allow extract information from a signal by knowing its transform coefficients. We will not only discuss the important technical tools and theories but also touch possible applications, as this motivates the results and work in later chapters.

Gabor Transforms and the Short Time Fourier Transform or in general time-frequency analysis have quite a wide range of applications and have been adopted in audio signal processing as a main tool. Throughout most of the text we will deal with one dimensional signals, at certain places we will make remarks about higher dimensions, which is interesting for e.g. image processing.

Before we start with the basics of Fourier theory we introduce some basic notation and definitions. The p -norm for $1 < p < \infty$ of a measurable function $f : S \rightarrow \mathbb{C}$ is explained by

$$\|f\|_p = \left(\int_S |f|^p d\mu \right)^{1/p}, \quad (1.1)$$

where (S, Σ, μ) is a measure space. This leads to the definition of the classical Lebesgue spaces

$$L^p(S) = \{f : f \text{ measurable}, \|f\|_p < \infty\}. \quad (1.2)$$

For the above excluded case $p = \infty$ we define

$$\|f\|_\infty = \text{ess sup } f \quad (1.3)$$

For a proper introduction of the Lebesgue integral as well as measure theoretic tools we refer to standard literature in that field, e.g. [99]. The Lebesgue spaces are Banach spaces, for $p = 2$ one even has an inner product with respect to which the space $L^2(S)$ becomes a Hilbert space. The inner product between two functions $f, g \in L^2(S)$ is defined in the usual sense

$$\langle f, g \rangle = \int_S f \bar{g} d\mu. \quad (1.4)$$

Throughout most of the text the measure space will be $\mathbb{R}^d$ equipped with the standard Lebesgue measure. In many places when there is no danger of confusion we also omit the explicit measure space and just write L^p .

For functions on the d -dimensional Euclidean space $\mathbb{R}^d$ we introduce the multi-index notation for derivatives. For some given multi-index $\alpha \in \mathbb{N}^d$ we define $|\alpha| = \sum_{j=1}^d \alpha_j$. We further use $x^\alpha = \prod_{j=1}^d x^{\alpha_j}$ and $D^\alpha = \prod_{j=1}^d \frac{\partial^{\alpha_j}}{\partial x^{\alpha_j}}$. For the factorials of a multi-index we define $\alpha! = \prod_{j=1}^d \alpha_j!$.

For vectors $x, \xi \in \mathbb{R}^d$ we denote their inner product by $x \cdot \xi$ and the Euclidean norm of a vector by $|x| = \|x\|_2 = (x \cdot x)^{1/2}$, which is a slight abuse of notation because the length of a multi-index is denoted in the same way. However, in the text no confusion arises from this inaccuracy. Analogous to the Lebesgue spaces one can define p -norms of vectors.

Throughout almost all of the first chapter we will refrain from giving proofs, unless the computations are necessary to find the correct normalizations. All the omitted proofs can be found in the given citations, the purpose of the first chapter is to collect a basis of results that will be used later in the text.

1.1 Motivation

The motivation behind the research presented in this thesis is the optimal construction of function systems for one-dimensional signal analysis. In this thesis I present the research done in an EU funded project called UNLocX (Uncertainty principles versus localization properties, function systems for efficient coding schemes). The project finished at the end of August of 2013. The University of Vienna had the task to build frames through discretization of phase space. While this project lead to many interesting collaborations, which allowed me to look into several interesting topics, the main core of this text will be devoted to the construction of frames. This thesis presents mainly two constructions, that are beyond the established state of the art: non-separable Gabor schemes and a flexible method to construct a wide range of transforms adapted to a frequency scale chosen by the user. While the first one has been considered in the literature and the contribution is mainly an improvement of the existing algorithms, the latter mentioned idea is fundamentally new and leads to transforms beyond the state of the art. Beyond the algorithmic discussion of non-separable Gabor schemes we will also look into optimality criteria for choosing sampling patterns in phase space.

1.2 The continuous Fourier Transform

The Fourier Transform is a widely used tool to tackle all kinds of problems in pure mathematics as well as in engineering. In the following we are mostly interested in the application of the Fourier Transform to signal processing but it is important to note that the importance of this transform is not limited to this field.

Definition 1.1 (Fourier Transform). Let $f \in L^1(\mathbb{R}^d)$. Then we define its Fourier Transform as

$$\hat{f}(\xi) = \mathcal{F}f(\xi) = \int_{\mathbb{R}^d} f(x) e^{-2\pi i x \cdot \xi} dx. \quad (1.5)$$

It is easy to see that $\|\hat{f}\|_\infty \leq \|f\|_1$. We would like to define the Fourier Transform as an operator acting on $L^2(\mathbb{R}^d)$, this extension is a standard procedure and is done via a density argument [26, 54]. The crucial statement behind the extension is Plancherel's theorem.

Theorem 1.2 (Plancherel). *If $f \in L^1(\mathbb{R}^d) \cap L^2(\mathbb{R}^d)$. Then*

$$\|f\|_2 = \|\hat{f}\|_2. \quad (1.6)$$

Since $L^1(\mathbb{R}^d) \cap L^2(\mathbb{R}^d)$ is dense in $L^2(\mathbb{R}^d)$ we can extend the Fourier Transform to a unitary operator on $L^2(\mathbb{R}^d)$. From now on we denote by $\mathcal{F}f = \hat{f}$ the Fourier transform on $L^2(\mathbb{R}^d)$, which respects the important formula

$$\langle f, g \rangle = \langle \hat{f}, \hat{g} \rangle. \quad (1.7)$$

The inverse Fourier Transform is then easily computed

Theorem 1.3. *Let $f, \hat{f} \in L^1(\mathbb{R}^d)$. Then*

$$f(t) = \int_{\mathbb{R}} \hat{f}(\xi) e^{2\pi i \xi \cdot t}. \quad (1.8)$$

The Fourier Transform allows to extract frequency information from a signal f . We call f the *time representation* and $\hat{f}$ the *frequency representation*.

An important space connected to the Fourier Transform is the space of *Schwartz functions*.

Definition 1.4. The class of *Schwartz functions* is defined as

$$\mathcal{S}(\mathbb{R}^d) = \left\{ f : \mathbb{R}^d \rightarrow \mathbb{C} : \sup_{x \in \mathbb{R}^d} |D^\alpha x^\beta f(x)| < \infty \right\}. \quad (1.9)$$

The Schwartz functions serve as test function space when defining *tempered distributions* as its dual space [48].

Definition 1.5. The space $\mathcal{S}'$, the dual space of $\mathcal{S}$, is called the space of *tempered distributions*.

A function $f \in L^2(\mathbb{R}^d)$ can be identified with a distribution via the functional

$$F_f(\varphi) = \langle f, \varphi \rangle. \quad (1.10)$$

This is the standard procedure to view the tempered distributions as a generalization of L^2 functions. Among the reasons to develop the theory of distributions in the first place is the point evaluation distribution

$$\delta_x \varphi = \varphi(x), \quad \text{for } \varphi \in \mathcal{S}. \quad (1.11)$$

In the following we will use tempered distributions in several places without describing the theory in detail. Most of the rules for manipulating distributions carry over from the usage of functions, however one has to be careful, as there

are some operations that are not properly definable (e.g. the product of distributions). All the operations that are used in the text are properly defined and justified in the framework of distributions.

The space of Schwartz functions has the convenient property that it is invariant under the Fourier Transform

$$\mathcal{F}(\mathcal{S}(\mathbb{R}^d)) = \mathcal{S}(\mathbb{R}^d). \quad (1.12)$$

This allows to define the Fourier Transform of distributions

$$\langle \mathcal{F}F, \varphi \rangle = \langle F, \mathcal{F}^{-1}\varphi \rangle, \quad (1.13)$$

for some $F \in \mathcal{S}'$. We mention now two important distributional Fourier transform pairs. The pure frequencies or characters are defined as

$$\chi_x(t) = e^{-2\pi i x \cdot t}. \quad (1.14)$$

This leads to the useful formulas

$$\mathcal{F}\delta_x = \chi_x \quad \text{and} \quad (1.15)$$

$$\mathcal{F}\chi_x = \delta_x. \quad (1.16)$$

1.2.1 The Fourier Transform and some basic operations

Some operations and their interplay with the Fourier Transform is of special interest. We start with the convolution of two functions.

Definition 1.6. Let $f, g \in L^1(\mathbb{R}^d)$. Then their convolution is

$$f * g(t) = \int_{\mathbb{R}^d} f(\tau)g(t - \tau) d\tau. \quad (1.17)$$

The convolution is well defined, as it satisfies the fundamental relation

$$\|f * g\|_1 \leq \|f\|_1 \|g\|_1. \quad (1.18)$$

Of great importance are also the translation and modulation operators, as they will represent time and frequency shifts, respectively in the definition of STFT and Gabor systems. We also define dilation operators as they naturally appear in many places.

Definition 1.7. Let $f \in L^2(\mathbb{R}^d)$ and $x, \xi \in \mathbb{R}^d$, $a \in \mathbb{R}^+$. Then we define the three unitary operators

$$T_x f(t) = f(t - x) \quad (1.19)$$

$$M_\xi f(t) = e^{2\pi i t \cdot \xi} f(t) \quad (1.20)$$

$$D_a f(t) = a^{-d/2} f(a^{-1}t). \quad (1.21)$$

Now that we have defined these important operators we collect some of their properties and investigate their interplay with the Fourier Transform.

Proposition 1.8. Let $f, g \in L^2(\mathbb{R}^d)$. Then

1. $(\hat{f})^\wedge(t) = f(-t)$.
2. $(D^\alpha f)^\wedge(\xi) = (2\pi i \xi)^\alpha \hat{f}(\xi)$ for $f \in \mathcal{S}(\mathbb{R}^d)$.
3. If f, g in $L^2(\mathbb{R}^d) \cap L^1(\mathbb{R}^d)$, then $(f * g)^\wedge = \hat{f} \cdot \hat{g}$.
4. $(M_\xi f)^\wedge = T_\xi \hat{f}$ and $(T_x f)^\wedge = M_{-x} \hat{f}$.
5. $(D_a f)^\wedge = D_{1/a} \hat{f}$.
6. $T_x M_\xi = e^{-2\pi i x \cdot \xi} M_\xi T_x$.

Another very important formula connected to the Fourier Transform is the *Poisson summation formula*, which relates periodization and sampling on the time and Fourier side. A space, which comes up naturally when treating periodizations is the Wiener space.

Definition 1.9. A function $g \in L^\infty(\mathbb{R}^d)$ belongs to the *Wiener space* $W(\mathbb{R}^d)$ if

$$\|g\|_W = \sum_{n \in \mathbb{Z}^d} \operatorname{ess\,sup}_{x \in [0,1]^d} |g(x+n)|. \quad (1.22)$$

Theorem 1.10. Assume that $f, \hat{f} \in W(\mathbb{R}^d)$, then

$$\sum_{n \in \mathbb{Z}^d} f(x+n) = \sum_{n \in \mathbb{Z}^d} \hat{f}(n) e^{2\pi i n \cdot x}. \quad (1.23)$$

The Poisson summation formula can easily be generalized to non-integer sampling.

Corollary 1.11. Assume that $f, \hat{f} \in W(\mathbb{R}^d)$, then

$$\sum_{n \in \mathbb{Z}^d} f(x + \alpha n) = \alpha^{-d} \sum_{n \in \mathbb{Z}^d} \hat{f}(n/\alpha) e^{2\pi i n \cdot x / \alpha}. \quad (1.24)$$

Remark 1.12. Poisson's summation formula can be formulated in the context of distributions

$$\left(\sum_{n \in \mathbb{Z}} \delta_{an} \right)^\wedge = \frac{1}{a^d} \sum_{n \in \mathbb{Z}} \delta_{n/a}. \quad (1.25)$$

In this context it can be interpreted as a relation between periodization, which is a convolution with a delta train, and sampling, which is pointwise multiplication with a delta train

$$\left(f * \left(\sum_{n \in \mathbb{Z}} \delta_{an} \right) \right)^\wedge = \frac{1}{a^d} \hat{f} \cdot \left(\sum_{n \in \mathbb{Z}} \delta_{n/a} \right). \quad (1.26)$$

1.2.2 Hermite functions

Since the Fourier Transform is a unitary operator on $L^2(\mathbb{R}^d)$ we know that it can only have eigenvalues with modulus 1.

Definition 1.13. Define the *standard normalized Gaussian* function as

$$g_0(t) = 2^{d/4} e^{-\pi t^2}, \quad (1.27)$$

which is normalized, such that $\|g_0\|_2 = 1$.

Lemma 1.14. Let g_0 denote the normalized standard Gaussian function. Then

$$g_0(t) = \widehat{g_0}(t). \quad (1.28)$$

The remainder of this section is devoted to the investigation of the so-called *Hermite functions*, which we will prove to be an ONB for $L^2(\mathbb{R}^d)$ consisting of eigenfunctions of the Fourier Transform. The derivation of the Hermite functions is a standard procedure and can be found in e.g. [45, 78]. We still derive the exact formulas as the Fourier Transform comes in different normalizations. In [45, 54] the Hermite functions are introduced as the pre-image of monomials under the Bargman transform, here we choose a more direct way.

Let us start with the definition of the Hermite polynomials and functions.

Definition 1.15. Let α be some d -dimensional multi-index. The *Hermite polynomials* on $\mathbb{R}^d$ are defined as

$$h_\alpha(t) = (-1)^{|\alpha|} e^{2\pi t^2} \frac{\partial^\alpha}{\partial t^\alpha} e^{-2\pi t^2}. \quad (1.29)$$

The *Hermite functions* on $\mathbb{R}^d$ are defined as

$$H_\alpha(t) = \left(\frac{1}{2\sqrt{\pi}} \right)^{|\alpha|} (\alpha!)^{-1/2} g_0(t) h_\alpha(t). \quad (1.30)$$

The above defined polynomials and functions have important properties. We first treat the one-dimensional case, the generalization to higher dimensions is straightforward.

Theorem 1.16. Let $\{H_n\}_{n \in \mathbb{N}_0}$ denote the Hermite functions and $\{h_n\}_{n \in \mathbb{N}_0}$ denote the Hermite polynomials for $d = 1$. Then

1. the Hermite polynomials satisfy the recurrence relation

$$h'_n(t) = 4\pi n h_{n-1}(t). \quad (1.31)$$

2. the Hermite polynomials satisfy the generating function relation

$$e^{2\pi(2t\tau - \tau^2)} = \sum_{n=0}^{\infty} \frac{h_n(t)}{n!} \tau^n \quad (1.32)$$

3. the Hermite functions form an ONB of $L^2(\mathbb{R})$.

4. the Hermite functions are eigenvectors of the one-dimensional Fourier Transform

$$\widehat{H}_n = (-i)^n H_n. \quad (1.33)$$

Proof. We start by showing (1.31). We use the product rule and the Leibnitz rule to compute

$$\begin{aligned} h'_n(t) &= (-1)^n e^{2\pi t^2} \left(4\pi t \frac{d^n}{dt^n} e^{-2\pi t^2} - \frac{d^n}{dt^n} (4\pi t e^{-2\pi t^2}) \right) \\ &= (-1)^{n-1} e^{2\pi t^2} 4\pi n \frac{d^{n-1}}{dt^{n-1}} e^{-2\pi t^2} \\ &= 4\pi n h_{n-1}. \end{aligned} \quad (1.34)$$

The generating function (1.32) is derived by computing the Taylor expansion at 0 of the function

$$f(\tau) = e^{-2\pi(t-\tau)^2} \quad (1.35)$$

Some computation shows that

$$\frac{d^n}{d\tau^n} f(0) = e^{-2\pi t^2} h_n(t), \quad (1.36)$$

which proves (1.32) since we can use Taylor's expansion to find

$$e^{-2\pi t^2} e^{2\pi(2t\tau - \tau^2)} = f(\tau) = e^{-2\pi t^2} \sum_{n=0}^{\infty} \frac{h_n(t)}{n!} \tau^n. \quad (1.37)$$

To show the ONB property we first compute the product of two Hermite functions

$$H_n(t) H_m(t) = \left(\frac{1}{2\sqrt{\pi}} \right)^{n+m} \sqrt{n!m!} 2^{d/2} (-1)^m h_n(t) \frac{d^m}{dt^m} e^{-2\pi t^2}. \quad (1.38)$$

We assume that w.l.o.g. $m < n$ and compute by partial integration and (1.31)

$$\int_{\mathbb{R}} h_n(t) \frac{d^m}{dt^m} e^{-2\pi t^2} dt = (4\pi)^n n! \int_{\mathbb{R}} \frac{d^{m-n}}{dt^{m-n}} e^{-2\pi t^2} dt = 0. \quad (1.39)$$

For $m = n$ the same computation leads to

$$\int_{\mathbb{R}} h_n(t) \frac{d^n}{dt^n} e^{-2\pi t^2} dt = (4\pi)^n \frac{n!}{\sqrt{2}}. \quad (1.40)$$

For showing completeness it suffices to show that if

$$\int_{\mathbb{R}} f(t) t^n e^{-\pi t^2} dt = 0 \quad (1.41)$$

for all $n \in \mathbb{N}_0$, then f has to be the zero function. To prove that we note that (1.41) is equivalent to

$$0 = \sum_{n=0}^{\infty} \frac{(-2\pi\xi)^n}{n!} \int_{\mathbb{R}} f(t) t^n e^{-\pi t^2} dt = (f \cdot e^{-\cdot^2})^\wedge(\xi) \quad (1.42)$$

The Fourier Transform $(f \cdot e^{-\cdot^2})$ vanishes if and only if $f = 0$.

We still have to proof relation (1.33). In the following we denote the coefficients in the definition of the Hermite functions (1.30) by c . Using the generating function (1.32) we find that

$$c e^{-\pi t^2} e^{2\pi(2t\tau - \tau^2)} = \sum_{n=0}^{\infty} \frac{H_n(t)}{n!} \tau^n. \quad (1.43)$$

We denote the left hand side of this equation by $G(t)$. Completing squares and using the Fourier invariance of the standard Gaussian function its Fourier Transform can be computed explicitly

$$\begin{aligned} \hat{G}(\xi) &= c e^{-\pi \xi^2} e^{2\pi(2i\tau\xi) - (-i\tau)^2} \\ &= \sum_{n=0}^{\infty} \frac{(-i)^n H_n(\xi)}{n!} \tau^n. \end{aligned} \quad (1.44)$$

On the other hand

$$\hat{G}(\xi) = \sum_{n=0}^{\infty} \frac{\widehat{H}_n(\xi)}{n!} \tau^n. \quad (1.45)$$

Equating coefficients in the above power series proves the Fourier invariance of the Hermite functions. $\square$

As an easy corollary we state the above result for higher dimensions.

Corollary 1.17. *Let $\{H_n\}_{\alpha \in \mathbb{N}_0^d}$ denote the multi-dimensional Hermite polynomials. Then*

1. *the Hermite functions form an ONB of $L^2(\mathbb{R}^d)$.*
2. *the Hermite functions are eigenvectors of the Fourier Transform*

$$\widehat{H}_\alpha = (-1)^{|\alpha|} H_\alpha. \quad (1.46)$$

Proof. The proof easily follows from the observation

$$H_\alpha(t) = \prod_{j=1}^d H_{\alpha_j}(t_j), \quad (1.47)$$

where H_{α_j} denote the one dimensional Hermite functions and t_j is the j -th component of t . $\square$

One of the applications of Hermite functions is the *fractional Fourier Transform*, which is a generalization of the standard FT [72]. It is possible to define it in several equivalent ways, we base ours on the orthonormal decomposition of $L^2(\mathbb{R}^d)$ into Hermite functions since they eigenfunctions of the standard Fourier Transform. We shall see later that this transform corresponds to rotations in phase space and is therefore an interesting operation for our purposes.

Definition 1.18. For some angle α the *fractional Fourier Transformation* of a function $f \in L^2(\mathbb{R}^d)$ is defined as

$$\mathcal{F}^\alpha = \sum_{n=0}^{\infty} e^{-i\alpha n} \langle f, H_n \rangle H_n. \quad (1.48)$$

As an easy consequence of the explicitly computed eigenvalues of the Hermite functions in connection with the Fourier Transform (1.46) we immediately see that $\mathcal{F}^{\pi/2} = \mathcal{F}$. Furthermore, the fractional Fourier Transform is unitary, satisfies $\mathcal{F}^\alpha \mathcal{F}^\beta = \mathcal{F}^{\alpha+\beta}$ and the Hermite functions are eigenfunctions.

1.2.3 Uncertainty

Uncertainty principles put a bound on the joint time and frequency concentration of a function. They appear in many areas such as quantum physics, harmonic analysis and signal processing. As general references for uncertainty principles we mention [22, 46, 54, 58, 85]. In this section the focus is on bounds in the fashion of the famous *Heisenberg uncertainty principle*, which is variance based. However, there are also other ways of bounding the concentration, e.g. by determining the maximal simultaneous decay of a signal and its Fourier Transform such as *Hardy's uncertainty principle* [58]. Heisenberg's uncertainty principle is closely related to the definition of the quantum mechanical *Hamiltonian*. We first define the position and the momentum operators.

Definition 1.19. On the d -dimensional Schwartz space we define the *position operator* and *momentum operator* in dimension $j = 1, \dots, d$ respectively as

$$X_j f(t) = t_j f(t), \quad (1.49)$$

$$P_j f(t) = \mathcal{F}^{-1} \xi_j \mathcal{F} f = \frac{1}{2\pi i} \frac{\partial}{\partial x_j} f. \quad (1.50)$$

In the one-dimensional case we omit the dimension and just write X and P .

These operators are self-adjoint in the $L^2(\mathbb{R}^d)$ sense.

Theorem 1.20 (Robertson-Schrödinger). *Let A, B be self-adjoint operators on some Hilbert space $\mathcal{H}$. Then for any $a, b \in \mathbb{R}$ and f in the intersection of the domains of AB and BA*

$$\|(A - a)f\| \|(B - b)f\| \geq \frac{1}{2} |\langle (AB - BA)f, f \rangle|. \quad (1.51)$$

As a corollary we can now state the classical Heisenberg uncertainty principle, the optimally concentrated function is a Gaussian function.

Corollary 1.21 (Heisenberg Uncertainty principle). *Let $f \in L^2(\mathbb{R}^d)$, then for all $a_j, b_j \in \mathbb{R}$*

$$\|(X_j - a_j)f\| \|(D_j - b_j)f\| \geq \frac{1}{4\pi} \|f\|^2. \quad (1.52)$$

Equality holds for all j if and only if

$$f = d T_a M_b D_c g_0, \quad (1.53)$$

for some $d \in \mathbb{R}$ and $c \in \mathbb{R}^+$.

Corollary 1.22. *Let $f \in L^2(\mathbb{R}^d)$, then*

$$\|X_j f\|^2 + \|P_j f\|^2 \geq \frac{1}{2\pi} \|f\|^2. \quad (1.54)$$

Equality for all j holds if and only if

$$f = c g_0, \quad (1.55)$$

for some $c \in \mathbb{R}$.

In one dimension

$$\|Xf\|^2 + \|Pf\|^2 = \langle (X^2 + P^2)f, f \rangle. \quad (1.56)$$

The operator $\mathbb{H} = X^2 + P^2$ is called (*quantum mechanical*) *Hamiltonian*. It turns out that the eigenfunctions of this operator are the Hermite functions [45], as they satisfy the relation

$$2\pi(X^2 + P^2)H_n = (2n + 1)H_n. \quad (1.57)$$

Therefore the Hermite functions can not only be seen as an orthonormal basis diagonalizing the Fourier Transform, but also as a basis that is optimally concentrated in the sense of (1.54).

1.3 The Discrete Fourier Transform

The Fourier Transform on $\mathbb{R}$ as introduced above has a counterpart on the finite discrete Hilbert space of vectors of length L given by $\mathbb{C}^L$. The reader should bear in mind that the argument of a vector is to be understood modulo the signal length L . In mathematical terms a finite discrete signal is a function $f : \mathbb{Z}/L\mathbb{Z} \rightarrow \mathbb{C}$, however in the following we denote this in a sloppy way by $f \in \mathbb{C}^L$.

Definition 1.23. The *Discrete Fourier Transform* (DFT) of a vector $f \in \mathbb{C}^L$ is defined as

$$\hat{f}(l) = \frac{1}{\sqrt{L}} \sum_{m=0}^{L-1} f(m) e^{-2\pi i m l / L}, \quad l = 0, \dots, L-1. \quad (1.58)$$

The normalization term makes the transform unitary and we use the same notation as for the continuous version of the Fourier Transform whenever there is no confusion arising from this abuse of notation. At places where a clear distinction between the two operations is necessary we denote by $\mathcal{F}_{\mathbb{R}}$ the Fourier Transform on the real line and by $\mathcal{F}_{\mathbb{C}}$ the finite discrete Fourier Transform. The inverse transform is given by

$$f(l) = \frac{1}{\sqrt{L}} \sum_{m=0}^{L-1} \hat{f}(m) e^{2\pi i m l / L}, \quad l = 0, \dots, L-1. \quad (1.59)$$

Often the DFT is referred to as FFT, which is technically incorrect. While the DFT is the transform, the FFT or Fast Fourier Transform is an algorithm for

computing the DFT. The FFT, popularized by Cooley and Tukey in [19], is the base of many algorithms that are used today. It is a divide and conquer algorithm, that allows to realize the DFT with a computational complexity of $\mathcal{O}(L \log L)$ instead of the straightforward implementation with a complexity of $\mathcal{O}(L^2)$.

We will revisit many of the statements that were made for the continuous FT and see that in many cases one can obtain similar statements, the main difference being mostly the normalization. Since we work with a unitary version of the DFT, we will show some of the proofs, as they lead to the correct normalization, which is hard to find in the literature. The definition of the DFT is only for one dimension, but since the higher dimensional Fourier Transform is just a tensor product of one dimensional transforms, generalization is an easy task. However, especially in the discrete case the treatment in one dimension avoids notational complication. The connection between $\mathcal{F}_{\mathbb{R}}$ and $\mathcal{F}_{\mathbb{C}}$ will be shortly explained in Section 1.6.1 later in this chapter.

1.3.1 The DFT and some basic operations

This section will largely follow Section 1.2.1, i.e. the corresponding section for the continuous Fourier Transform. To avoid complicated notation we will keep the same letters and signs for the operations, from the context it will always be clear which operation is needed.

Definition 1.24. Let $f, g \in \mathbb{C}^L$. Their discrete convolution is

$$f * g(l) = \sum_{m=0}^{L-1} f(m)g(l-m). \quad (1.60)$$

Apart from the convolution of two vectors we can also define translations and modulations in the same spirit as their continuous counterparts.

Definition 1.25. Let $f \in \mathbb{C}^L$. Then we define the finite discrete versions of translation and modulation respectively as

$$T_x f(l) = f(l-x), \quad \text{and} \quad (1.61)$$

$$M_{\xi} f(l) = e^{2\pi i l \xi / L} f(l). \quad (1.62)$$

Proposition 1.26. Let $f, g \in \mathbb{C}^L$. Then

1. $(\hat{f})^{\wedge}(l) = f(-l)$.
2. $(f * g)^{\wedge} = \sqrt{L} \hat{f} \cdot \hat{g}$.
3. $(M_{\xi} f)^{\wedge} = T_{\xi} \hat{f}$ and $(T_x f)^{\wedge} = M_{-x} \hat{f}$.
4. $T_x M_{\xi} = e^{-2\pi i x \xi / L} M_{\xi} T_x$.

Some of the statements of Proposition 1.8 cannot be translated directly to the discrete domain because the involved operators are not defined in a meaningful way.

The Poisson summation formula can be carried over to the finite discrete setting. It is essentially the same statement as in Theorem 1.10 but with different normalization factors. We formulate the result using finite discrete delta functions in $\mathbb{C}^L$, which we define for some $a = 0, \dots, L-1$ as

$$\delta_a(l) = \begin{cases} 1 & \text{if } l = a \\ 0 & \text{else.} \end{cases} \quad (1.63)$$

Theorem 1.27. *Given some signal length $L \in \mathbb{N}$ and a period $a|L$. Then*

$$\left(\sum_{l=0}^{L/a-1} \delta_{al} \right)^\wedge = \frac{\sqrt{L}}{a} \sum_{l=0}^{a-1} \delta_{lL/a}. \quad (1.64)$$

Proof. It is well known that

$$\sum_{l=0}^{L/a-1} e^{2\pi i a l \xi / L} = \begin{cases} \frac{L}{a} & \text{if } \frac{L}{a} | \xi \\ 0 & \text{else.} \end{cases} \quad (1.65)$$

Using this identity the assertion easily follows using the identity

$$\hat{\delta}_{an}(\xi) = \frac{1}{\sqrt{L}} e^{2\pi i a n \xi / L}. \quad (1.66)$$

□

We can write Poisson's summation formula also in the following form for a closer resemblance to Theorem 1.10

$$\sum_{n=0}^{L/a-1} f(l + an) = \frac{\sqrt{L}}{a} \sum_{n=0}^{a-1} \hat{f}(Ln/a) e^{2\pi i n l / a}. \quad (1.67)$$

This is an easy consequence from Equation (1.64) above by applying an inverse Fourier Transform to

$$\left(\sum_{n=0}^{L/a-1} f(l + an) \right)^\wedge = \sum_{n=0}^{L/a-1} (f * \delta_{an})^\wedge(l) = \frac{L}{a} \sum_{n=0}^{a-1} \hat{f}(Ln/a) \delta_{Ln/a}. \quad (1.68)$$

1.3.2 Discrete Hermite Functions

For the definition of the discrete Hermite functions we refer to [13, 72], where these functions are used to define the fractional Fourier Transform. The two main properties of the continuous Hermite functions are the facts that they diagonalize the Fourier Transform and that they are orthogonal. Any sensible definition for the finite discrete setting should therefore respect these two properties with respect to the DFT. Later in the text we will describe another method to obtain a choice of discrete Hermite functions based on STFT multipliers, c.f. Section 3.3.

Of particular interest is of course the discrete version of the standard Gaussian, i.e. the 0-th Hermite function. It should be an eigenvector of the Fourier

Transform, this can be obtained by periodizing and sampling the continuous standard Gaussian function, as we will detail later in Section 1.6.1. It is then of course possible to do the same for the continuous Hermite functions, this allows to obtain a discrete set of Hermite functions that preserves the eigenvector property (1.33). However, these functions will not be orthogonal to each other anymore. In the following we mention another way of defining discrete Hermite functions through the discretization of the quantum mechanical Hamiltonian operator defined in 1.2.3. The computation of these functions, however, involves the diagonalization of an $L \times L$ matrix. This is prohibitive for large L , therefore also the sampled periodized Hermite functions have been used for obtaining a fast (but more sloppy) version of the fractional Fourier transform [?].

Since the DFT is a linear operator it has a matrix representation. However, diagonalizing this matrix directly will not give a satisfactory result because one gets four eigenspaces that correspond to the eigenvalues of the form $(-i)^n$. As opposed to the Fourier Transform, the Hamiltonian has distinct eigenvalues and would therefore be an ideal candidate for the definition of discrete Hermite functions. For definition of the discrete Hamiltonian we replace the continuous position operator P^2 by the central difference operator.

Definition 1.28. On $\mathbb{C}^L$ the squares of the *discrete position operator* and the *discrete momentum operator* are defined respectively as

$$P^2 f(l) = f(l+1) - 2f(l) + f(l-1), \quad \text{and} \quad (1.69)$$

$$X^2 f(l) = \mathcal{F}^{-1} P^2 \mathcal{F} f(l). \quad (1.70)$$

The discrete Hamiltonian is defined as $\mathbb{H} = P^2 + X^2$.

For some vector $f \in \mathbb{C}^L$ the quantities $\langle P^2 f, f \rangle$ and $\langle X^2 f, f \rangle$ are a measure of the variance of f in position and momentum. The next step is to investigate the commutation relations between the DFT and the discrete Hamiltonian, the following two statements are taken from [13].

Theorem 1.29. *The DFT commutes with the discrete Hamiltonian operator.*

The theorem above shows that the discrete Hamiltonian and the DFT have a common set of eigenvectors.

Theorem 1.30. *If L is not a multiple of 4, the eigenvalues of the discrete Hamiltonian are distinct. Otherwise, there is only one eigenvector with degeneracy 2 and the others are distinct.*

The above presented theory justifies the use of the eigenfunctions of the discrete Hamiltonian as the equivalent of the continuous version. The discrete Hamiltonian is a banded matrix with only the lower and upper side diagonal different from 0. Therefore, techniques for sparse matrices can be harnessed to diagonalize $\mathbb{H}$.

The construction of a discrete counterpart of the Hamiltonian operator allows us to adopt some kind of notion of uncertainty for this discrete case. Analogous to the continuous case the sum of the variances in position and momentum of a vector $f \in \mathbb{C}^L$ described by $\langle \mathbb{H} f, f \rangle$ can be bounded from below by a constant, i.e. the smallest eigenvalue of $\mathbb{H}$. For discrete uncertainty principles we refer to [75], where the author also discusses an uncertainty principle formulated via the product of the two variances. This notion of uncertainty is closer to the classical Heisenberg uncertainty for the continuous case stated in Corollary 1.21.

1.4 Frames

Over the past decades frames became an important and widely used tool in mathematics. Frames are a generalization of bases and offer more flexibility. A frame can be described as a possibly overcomplete basis, that still spans the space in a stable way. While we require basis elements to be linearly independent, for a frame this is not a condition. For a thorough treatment we refer to the book of Ole Christensen [15], excellent sections on frames can also be found in [26, 54].

We define frames for a general separable Hilbert space, which we always denote by $\mathcal{H}$ in the remainder of this section. The reason for formulating the theory in this rather abstract setting is that we need the results both for the continuous space $L^2(\mathbb{R}^d)$ as well as for the discrete space $\mathbb{C}^L$, as well as $L^2(D)$ for a (possibly unbounded) interval D in the last chapter. In this section we will omit the Hilbert space when denoting the norm and the scalar product in several places. We write $\|f\|_{\mathcal{H}} = \|f\|$ and $\langle f, g \rangle_{\mathcal{H}} = \langle f, g \rangle$.

Frames generalize the concept of bases, in particular of orthonormal bases (ONB). Given an orthonormal basis $\{g_i\}_{i \in I}$ it is well known and easy to see that any element $f \in \mathcal{H}$ can be expanded as

$$f = \sum_{i \in I} \langle f, g_i \rangle g_i. \quad (1.71)$$

We call the family $\{\langle f, g_i \rangle\}_{i \in I}$ *expansion coefficients* with respect to $\{g_i\}_{i \in I}$. A specific property of ONBs is Parseval's identity

$$\sum_{i \in I} |\langle f, g_i \rangle|^2 = \|f\|^2, \quad (1.72)$$

which means that the l^2 -norm of the expansion coefficients is equivalent to the Hilbert space norm. Furthermore, if one element is removed from an ONB not every element from the Hilbert space has an expansion of the form (1.71) anymore. Such a system is *not* overcomplete and in the following we will call such a family *exact*. The concept of ONB is rather restrictive and our goal will be to relax the conditions to allow overcomplete systems, while keeping the property that vectors can be reconstructed from their expansion coefficients.

Equation (1.72) serves as an inspiration for the generalization of ONBs.

Definition 1.31. A family $\{g_i\}_{i \in I}$ in $\mathcal{H}$ is called a *frame* if and only if there exist $A > 0$ and $B < \infty$, such that

$$A\|f\|^2 \leq \sum_{i \in I} |\langle f, g_i \rangle|^2 \leq B\|f\|^2. \quad (1.73)$$

A frame is called *tight* if and only if $A = B$. The ratio of the frame bounds B/A is called *condition* of the frame. If only an upper frame bound B exists we call the family *Bessel sequence*.

1.4.1 Important operators

Connected to the concept of frames there are some important operators. We define

$$C : \mathcal{H} \rightarrow l^2(I), \quad f \mapsto \{\langle f, g_i \rangle\}_{i \in I}, \quad (1.74)$$

$$D : l^2(I) \rightarrow \mathcal{H}, \quad c \mapsto \sum_{i \in I} c_i g_i, \quad (1.75)$$

$$S : \mathcal{H} \rightarrow \mathcal{H}, \quad f \mapsto DCf = \sum_{i \in I} \langle f, g_i \rangle g_i. \quad (1.76)$$

We call C , D and S *analysis operator*, *synthesis operator* and *frame operator*, respectively. The next Lemma not only justifies the definitions from above, but also clarifies why the mentioned operators are important for the treatment of frames.

Lemma 1.32. *Let $\{g_i\}_{i \in I}$ be a frame, then*

1. *the analysis operator C is a bounded linear operator,*
2. *the synthesis and analysis operators are adjoint to each other $C^* = D$,*
3. *the frame operator S is a bounded invertible operator on $\mathcal{H}$.*

When working with frames it is of great interest to invert the frame operator as this gives a method to reconstruct from the coefficients

$$f = \sum_{i \in I} \langle f, g_i \rangle S^{-1} g_i = \sum_{i \in I} \langle f, S^{-1} g_i \rangle g_i. \quad (1.77)$$

This equation shows that the family $\{S^{-1} g_i\}_{i \in I}$ can be used to reconstruct vectors in the Hilbert space. It is called *canonical dual frame*.

Theorem 1.33. *Let $\{g_i\}_{i \in I}$ be a frame with frame bounds A, B . Then $\{S^{-1} g_i\}_{i \in I}$ is a frame with frame bounds B^{-1}, A^{-1} .*

While for an ONB the expansion coefficients are the unique set of numbers to expand f in the form

$$f = \sum_{i \in I} c_i g_i, \quad (1.78)$$

for a non-exact frame the coefficients are not unique. In other words, $\ker D \neq \emptyset$ and as a consequence there also exist other dual frames than the canonical, which have the property (1.77).

1.5 The Short-Time Fourier Transform

The Short-Time Fourier Transform (STFT) is a transform defined for finite energy signals $f \in L^2(\mathbb{R}^d)$ in the continuous setting and for finite discrete signals $f \in \mathbb{C}^L$ in a very similar manner. In the following we will make the definitions and statements for the continuous case and denote in brackets if the same can be formulated for the discrete counterpart. As opposed to the Fourier Transform,

the STFT is designed to extract temporal and frequency information at the same time, i.e. time-frequency analysis. The basic idea behind this transform is windowing, i.e. pointwise multiplication with some translation of a prototype window function $T_x g$. If one assumes g to be concentrated around 0, then this translate is concentrated around x . Therefore this pointwise multiplication cuts out the information of the signal f around x . This motivates the definition of the STFT as a function of time and frequency.

Definition 1.34 (STFT). Let $f, g \in L^2(\mathbb{R}^d) (\in \mathbb{C}^L)$. Then the Short-Time Fourier Transform of f with respect to g evaluated at $x, \xi \in \mathbb{R}^d (\in 0, \dots, L-1)$ is defined as

$$V_g f(x, \xi) = \langle f, M_\xi T_x g \rangle. \quad (1.79)$$

The STFT is a function $V_g f : \mathbb{R}^d \times \mathbb{R}^d \rightarrow \mathbb{C}$, from now on we call the domain of this function the *time-frequency plane* (TF-plane) or *phase space*. Furthermore, we frequently use the notation $\lambda = (x, \xi)$ for elements of the time-frequency plane. For time-frequency shifts we adopt the notation $\pi(\lambda) = M_\xi T_x$. The same notations can be used for the finite discrete case by just replacing $\mathbb{R}^d$ by $\mathbb{C}^L$. In this setting the phase space is $\mathbb{Z}_L \times \mathbb{Z}_L = \mathbb{Z}_L^2$, we denote by $\mathbb{Z}_L = \mathbb{Z}/L\mathbb{Z}$ the quotient group in the usual sense. Proofs for all the statements in this section can be found in [54].

Theorem 1.35 (Orthogonality Relations). Let $f_1, f_2, g_1, g_2 \in L^2(\mathbb{R}^d)$. Then

$$\langle V_{g_1} f_1, V_{g_2} f_2 \rangle = \langle f_1, f_2 \rangle \overline{\langle g_1, g_2 \rangle}. \quad (1.80)$$

If $f_1, f_2, g_1, g_2 \in \mathbb{C}^L$, Then

$$\langle V_{g_1} f_1, V_{g_2} f_2 \rangle = L \langle f_1, f_2 \rangle \overline{\langle g_1, g_2 \rangle}. \quad (1.81)$$

These relations give rise to an inversion formula for the STFT. Before we formulate the explicit expression for the continuous case it is necessary to explain vector valued integrals of the form

$$f = \int_{\mathbb{R}^d} g(x) dx, \quad (1.82)$$

where $g : \mathbb{R}^d \rightarrow B$, where B denotes some Banach space. We say that (1.82) holds in the weak sense if and only if for all $h \in B^*$

$$\langle f, h \rangle = \int_{\mathbb{R}^d} \langle g(x), h \rangle dx. \quad (1.83)$$

Using these conventions, we can define the adjoint operator to $V_g : L^2(\mathbb{R}^d) \rightarrow L^2(\mathbb{R}^{2d})$.

Definition 1.36. Let $g \in L^2(\mathbb{R}^d)$. Then we weakly define the operator

$$V_g^* : L^2(\mathbb{R}^{2d}) \rightarrow L^2(\mathbb{R}^d), F \mapsto \int_{\mathbb{R}^{2d}} F(\lambda) \pi(\lambda) g d\lambda. \quad (1.84)$$

In the finite discrete case we define analogously for $g \in \mathbb{C}^L$

$$V_g^* : \mathbb{C}^{L \times L} \rightarrow \mathbb{C}^L, F \mapsto \sum_{\lambda \in \mathbb{Z}_L^2} F(\lambda) \pi(\lambda) g. \quad (1.85)$$

In the following we call V_g *analysis* and V_g^* *synthesis* operator.

Corollary 1.37 (Orthogonality Relations). *Let $f, g, \tilde{g} \in L^2(\mathbb{R}^d)$ with $\langle g, \tilde{g} \rangle \neq 0$. Then*

$$f = \frac{1}{\langle g, \tilde{g} \rangle} \int_{\mathbb{R}^{2d}} V_g f(x, \xi) M_{\xi} T_x \tilde{g} d\xi dx = \frac{1}{\langle g, \tilde{g} \rangle} V_{\tilde{g}}^* V_g f. \quad (1.86)$$

If $f, g, \tilde{g} \in \mathbb{C}^L$ with $\langle g, \tilde{g} \rangle \neq 0$, Then

$$f = \frac{1}{L \langle g, \tilde{g} \rangle} \sum_{x=0}^{L-1} \sum_{\xi=0}^{L-1} V_g f(x, \xi) M_{\xi} T_x \tilde{g} d\xi dx = \frac{1}{L \langle g, \tilde{g} \rangle} V_{\tilde{g}}^* V_g f. \quad (1.87)$$

In the continuous case we see that $f = V_g^* V_g f$ for unit norm g , in the finite discrete case the same formula holds true with an additional factor.

Proposition 1.38. *Let $f, g \in L^2(\mathbb{R}^d)$. Then $V_g f : \mathbb{R}^{2d} \rightarrow \mathbb{C}$ is a uniformly continuous function.*

In particular, the above proposition shows that the image of the operator $V_g : L^2(\mathbb{R}^d) \rightarrow L^2(\mathbb{R}^{2d})$ is only a small part of $L^2(\mathbb{R}^{2d})$. We now derive a formula for projecting an arbitrary function F on phase space to the range of the STFT. First of all, it is important to note that the set of time-frequency shifts do not form a group, since the composition

$$\pi(x_1, \xi_1) \pi(x_2, \xi_2) = e^{-2\pi i x_1 \cdot \xi_2} \pi(x_1 + x_2, \xi_1 + \xi_2), \quad (1.88)$$

is not a time-frequency shift anymore. The additional phase factor is often called *cocycle*. A more detailed investigation of time-frequency shifts leads to the structure of the so-called Heisenberg group [45]. In the following we work without this group structure and have to take care of the phase factors from interchanging time and frequency shifts, as a consequence we need an adapted notion of convolution.

Definition 1.39. Let $F \in L^2(\mathbb{R}^{2d})$ and $G \in L^1(\mathbb{R}^{2d})$ be two functions on phase space. Then the *twisted convolution* is given by

$$F \natural G(\lambda) = \int_{\mathbb{R}^{2d}} F(\mu) G(\lambda - \mu) e^{2\pi i \mu_1 \cdot (\mu_2 - \lambda_2)} d\mu. \quad (1.89)$$

The twisted convolution is, in contrast to the regular one, not commutative. However, we can use this operation to realize the projection onto the range of the STFT operator V_g . The calculations for the finite discrete case are very similar, the only difference being normalization constants and the division by L in the exponent of the complex exponentials. In order to keep the material well readable we do not explicitly formulate all the results for the finite discrete case separately, even though it would be an easy task.

Proposition 1.40. *Let $\|g\| = 1$, such that $V_g g \in L^1(\mathbb{R}^{2d})$, then*

1. $V_g f \natural V_g g = V_g f$ for all $f \in L^2(\mathbb{R}^d)$.
2. Let $F \in L^2(\mathbb{R}^{2d})$, then the twisted convolution $F \natural V_g g$ is the orthogonal projection of F onto the range of V_g .

$$3. \quad V_g V_g^* G = G \natural V_g g$$

Proof. Using formula (1.88), the orthogonality relations (1.80) and the fact that $\pi(-\mu)^{-1} = e^{-2\pi i \mu_1 \mu_2} \pi(\mu)$ we compute

$$\begin{aligned} V_g f \natural V_g g(\lambda) &= \int_{\mathbb{R}^{2d}} \langle f, \pi(\mu)g \rangle \overline{\langle \pi(\lambda)g, \pi(\mu)g \rangle} d\mu \\ &= V_g f(\lambda). \end{aligned} \quad (1.90)$$

From the equality above it is easy to see that the twisted convolution with $V_g g$ is self-adjoint

$$\begin{aligned} \langle G, F \natural V_g g \rangle &= \int_{\mathbb{R}^{2d}} \int_{\mathbb{R}^{2d}} G(\lambda) \overline{F(\mu)} V_g g(\mu - \lambda) e^{2\pi i \lambda_1 (\lambda_2 - \mu_2)} d\mu d\lambda \\ &= \langle G \natural V_g g, F \rangle. \end{aligned} \quad (1.91)$$

This finishes the proof of the second statement, since

$$\langle F - F \natural V_g g, F \natural V_g g \rangle = 0. \quad (1.92)$$

The third statement easily follows by simply computing the composition of V_g and V_g^*

$$V_g V_g^* G(\lambda) = \int_{\mathbb{R}^d} G(\mu) \langle \pi(\mu)g, \pi(\lambda)g \rangle d\mu = G \natural V_g g(\lambda). \quad (1.93)$$

□

In the above formulas it becomes apparent that the *ambiguity function* of a window g defined by $V_g g$ plays an important role, in Section 3.3 we will construct windows, such that their ambiguity functions satisfies some optimal concentration properties.

Corollary 1.41. *Let $\|g\| = 1$. Then V_g is an isometry from $L^2(\mathbb{R}^d)$ into $L^2(\mathbb{R}^{2d})$. The adjoint operator V_g^* is an isometry from $V_g(L^2(\mathbb{R}^d))$ onto $L^2(\mathbb{R}^d)$. Furthermore, the operator $V_g^* : L^2(\mathbb{R}^{2d}) \rightarrow L^2(\mathbb{R}^d)$ has norm 1.*

The proof is an easy consequence of the collected properties of V_g^* and the orthogonality relations.

1.5.1 Localization operators

The general idea of localization operators is to multiply the STFT of a function f with some weight function on phase space before doing the re-synthesis.

Definition 1.42. Given $m \in L^\infty(\mathbb{R}^2)(\in \mathbb{C}^{L \times L})$ and $g \in L^2(\mathbb{R}^d)(\in \mathbb{C}^L)$. Then define the corresponding *STFT multiplier*

$$M_{m,g} = V_g^* \circ \mathcal{M}_m \circ V_g, \quad (1.94)$$

where $\mathcal{M}_m$ denotes the multiplication with the weight function m .

Localization operators have received considerable attention in the literature [20, 109]. It is easy to see that for a bounded weight function and finite energy window the corresponding multiplier is a bounded operator $M_{m,g} : L^2(\mathbb{R}^d) \rightarrow L^2(\mathbb{R}^d)$. Later in the text, compact localization operators will play an important role, in the literature results can be found in [6, 43].

Definition 1.43. Let $\mathcal{H}$ be a Hilbert space. A sequence $f_n \in \mathcal{H}$ converges to 0 weakly if

$$\langle f_n, h \rangle \rightarrow 0, \quad \text{for all } h \in \mathcal{H}. \quad (1.95)$$

Definition 1.44. Let $\mathcal{H}_1, \mathcal{H}_2$ be Hilbert spaces. Then the operator $T : \mathcal{H}_1 \rightarrow \mathcal{H}_2$ is called *compact* if for any weakly convergent sequence $f_n \in \mathcal{H}_1$

$$\|Tf_n\| \rightarrow 0, \quad \text{as } n \rightarrow \infty. \quad (1.96)$$

Compact operators have an important and useful property.

Theorem 1.45 (Spectral theorem). *Let $\mathcal{H}$ be a separable Hilbert space and $T : \mathcal{H} \rightarrow \mathcal{H}$ be a self-adjoint compact operator. Then there exists an orthonormal basis of $\mathcal{H}$ consisting of eigenvectors of T , such that the corresponding eigenvalues are real and have no accumulation points away from 0.*

The following theorem is well known in the community, nevertheless we include the proof since it is rather easy.

Theorem 1.46. *Let $g \in L^2(\mathbb{R}^d)$ and $m \in L^\infty(\mathbb{R}^{2d})$, such that $\lim_{|\lambda| \rightarrow \infty} m(\lambda) = 0$. Then the multiplier $M_{m,g}$ is compact.*

Proof. Let us first assume that $\text{supp } m \subseteq C$ for some compact set C . Let $f_n \in L^2(\mathbb{R}^d)$ be weakly convergent to 0, then we estimate

$$\begin{aligned} |(\mathcal{M}_m \circ V_g)f_n(\lambda)| &= |\langle f_n, \pi(\lambda)g \rangle m(\lambda)| \\ &\leq \|f_n\|_2 \|g\|_2 \|m\|_\infty. \end{aligned} \quad (1.97)$$

Therefore, using the dominated convergence theorem we find that $\|(\mathcal{M}_m \circ V_g)f_n\| \rightarrow 0$, as $n \rightarrow \infty$. Furthermore, the operator V_g^* has unit operator norm, therefore we can conclude that

$$\|M_{m,g}f_n\| \rightarrow 0, \quad n \rightarrow \infty, \quad (1.98)$$

and therefore $M_{m,g}$ is compact.

Now let $m \in L^\infty(\mathbb{R}^{2d})$. We split the weight function $m = m_1 + m_2$, such that $\text{supp } m_1 \subseteq C$, where C is compact, such that $\|(\mathcal{M}_{m_2} \circ V_g)f_n\| < \varepsilon/2$ for some given $\varepsilon > 0$ and for all $n \in \mathbb{N}$. Choosing such a set C is possible since m is essentially bounded and the norm of $(\mathcal{M}_m \circ V_g)f_n$ is bounded on the TF-plane. Therefore the norm of $(\mathcal{M}_{m_2} \circ V_g)f_n$ can be made arbitrarily small for large enough C . Using the arguments for a compactly supported weight function we can find some n sufficiently large, such that

$$\|(\mathcal{M}_{m_2} \circ V_g)f_n\| = \|(\mathcal{M}_{m_1} \circ V_g)f_n\| + \|(\mathcal{M}_{m_2} \circ V_g)f_n\| < \varepsilon/2 + \varepsilon/2. \quad (1.99)$$

Using the fact that V_g^* has norm 1 we find that $M_{m,g}$ maps weakly convergent sequences to norm convergent sequences. $\square$

1.5.2 STFT of distributions

Up to now we only considered the STFT of functions $f \in L^2(\mathbb{R}^d)$. However, the defining scalar product can easily be extended to tempered distributions assuming that the window $g \in \mathcal{S}(\mathbb{R}^d)$. This leads to a pointwise definition of the STFT for tempered distributions without control over the norm of this function on phase space. Controlling this norm leads to the definition of *modulation spaces* [54].

Definition 1.47 (Modulation spaces). For some $g \in \mathcal{S} \setminus \{0\}$ and $1 \leq p \leq \infty$ the modulation space norm for some $f \in \mathcal{S}'(\mathbb{R}^d)$ is defined as

$$\|f\|_{M^p(\mathbb{R}^d)} = \|V_g f\|_{L^p(\mathbb{R}^d)}. \quad (1.100)$$

It is important to note that the definition of the modulation spaces is independent of the chosen window function. The modulation space $M^1(\mathbb{R}^d) = \mathcal{S}_0(\mathbb{R}^d)$ plays a special role, it is also called *Feichtinger's algebra* [37]. It is the minimal Fourier invariant Banach space and can therefore be used to establish a theory of distributions via duality that includes the Fourier Transform but is topologically easier than the space of tempered distributions. Most of the results from $L^2(\mathbb{R}^d)$ can be carried over to the modulation space setting including the inversion formula, c.f. [54] for an excellent introduction.

1.6 Gabor frames

In Section 1.5 we defined the STFT of a signal $f \in L^2(\mathbb{R}^d)$ as inner products with the continuous family $\{M_\xi T_x g\}_{\xi, x \in \mathbb{R}^d}$, $g \in L^2(\mathbb{R}^d)$. While we showed that this leads to an invertible transform, these coefficients are highly redundant, as $L^2(\mathbb{R}^d)$ is a separable Hilbert space and countably many coefficients should suffice. In this section we discretize the STFT and ask the question whether it is possible or not to reconstruct $f \in L^2(\mathbb{R}^d)$ from the discrete set of coefficients

$$\{\langle f, \pi(\lambda)g \rangle\}_{\lambda \in \Lambda}, \quad (1.101)$$

where Λ is some discrete subset of the time-frequency plane. We note here that this set of coefficients can be seen as $V_g f$ sampled at the points Λ . The coefficients from (1.101) are called *Gabor coefficients* of f with window g . For an introduction into Gabor analysis we refer to [15, 54], where the details and proofs for the results discussed in the following can be found.

As in the previous section we formulate the results for the continuous case and the finite discrete case. In the literature most of the results are originally formulated in the continuous context, for discrete treatment and algorithmic aspects we refer to [42]. A connection of the continuous theory with its discrete counterpart can be found in [41]. We show some of the proofs for the finite discrete setting to find the correct normalization factors.

The original idea for Gabor transformations was outlined by the Hungarian engineer and Nobel price winner Dénes Gábor in 1946 [50]. He originally proposed to use the family

$$\{M_k T_l g_0\}_{k, l \in \mathbb{Z}}, \quad (1.102)$$

where g_0 denotes the standard Gaussian function as defined in (1.27).

Definition 1.48. Let $g \in L^2(\mathbb{R}^d)(\in \mathbb{C}^L)$ and Λ a discrete subset of $\mathbb{R}^{2d}(\mathbb{C}^{2L})$. Then the corresponding *Gabor family* is defined as

$$\mathcal{G}(g, \Lambda) = \{\pi(\lambda)g\}_{\lambda \in \Lambda}. \quad (1.103)$$

The *Gabor frame operator* is given by

$$S_{g,\Lambda}f = \sum_{\lambda \in \Lambda} \langle f, \pi(\lambda)g \rangle \pi(\lambda)g. \quad (1.104)$$

Furthermore, the corresponding analysis and synthesis operators are given by

$$C_{g,\Lambda}f = \{\langle f, \pi(\lambda)g \rangle\}_{\lambda \in \Lambda} \quad (1.105)$$

$$D_{g,\Lambda}c = \sum_{\lambda \in \Lambda} c_\lambda \pi(\lambda)g \quad (1.106)$$

If a Gabor family forms a frame in the for $L^2(\mathbb{R}^d)$ or $\mathbb{C}^L$ we call it *Gabor frame*. Up to here the set Λ did in general not satisfy any restrictions, therefore we call the resulting systems *irregular Gabor families*. There are necessary conditions for such systems to form frames [83]. From now on we restrict ourselves to the special case where Λ is a discrete subgroup of phase space. Recall that the finite discrete time-frequency plane is given by $\mathbb{Z}_L \times \mathbb{Z}_L = \mathbb{Z}_L^2$.

Definition 1.49. A discrete subgroup $\Lambda \subseteq \mathbb{R}^{2d}(\mathbb{Z}_L^2)$ is called *lattice* and is denoted by $\Lambda \leq \mathbb{R}^{2d}(\mathbb{Z}_L^2)$.

It turns out that Gabor systems are particularly rich in structure when the sampling set Λ is a lattice. As discussed in Section 1.4, the frame operator of a system plays a crucial role. It is invertible if and only if the Gabor family $\mathcal{G}(g, \Lambda)$ is a frame. The canonical dual frame of the Gabor system has a property that makes it attractive from a structural and applied point of view.

Proposition 1.50. Let $\Lambda \leq \mathbb{R}^{2d}(\mathbb{Z}_L^2)$ and $g \in L^2(\mathbb{R}^d)(\in \mathbb{C}^L)$. Then

$$S_{g,\Lambda}\pi(\lambda)f = \pi(\lambda)S_{g,\Lambda}f, \quad \text{for all } \lambda \in \Lambda. \quad (1.107)$$

An important consequence of this proposition is that the canonical dual of a Gabor frame is a Gabor frame itself and given by $\mathcal{G}(S_{g,\Lambda}^{-1}g, \Lambda)$. This reduces the computational effort tremendously, as one only needs to apply the inverse frame operator to the original window function instead of the whole system.

Before we turn our attention to general lattices in the TF-plane in Chapter 2, we assume Λ to be a separable discrete subgroup of the phase space, i.e. $\Lambda = a\mathbb{Z} \times b\mathbb{Z}$, where $a, b > 0$ in the continuous case and $\Lambda = a\mathbb{Z}_L \times b\mathbb{Z}_L$, for $a, b|L$ in the finite discrete case. Furthermore, we use the notation $S_{g,a,b}$ for the Gabor frame operator connected to the system defined by the window g and the step parameters a and b .

Definition 1.51. Let $\Lambda = a\mathbb{Z} \times b\mathbb{Z}$ be a lattice in $\mathbb{R}^2$ for some $a, b > 0$, then define

1. the *volume* as $\text{vol}(\Lambda) = ab$.
2. the *redundancy* as $\text{red}(\Lambda) = \text{vol}(\Lambda)^{-1}$.

Let $\Lambda = a\mathbb{Z}_L \times b\mathbb{Z}_L$ for $a, b|L$ be a lattice in $\mathbb{Z}_L^2$, then define

1. the *volume* as $\text{vol}(\Lambda) = ab/L$.
2. the *redundancy* as $\text{red}(\Lambda) = \text{vol}(\Lambda)^{-1}$.

Before discussing discretization and structure of Gabor frame operators we mention an important result for Gabor systems with standard Gaussian windows.

Theorem 1.52. *Let $a, b > 0$. Then the family $\mathcal{G}(g_0, a, b)$ is a frame for $L^2(\mathbb{R}^d)$ if and only if $ab < 1$.*

1.6.1 From continuous to discrete Gabor frames

Results about the relations between the finite discrete case and the continuous case can be found in [66, 70, 90], in the following we collect some results from the two latter mentioned research papers. The main ingredient in going from continuous to finite discrete is *sampling* and *periodization*.

Definition 1.53. We define the *sampling* and *periodization* operator respectively as

$$\mathcal{P}_L : l^1(\mathbb{Z}) \rightarrow \mathbb{C}^L, \quad \mathcal{P}_L g(l) = \sum_{m \in \mathbb{Z}} g(l + mL), \quad (1.108)$$

$$\mathcal{S}_\alpha : \mathcal{S}_0(\mathbb{R}) \rightarrow l^1(\mathbb{Z}), \quad \mathcal{S}_\alpha(l) = f(l\alpha). \quad (1.109)$$

Of particular interest is the composition of the operators defined above

$$\mathcal{P}_L \mathcal{S}_{\sqrt{L}} f(l) = \sum_{m \in \mathbb{Z}} f(l/\sqrt{L} + m\sqrt{L}), \quad (1.110)$$

which shows the transition from $\mathcal{S}_0$ to $\mathbb{C}^L$. The specific sampling rate of $\sqrt{L}$ makes sure that sampling and periodization are simultaneously increased at the same rate when increasing L .

Proposition 1.54. *Given a signal length $L \in \mathbb{N}$ and a function $f \in \mathcal{S}_0$. Then*

$$\mathcal{P}_L \mathcal{S}_{\sqrt{L}} \mathcal{F}_{\mathbb{R}} f = \mathcal{F}_{\mathbb{C}} \mathcal{P}_L \mathcal{S}_{\sqrt{L}} f. \quad (1.111)$$

The proof is an application of Poisson's summation formula (1.23), the technical details can be found in [70]. It is also possible to relate Gabor frames of continuous functions with Gabor frames in the finite discrete setting.

Theorem 1.55. *Let $\alpha, \beta > 0$, $g \in \mathcal{S}_0$ and $L \in \mathbb{N}$ be the signal length. Furthermore, choose $a, b|L$, such that $\alpha\beta = ab/L$. Furthermore assume that the Gabor system $\mathcal{G}(g, \alpha, \beta)$ establishes a frame for $L^2(\mathbb{R})$ with canonical dual window $\tilde{g}$. Then the finite discrete system $\mathcal{G}(\sqrt{\alpha/a} \mathcal{P}_L \mathcal{S}_{\alpha/a} g, a, b)$ establishes a frame for $\mathbb{C}^L$ with the same frame bounds and with canonical dual window $\sqrt{\alpha/a} \mathcal{P}_L \mathcal{S}_{\alpha/a} \tilde{g}$.*

Given continuous step parameters α and β , the choice of the discrete sampling parameters $a = \sqrt{L}\alpha$ and $b = \sqrt{L}\beta$, similar to (1.110), is of particular

interest, as it treats time and frequency step equally. This choice admits a sampling pattern that is geometrically the same as the continuous one and has the same redundancy.

These sampling and periodization techniques motivate the definition of the finite discrete Gaussian as a sampled periodized version of g_0 , as already mentioned in 1.3.2. In Chapter 3 we will discuss another method to obtain a discrete set of Hermite functions based on STFT multipliers. In Section 3.3.3 we will also show an expansion of the periodized and sampled Gaussian function in terms of the discrete Hermite functions defined above.

1.6.2 Structure of Gabor systems

Gabor frames and in particular the Gabor frame operator have a lot of structure. We list the most important results for separable lattices in the following, a more in depth discussion and the proofs can be found in [54].

Theorem 1.56 (Walnut representation). *Let $g \in W(\mathbb{R}^d)$. Then for any $a, b > 0$ the continuous Gabor frame operator satisfies*

$$S_{g,a,b}f(t) = \frac{1}{b^d} \sum_{n \in \mathbb{Z}^d} \left(\sum_{k \in \mathbb{Z}^d} \bar{g}(t - n/b - ak)g(t - ak) \right) T_{n/b}f(t). \quad (1.112)$$

Let $g \in \mathbb{C}^L$. Then for any $a, b|L$ the discrete Gabor frame operator satisfies

$$S_{g,a,b}f(l) = \frac{L}{b} \sum_{n=0}^{b-1} \left(\sum_{k=0}^{L/a-1} \bar{g}(l - ak - Ln/b)g(l - ak) \right) T_{Ln/b}f(l). \quad (1.113)$$

We only give a proof for the discrete formula as it contains the crucial idea and yields the correct normalization. The technical details for the continuous version can be found in the original research paper [106] or in [54].

Proof of the discrete Walnut representation. By the definition of the discrete Gabor frame operator and the discrete Poisson summation formula we find

$$\begin{aligned} S_{g,a,b}f(l) &= \sum_{k=0}^{L/a-1} \sum_{n=0}^{L/b-1} \langle f, M_{bn}T_{ak}g \rangle M_{bn}T_{ak}g(l) \\ &= \sum_{k=0}^{L/a-1} \sum_{n=0}^{L/b-1} \sqrt{L} \mathcal{F}(f T_{ak}\bar{g})(bn) e^{2\pi i b n l / L} T_{ak}g(l) \\ &= \frac{L}{b} \sum_{k=0}^{L/a-1} \sum_{n=0}^{b-1} T_{Ln/b}f(l) T_{ak+Ln/b}\bar{g}(l) T_{ak}g. \end{aligned} \quad (1.114)$$

The assertion follows by grouping the right terms in the sum above. $\square$

The terms that appear in the Walnut representation are periodic

$$G_n(l) = \sum_{k=0}^{L/a-1} \bar{g}(l - ak - Ln/b)g(l - ak). \quad (1.115)$$

Due to this periodicity we can find more structure in the frame operator that makes the representation more symmetric with respect to time and frequency shifts.

Theorem 1.57 (Janssen representation). *Let $a, b \in \mathbb{R}^+$ ($a, b|L$) and $g \in \mathcal{S}_0(\mathbb{R}^d)(\in \mathbb{C}^L)$. Then the frame operator assumes the form*

$$S_{g,a,b} = \text{vol}(\Lambda)^{-1} \sum_{n,k \in \mathbb{Z}^d} \langle g, M_{k/a} T_{n/b} g \rangle M_{k/a} T_{n/b} \quad (1.116)$$

As above we only present the proof for the finite discrete case. The continuous case has essentially the same proof, we need the technical assumption $g \in \mathcal{S}_0$ to satisfy the condition of Tolimieri-Orr [100].

Proof of the discrete Janssen representation. We first compute the Fourier Transform of the terms defined in (1.115) using Poisson summation

$$\begin{aligned} \hat{G}_n(l) &= \left((g T_{Ln/b} \bar{g}) * \left(\sum_{k=0}^{L/a-1} \delta_{ak} \right) \right)^\wedge(l) \\ &= \frac{\sqrt{L}}{a} \langle g, M_l T_{Ln/b} g \rangle \left(\sum_{k=0}^{a-1} \delta_{Lk/a}(l) \right). \end{aligned} \quad (1.117)$$

Therefore, the inverse Fourier Transform yields

$$G_n(l) = \frac{1}{a} \sum_{k=0}^{a-1} \langle g, M_{Lk/a} T_{Ln/b} g \rangle \cdot e^{2\pi i l k / a}, \quad (1.118)$$

and by plugging this into the Walnut representation we obtain the Janssen representation. $\square$

Theorem 1.58 (Wexler-Raz). *Let $a, b \in \mathbb{R}^+$ ($a, b|L$) and denote by Λ the corresponding separable lattice. Furthermore, assume that the synthesis operators $D_{g,\Lambda}$ and $D_{\gamma,\Lambda}$ are bounded. Then the following are equivalent*

1. $\sum_{\lambda \in \Lambda} \langle f, \pi(\lambda)g \rangle \pi(\lambda)\gamma = f$ for all $f \in L^2(\mathbb{R}^d)$.
2. $\text{red}(\Lambda) \langle \gamma, M_{l/a} T_{n/b} g \rangle = \delta(l)\delta(n)$ for $l, n \in \mathbb{Z}^d$.

An immediate consequence of this theorem is the following characterization of tight frames.

Corollary 1.59. *Let $a, b \in \mathbb{R}^+$ ($a, b|L$) be given. Then the Gabor system $\mathcal{G}(g, a, b)$ is a tight frame if and only if $\mathcal{G}(g, 1/b, 1/a)$ is an orthogonal system.*

1.6.3 Multi-window Gabor frames

This section only shortly mentions multiwindow Gabor frames. Most of the results from Gabor frames can be used in this framework, results and computational aspects can be found in [111–114].

Definition 1.60. Given a finite number of windows $\mathbf{g} = \{g_r\}_{1 \leq r \leq R}$ in the signal space $L^2(\mathbb{R}^d) (\mathbb{C}^L)$ and Λ a lattice in the TF plane, then the *multiwindow Gabor system* is defined as

$$\mathcal{G}(\mathbf{g}, \Lambda) = \{\pi(\lambda)g_r : 1 \leq r \leq R, \lambda \in \Lambda\}. \quad (1.119)$$

It is then easy to see that the frame operator associated with the multi-window Gabor system is just a sum of single window Gabor frame operators

$$S_{\mathbf{g}, \Lambda} = \sum_{r=1}^R S_{g_r, \Lambda}. \quad (1.120)$$

This observation shows that multi-window Gabor frame operators inherit most of the structure from the single window case. We will use multiwindow Gabor systems later in Section 2.3.1 to treat Gabor transforms on non-separable lattices.

Chapter 2

Gabor transforms on general lattices

In Section 1.6 we described the Gabor transform on product lattices. Now we want to generalize these results to any given subgroup of the TF-plane. We will harness the theory of metaplectic operators and symplectic matrices to reduce the general case to the known statements. There are differences in several places when switching between the continuous and the discrete theory. The reason for the big difference in the structure is the inherent unit step, that one has in the finite discrete case as well as the circular structure of the finite discrete TF-plane.

2.1 Lattices in the TF-plane

We start by looking at lattices in the continuous phase space, which is rather easy.

Proposition 2.1. *Given a lattice $\Lambda \leq \mathbb{R}^{2d}$. Then there exists a matrix $A \in \mathbb{R}^{2d \times 2d}$, such that $\Lambda = A \cdot \mathbb{Z}^{2d}$.*

This proposition follows easily from the simple fact, that $\Lambda \leq \mathbb{R}^{2d}$ is generated by at most $2d$ vectors in $\mathbb{R}^{2d}$.

Definition 2.2. Given a lattice $\Lambda \leq \mathbb{R}^{2d}$ with generator matrix A . Then

1. we call $1/\det(A)$ the *redundancy* of the Λ .
2. we call $\det(A) = \text{vol}(\Lambda)$ the *volume* of the Λ .

The following proposition allows us to view any lattice in $\mathbb{R}^{2d}$ as a rotation/inflection of a lattice generated by a lower triangular matrix.

Proposition 2.3. *For any lattice $\Lambda \leq \mathbb{R}^{2d}$ there exists a lower triangular matrix L and an orthogonal matrix Q , such that $\Lambda = LQ\mathbb{Z}^{2d}$.*

Proof. The proof is an easy consequence from the fact that any matrix A can be written as $A = LQ$ with lower triangular matrix L and an orthogonal matrix Q by Gram-Schmidt orthogonalization. $\square$

Now we turn our attention to subgroups of the cyclic group $\mathbb{Z}_L \times \mathbb{Z}_L$, i.e. the TF-plane in the finite discrete case. We largely follow [57], where my co-authors and I investigate subgroups of the form $\mathbb{Z}_n \times \mathbb{Z}_m$ for some given m, n are investigated. We restrict ourselves here to the case where $m = n = L$, which is slightly more restrictive but the general ideas are the same. For treating Gabor transforms on non-separable lattices it suffices to consider lattices in $\mathbb{Z}_L^2$. The next theorem describes the structure of lattices in $\mathbb{Z}_L^2$, an illustration can be found in Figure 2.1.

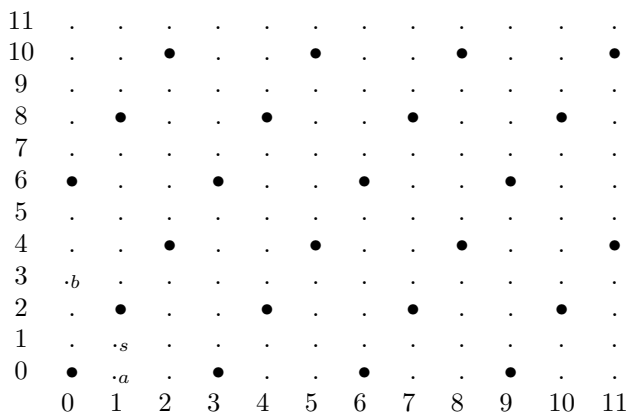


Figure 2.1: A subgroup of $\mathbb{Z}_{12}^2$ with $a = 1$, $b = 6$ and $s = 2$

Theorem 2.4. *Let $L \in \mathbb{N}$ and define*

$$I_L := \{(a, b, t) \in \mathbb{N}^2 \times \mathbb{N}_0 : a, b | L, 0 \leq t \leq \gcd(L/a, b) - 1\}. \quad (2.1)$$

For $(a, b, t) \in I_L$ define

$$\Lambda_{a,b,t} := \{(ia, itb/\gcd(L/a, b) + jb)^T : 0 \leq i \leq L/a - 1, 0 \leq j \leq L/b - 1\}, \quad (2.2)$$

then $\Lambda_{a,b,t}$ is a subgroup of order $\frac{L^2}{ab}$ of $\mathbb{Z}_L^2$ and the map $(a, b, t) \mapsto \Lambda_{a,b,t}$ is a bijection between the set I_L and the set of subgroups of $\mathbb{Z}_L^2$.

Proof. We show the one to one mapping between the set I_L and the subgroups of $\mathbb{Z}_L^2$. Let us first consider a subgroup $\Lambda \leq \mathbb{Z}_L^2$ and the natural projection $\pi_1 : \mathbb{Z}_L^2 \rightarrow \mathbb{Z}_L$ given by $\pi_1(x, y) = x$. Then $\pi_1(\Lambda)$ is a subgroup of $\mathbb{Z}_L$ and there exists a unique divisor a of L such that $\pi_1(\Lambda) = \langle a \rangle := \{ia : 0 \leq i \leq L/a - 1\}$. Let $s \geq 0$ be minimal such that $(a, s) \in \Lambda$.

Furthermore, consider the natural inclusion $\iota_2 : \mathbb{Z}_L \rightarrow G$ given by $\iota_2(y) = (0, y)$. Then $\iota_2^{-1}(\Lambda)$ is a subgroup of $\mathbb{Z}_L$ and there exists a unique divisor b of L such that $\iota_1^{-1}(\Lambda) = \langle b \rangle$.

We show that $\Lambda = \{(ia, is + jb)^T : i, j \in \mathbb{Z}\}$. Indeed, for every $i, j \in \mathbb{Z}$, $(ia, is + jb)^T = i(a, s)^T + j(0, b)^T \in \Lambda$. On the other hand, for every $(u, v)^T \in \Lambda$ one has $u \in \pi_1(\Lambda)$ and hence there is $i \in \mathbb{Z}$ such that $u = ia$. We obtain $(0, v - is)^T = (u, v)^T - i(a, s)^T \in \Lambda$, $v - is \in \iota_2^{-1}(\Lambda)$ and there is $j \in \mathbb{Z}$ with $v - is = jb$.

Here a necessary condition is that $(0, sL/a)^T \in \Lambda$ (obtained for $i = L/a$, $j = 0$), that is $b \mid sL/a$, equivalent to $b/\gcd(L/a, b) \mid s$. Clearly, if this is

verified, then for the above representation of Λ it is enough to take the values $0 \leq i \leq L/a - 1$ and $0 \leq j \leq L/b - 1$.

Also, dividing s by b we have $s = bq + r$ with $0 \leq r < b$ and $(a, r) = (a, s) - q(0, b) \in \Lambda$, showing that $s < b$, by its minimality. Hence $s = tb/\gcd(L/a, b)$ with $0 \leq t \leq \gcd(L/a, b) - 1$. Thus we obtain the given representation.

Conversely, every $(a, b, t) \in I_L$ generates a subgroup $\Lambda_{a,b,t}$ of order $L^2/(ab)$ of $\mathbb{Z}_L \times \mathbb{Z}_L$ and the proof is complete. $\square$

Remark 2.5. 1. The bijection in the theorem above allows us to construct all the subgroups of the finite discrete phase space in a very convenient manner considering elements of the set I_L . Furthermore, it allows us to construct all the subgroups of a given order.

2. The *redundancy* of the subgroup $\Lambda_{a,b,t}$ is given by $L/(ab)$.
3. The *volume* of the subgroup $\Lambda_{a,b,t}$ is given by ab/L .
4. The above theorem implies that for any given lattice $\Lambda_{a,b,t}$ there exists a uniquely determined lower triangular matrix A , such that

$$\Lambda_{a,b,t} = A \cdot \mathbb{Z}_L^2 = \begin{pmatrix} a & 0 \\ s & b \end{pmatrix} \cdot \mathbb{Z}_L^2, \quad (2.3)$$

where

$$s = tb/\gcd(L/a, b). \quad (2.4)$$

5. The normal form is closely related to the *Hermite normal form* for matrices over the integers. Any matrix $A \in \mathbb{Z}^{2 \times 2}$ can be written as $A = HU$, where $|\det U| = 1$ and H is lower triangular with positive diagonal entries and $0 \leq H_{2,1} < H_{2,2}$.

From now on we call s the shear, as it determines by which amount the two neighboring non-empty columns of the lattice have to be shifted with respect to each other. The normal form established above can also be changed into upper triangular form. The next statement shows how to interchange between these two normal forms.

Proposition 2.6. *Given a subgroup $\mathbb{Z}_L^2$ in normal form, i.e. given a, b and s . Then the following representations are equivalent*

$$\begin{pmatrix} a & 0 \\ s & b \end{pmatrix} \cdot \mathbb{Z}_L^2 = \begin{pmatrix} \tilde{a} & \tilde{s} \\ 0 & \tilde{b} \end{pmatrix} \cdot \mathbb{Z}_L^2, \quad (2.5)$$

where $\tilde{b} = \gcd(b, s)$, $\tilde{a} = ab/\gcd(b, s)$. Furthermore, we use Bézout's identity to represent $k_1s + k_2b = \gcd(b, s)$, then $\tilde{s} = k_1a$.

Proof. By computation one can verify that

$$\begin{pmatrix} a & 0 \\ s & b \end{pmatrix} \cdot \begin{pmatrix} b/\gcd(b, s) & k_1 \\ -s/\gcd(b, s) & k_2 \end{pmatrix} = \begin{pmatrix} \tilde{a} & \tilde{s} \\ 0 & \tilde{b} \end{pmatrix} \quad (2.6)$$

The second matrix has determinant 1 and therefore is invertible. The assertion follows because $Q \cdot \mathbb{Z}_L^2 = \mathbb{Z}_L^2$ for any invertible matrix. $\square$

So far we know only that any subgroup can be written as a linear combination of two basis vectors. We now establish a result, that sheds more light on the subgroups and is mostly interesting from a group theoretic point of view.

Definition 2.7. Let Λ be a finite group. Then define the *exponent* as

$$\exp \Lambda = \min \{n \in \mathbb{N} : \lambda^n = 0 \quad \forall \lambda \in \Lambda\}. \quad (2.7)$$

Theorem 2.8. 1. The exponent of the subgroup $\Lambda_{a,b,t}$ is given by

$$\exp \Lambda_{a,b,t} = \frac{L}{\gcd(b, a, s)}. \quad (2.8)$$

2. The subgroup $\Lambda_{a,b,t}$ is cyclic if and only if $\gcd(L/a, L/b, Ls/ab) = 1$.

3. The subgroup $\Lambda_{a,b,t}$ is separable if and only if $t = 0$ and $\Lambda_{a,b,0} = \mathbb{Z}_{L/a} \times \mathbb{Z}_{L/b}$. Here $\Lambda_{a,b,0}$ is cyclic if and only if $\gcd(L/a, L/b) = 1$.

Proof. The first assertion follows from the fact that the group is generated by the two elements $(a, s)^T$ and $(0, b)^T$. Therefore the exponent of the group is the least common multiple of the orders of these two elements. The order of $(0, b)^T$ is L/b . To compute the order of $(a, s)^T$ note that $L|rs$ if and only if $L/\gcd(L, s)|r$. Thus the order of $(a, s)^T$ is $\text{lcm}(L/\gcd(L, s), L/a)$ and therefore the formula for the exponent is

$$\begin{aligned} \exp \Lambda_{a,b,t} &= \text{lcm} \left(\frac{L}{b}, \frac{L}{a}, \frac{L}{\gcd(L, s)} \right) \\ &= \frac{L}{\gcd(a, b, s, L)} = \frac{L}{\gcd(a, b, s)}. \end{aligned} \quad (2.9)$$

To proof the second statement we note that $\Lambda_{a,b,t}$ is cyclic if and only if its exponent equals its order, that is $\frac{L}{\gcd(a, b, s)} = \frac{L^2}{ab}$. This is equivalent to $\gcd(L/a, L/b, Ls/ab) = 1$.

Finally, the third statement is a direct consequence of Theorem 2.4. $\square$

2.1.1 Counting subgroups of $\mathbb{Z}_L^2$

The normal form derived in Theorem 2.4 can not only be used to obtain the group theoretical results from above but can also be used to count the number of subgroups of a given order. Subgroups of $\mathbb{Z} \times \mathbb{Z}$ (sublattices of the two dimensional lattice) and associated counting functions were considered by several authors in pure and applied mathematics. It is known, for example, that the number of subgroups of index n in $\mathbb{Z} \times \mathbb{Z}$ is $\sigma(n)$, the sum of the (positive) divisors of n . See, e.g., [52], [115].

From now on we denote the number of subgroups of a group Λ by $s(\Lambda)$ and the number of subgroups of given order k by $s_k(\Lambda)$. By $s(m, n)$ and $s_k(m, n)$ we denote the number of subgroups of the group $\mathbb{Z}_m \times \mathbb{Z}_n$ and the number of subgroups of order k of this group, respectively. It is known that for every finite Abelian group the problem of counting all subgroups and the subgroups of a given order reduces to p -groups, which follows from the properties of the subgroup lattice of the group (see R. Schmidt [87], M. Suzuki [97]). In

particular, for $G = \mathbb{Z}_m \times \mathbb{Z}_n$ this can be formulated as follows. G is an Abelian group of rank two since $G \simeq \mathbb{Z}_u \times \mathbb{Z}_v$, where $u = \text{lcm}(m, n)$, $v = \text{gcd}(m, n)$. Let $u = p_1^{a_1} \cdots p_r^{a_r}$ and $v = p_1^{b_1} \cdots p_r^{b_r}$ be the prime power factorizations of u and v , respectively, where $a_j \geq b_j \geq 0$ ($1 \leq j \leq r$). Then

$$s(m, n) = \prod_{j=1}^r s(\mathbb{Z}_{p_j^{a_j}} \times \mathbb{Z}_{p_j^{b_j}}), \quad (2.10)$$

and

$$s_k(m, n) = \prod_{j=1}^r s_{k_j}(\mathbb{Z}_{p_j^{a_j}} \times \mathbb{Z}_{p_j^{b_j}}), \quad (2.11)$$

where $k = k_1 \cdots k_r$ and $k_j = p_j^{c_j}$ with some exponents $0 \leq c_j \leq a_j + b_j$ ($1 \leq j \leq r$).

Now let $\mathbb{Z}_{p^a} \times \mathbb{Z}_{p^b}$ be a p -group of rank two, where $1 \leq b \leq a$. Then one has the simple explicit formulae:

$$s(\mathbb{Z}_{p^a} \times \mathbb{Z}_{p^b}) = \frac{(a - b + 1)p^{b+2} - (a - b - 1)p^{b+1} - (a + b + 3)p + (a + b + 1)}{(p - 1)^2}, \quad (2.12)$$

$$s_{p^c}(\mathbb{Z}_{p^a} \times \mathbb{Z}_{p^b}) = \begin{cases} \frac{p^{c+1} - 1}{p - 1}, & c \leq b \leq a, \\ \frac{p^{b+1} - 1}{p - 1}, & b \leq c \leq a, \\ \frac{p^{a+b-c+1} - 1}{p - 1}, & b \leq a \leq c \leq a + b. \end{cases} \quad (2.13)$$

Formula (2.12) was derived by G. Călugăreanu [12, Sect. 4] and recently by J. Petrillo [81, Prop. 2] using Goursat's lemma for groups. M. Tărnăuceanu [98, Th. 3.3] deduced (2.12) and (2.13) by a method based on properties of certain attached matrices.

Therefore, $s(m, n)$ and $s_k(m, n)$ can be computed using (2.10), (2.12) and (2.11), (2.13), respectively. We now give a more compact form for the special case $L = m = n$. We make the convention $s(L) := s(L, L)$ and $s_k(L) := s_k(L, L)$. In [57] the more general formulas can be found.

Theorem 2.9. *Let $L \in \mathbb{N}$. Then*

1. *The number of subgroups of $\mathbb{Z}_L^2$ is given by*

$$s(L) = \sum_{a, b | L} \text{gcd}(a, b). \quad (2.14)$$

2. *The number of subgroups of $\mathbb{Z}_L^2$ of given order $k | L^2$ is given by*

$$s_k(L) = \sum_{\substack{a, b | L \\ La/b = k}} \text{gcd}(a, b). \quad (2.15)$$

Proof. For the first statement we simply count the number of elements in the set $I_{a, b, t}$. This leads to

$$s(L) = \sum_{a, b | L} \sum_{0 \leq t \leq \text{gcd}(L/a, b) - 1} 1 = \sum_{a, b | L} \text{gcd}(L/a, b) = \sum_{a, b | L} \text{gcd}(a, b). \quad (2.16)$$

The second statement is proved similarly by reordering the sum

$$\sum_{\substack{a,b|L \\ L^2/(ab)}} \gcd(L/a, b) = \sum_{\substack{a,b|L \\ La/b}} \gcd(a, b). \quad (2.17)$$

□

2.2 Metaplectic operators

Metaplectic operators allow to generalize results that were obtained for separable lattices to a more general class of lattices, i.e. the *symplectic lattices*. We will not give a full introduction to the subject here, as it relies on representation theory, which is beyond the scope of this text. The background can be found in [54] for the continuous case and in [38, 71] for the discrete case. A metaplectic operator, loosely speaking, is the signal domain counterpart to a symplectic transform of the lattice on phase space (i.e. the TF plane). The treatment of the finite discrete case is different from its continuous counterpart, a fact that we have already observed when describing lattices in phase space in the previous section.

2.2.1 Continuous

We first define the standard symplectic form, which arises from the representation theoretical background that we will not discuss, for proofs and a more thorough treatment we refer to [45, 54]. In the following we make use of the block matrix

$$\mathcal{J} = \begin{pmatrix} 0 & -I_d \\ I_d & 0 \end{pmatrix}, \quad (2.18)$$

where I_d denotes the $d \times d$ identity matrix.

Definition 2.10. The *symplectic group* is the group of all invertible matrices $A \in \mathbb{R}^{2d \times 2d}$, that satisfy

$$z^T A^T \mathcal{J} A z' = z^T \mathcal{J} z'. \quad (2.19)$$

The reason for this definition is the fact that for any symplectic deformation A of the TF-plane we can find an operator acting on signal space that is associated to it.

Proposition 2.11. Let $f, g \in L^2(\mathbb{R}^d)$ and A a symplectic matrix. Then there exists a unitary operator U_A and a phase factor ϕ_A , such that for all $\lambda \in \mathbb{R}^{2d}$

$$U_A \pi(\lambda) = \phi_A(\lambda) \pi(A\lambda) U_A. \quad (2.20)$$

Every symplectic matrix can be decomposed into a finite product of elements of 3 subgroups of the symplectic group [45].

Theorem 2.12. The symplectic group is generated by the following subgroups

1. $\{\mathcal{J}^k : k = 0, \dots, 3\},$
2. $\left\{ \begin{pmatrix} B & 0 \\ 0 & B_*^{-1} \end{pmatrix} : B \in GL(d) \right\},$

$$3. \left\{ \begin{pmatrix} I_d & 0 \\ C & I_d \end{pmatrix} : C = C^* \right\}.$$

For these subgroups it is possible to explicitly write down the formulas for the corresponding metaplectic operators.

Corollary 2.13. *The elementary matrices admit the following metaplectic operators*

$$1. U_{\mathcal{F}} = \mathcal{F}.$$

$$2. B = \begin{pmatrix} B^{*-1} & 0 \\ 0 & B \end{pmatrix} \text{ for some } B \in GL(d), \text{ then } U_B f(t) = (\det B)^{-1/2} f(B^{*-1}t).$$

$$3. C = \begin{pmatrix} I & 0 \\ C & I \end{pmatrix} \text{ for some } C = C^*, \text{ then } U_C f(t) = e^{-i\pi t \cdot C t} f(t).$$

Remark 2.14. Rotations matrices are not part of the generating subgroups in Theorem 2.12, however as they have determinant 1 they are symplectic in the one dimensional setting. One elegant way of seeing this is to write a rotation as a product of shear matrices [79]. This technique comes from computer graphics, as shears are computationally cheaper to realize than rotations. A general rotation matrix can be decomposed as

$$R(\alpha) = \begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix} = \begin{pmatrix} 1 & -\tan \alpha/2 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & 0 \\ \sin \alpha & 1 \end{pmatrix} \begin{pmatrix} 1 & -\tan \alpha/2 \\ 0 & 1 \end{pmatrix}, \quad (2.21)$$

for some rotation angle α . The corresponding metaplectic operator is the fractional Fourier Transform

$$U_{R(\alpha)} = \mathcal{F}^\alpha, \quad (2.22)$$

which we discussed at the end of Section 1.2.2 on Hermite functions.

With these results at hand we can turn our attention to non-separable lattices.

Definition 2.15. Let $A \in \mathbb{R}^{2d \times 2d}$. Then the generated lattice $\Lambda = A\mathbb{Z}^{2d}$ is called *symplectic* if there exists a diagonal matrix D and a symplectic matrix P , such that

$$A = PD. \quad (2.23)$$

A symplectic lattice is the image of a separable lattice under a symplectic transformation. For these lattices the Gabor transform and the Gabor frame operator can be reduced to the separable case.

Corollary 2.16. *Let $\Lambda = P\tilde{\Lambda}$ with a symplectic matrix P and U_P the corresponding metaplectic operator. By setting $\tilde{g} = U_P^{-1}g$*

$$S_{g,\Lambda} = U_P S_{\tilde{g},\tilde{\Lambda}} U_P^{-1}. \quad (2.24)$$

Furthermore, the Gabor coefficients can be computed using the identity

$$\langle f, \pi(\lambda)g \rangle = \phi_P(\lambda) \langle U_P^{-1}f, \pi(P^{-1}\lambda)\tilde{g} \rangle, \quad (2.25)$$

for all $\lambda \in \Lambda$.

Remark 2.17. 1. In one dimension every matrix with determinant 1 is symplectic

2. In one dimension every lattice is symplectic as any matrix A can be decomposed as $A = (\det(A)^{-1}A) \det(A)I$.

The above corollary allows to generalize the structural results from Section 1.6.2 to symplectic lattices. In one dimension this even means to all possible lattices. For simplicity we focus on the one dimensional setting from now on, it is a non-trivial problem to characterize the symplectic lattices in higher dimensions exactly, as the symplectic group has a more complicated structure.

Definition 2.18. Given a lattice Λ in the one-dimensional phase space. Then we denote the *adjoint lattice* as

$$\Lambda^0 = \{(x, \xi) : \pi(x, \xi)\pi(\lambda) = \pi(\lambda)\pi(x, \xi) \text{ for all } \lambda \in \Lambda\}. \quad (2.26)$$

Proposition 2.19. Let the lattice Λ be generated by the matrix $A \in \mathbb{R}^{2 \times 2}$. Then the adjoint lattice is generated by the matrix $\det(A)^{-1}A$.

Proof. With the notation $\lambda = (x, \xi)^T$ and $\lambda^0 = (x^0, \xi^0)^T$ we find that

$$\pi(\lambda)\pi(\lambda^0) = e^{2\pi i(x\xi^0 - x^0\xi)}\pi(\lambda^0)\pi(\lambda). \quad (2.27)$$

Therefore the two time-frequency shifts commute if and only if $x\xi^0 - x^0\xi \in \mathbb{Z}$ for all $(x, \xi) \in \Lambda$. If the lattice is separable, i.e. $\Lambda = D\mathbb{Z}^2$ for some diagonal matrix D it is easy to see that this condition is satisfied if and only if $(x^0, \xi^0) \in \Lambda^0 = \det(D)^{-1}D\mathbb{Z}^2$.

If the lattice is non-separable we can write it as the image of a symplectic matrix of a separable lattice, i.e. $\Lambda = P\tilde{\Lambda}$, where $\tilde{\Lambda}$ is separable (c.f. Remark 2.17). Then we use Proposition 2.11 to compute for all $\lambda \in \tilde{\Lambda}$

$$\begin{aligned} \pi(P\lambda)\pi(\lambda^0) &= U_P\pi(\lambda)\pi(P^{-1}\lambda^0)U_P \\ &= U_P\pi(\lambda)\pi(P^{-1}\lambda^0)U_P^{-1} \cdot \phi_P^{-1}(\lambda)\phi_P^{-1}(P^{-1}\lambda^0). \end{aligned} \quad (2.28)$$

The two time frequency shifts $\pi(\lambda)$ and $\pi(P^{-1}\lambda^0)$ in the above equation commute if and only if $P^{-1}\lambda^0 \in \tilde{\Lambda}^0$. One can revert the computations in (2.28) to conclude that $\Lambda^0 = P\tilde{\Lambda}^0$. Since $\tilde{\Lambda}^0 = \det(D)^{-1}D\mathbb{Z}^2$ and $\det(A)^{-1}A = \det(D)^{-1}PD$ the proof is finished. $\square$

Adjoint lattices have already played an important role when formulating the Janssen representation and the Wexler-Raz relations, as the adjoint lattice of the separable lattice $a\mathbb{Z} \times b\mathbb{Z}$ is given by $\frac{1}{b}\mathbb{Z} \times \frac{1}{a}\mathbb{Z}$. Now we can generalize the Janssen representation and the Wexler-Raz relations to arbitrary lattices in $\mathbb{R}^2$.

Corollary 2.20. Let Λ be a lattice generated by $A \in \mathbb{R}^{2 \times 2}$ and $g \in L^2(\mathbb{R})$ be a window function that satisfies

$$\sum_{\lambda^0 \in \Lambda^0} |V_g g(\lambda^0)| < \infty. \quad (2.29)$$

Then the frame operator assumes the form

$$S_{g, \Lambda} = \det(A)^{-1} \sum_{\lambda^0 \in \Lambda^0} V_g g(\lambda^0) \pi(\lambda^0). \quad (2.30)$$

Corollary 2.21. *Let Λ be a lattice generated by $A \in \mathbb{R}^{2 \times 2}$ and assume that $D_{g,\Lambda}$ and $D_{\gamma,\Lambda}$ are both bounded operators. Then the following are equivalent*

1. $\sum_{\lambda \in \Lambda} \langle f, \pi(\lambda)g \rangle \pi(\lambda)\gamma = f$ for all $f \in L^2(\mathbb{R}^d)$.
2. $\det(A)^{-1} \langle \gamma, \pi(\lambda^0)g \rangle = \delta(\lambda^0)$ for $\lambda^0 \in \Lambda^0$.

Using this theory of metaplectic operators we can also generalize Theorem 1.52.

Corollary 2.22. *The Gabor system $\mathcal{G}(g_0, \Lambda)$ is a frame for $L^2(\mathbb{R})$ if and only if $\text{vol}(\Lambda) < 1$.*

2.2.2 Discrete

The discrete counterpart of the theory of metaplectic operators is fundamentally different from the continuous case. The discrete phase space has a different structure and not all symplectic operations from the continuous theory can be carried over to the finite discrete case. For example, rotations can not be defined in general on a purely discrete phase space. Also dilations, which are part of the continuous theory can only be found in the discrete case for very special choices of dilation factors. We will follow the summary of [108] for an introduction into the theory of finite discrete metaplectic operators. Since we will need the metaplectic and symplectic operators in more explicit form when we discuss algorithms for the finite discrete Gabor transform later in Section 2.3 we start by defining some elementary symplectic transformations

Definition 2.23. For $c \in \mathbb{Z}_L$ and $a \in \mathbb{Z}_L$ invertible (i.e. $\gcd(a, L) = 1$) define the following matrices

$$\begin{aligned} F &= \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}, \quad S_c = \begin{pmatrix} 1 & 0 \\ c & 1 \end{pmatrix}, \\ D_a &= \begin{pmatrix} a & 0 \\ 0 & a^{-1} \end{pmatrix}, \end{aligned} \tag{2.31}$$

As in the continuous setting, it is possible to decompose all matrices with determinant 1 into these basic building blocks.

Proposition 2.24 (Feichtinger et al. (2008) [38]). *Let $A = \begin{pmatrix} a & b \\ c & d \end{pmatrix} \in \mathbb{Z}_L^2$ with $\det(A) = 1$, then there exists $m \in \mathbb{Z}$ such that $a_0 = a + mb$ is invertible in $\mathbb{Z}_L$. Let $c_0 = c + md$, then*

$$A = S_{c_0 a_0^{-1}} D_{a_0} F^{-1} S_{-a_0^{-1} b} F S_{-m}. \tag{2.32}$$

The proof is based on Weil's decomposition of arbitrary symplectic matrices.

Lemma 2.25. *For the elementary matrices of Definition 2.23 we define the corresponding metaplectic operators as follows*

$$\begin{aligned} F &\mapsto U_F = \mathcal{F} \\ S_c &\mapsto U_{S_c} = (f(\cdot) \mapsto f(\cdot) \exp(\pi i c \cdot^2 (L+1)/L)) \\ D_a &\mapsto U_{D_a} = (f(\cdot) \mapsto f(a^{-1} \cdot)). \end{aligned}$$

With these transformations the following hold for all $\lambda \in \mathbb{Z}_L^2$

$$\begin{aligned} U_F \pi(\lambda) &= \phi_F(\lambda) \pi(F\lambda) U_F \\ U_{S_c} \pi(\lambda) &= \phi_{S_c}(\lambda) \pi(S_c \lambda) U_{S_c} \\ U_{D_a} \pi(\lambda) &= \phi_{D_a}(\lambda) \pi(D_a \lambda) U_{D_a}, \end{aligned}$$

where ϕ_F , ϕ_{S_c} and ϕ_{D_a} are phase factors.

Proof. Some simple calculations are sufficient to establish the result:

$$\begin{aligned} U_F \pi(\lambda) f &= \mathcal{F} M_\xi T_x f = T_\xi M_{-x} \hat{f} \\ &= e^{-2\pi i x \xi / L} M_{-x} T_\xi \hat{f} = e^{-2\pi i x \xi / L} \pi(F\lambda) U_F f, \end{aligned}$$

$$\begin{aligned} U_{S_c} \pi(\lambda) f &= M_\xi T_x e^{\pi i c(\cdot+x)^2(L+1)/L} f \\ &= e^{\pi i c x^2(L+1)/L} M_{\xi+cx} T_x U_{S_c} f \\ &= e^{\pi i c x^2(L+1)/L} \pi(S_c \lambda) U_{S_c} f \end{aligned}$$

and

$$\begin{aligned} U_{D_a} \pi(\lambda) f &= M_{a^{-1}\xi} f(a \cdot -x) = M_{a^{-1}\xi} f(a^{-1}(\cdot - ax)) \\ &= M_{a^{-1}\xi} T_{ax} U_{D_a} f = \pi(D_a \lambda) U_{D_a} f. \end{aligned}$$

□

Remark 2.26. The proper definition of discrete finite chirps, which is the operator associated to the operation of shearing on the lattice has been the subject of some discussion, see e.g. [14]. While the naive linear chirp $\exp(2\pi i s t^2 / L)$ is still used by Bastiaans and van Leest [103, 104], a more appropriate definition has been proposed by Kaiblinger et al. [38, 69], constituting a second degree character. While these second degree characters are the correct ones to establish the commutation properties with time-frequency shifts as in Lemma 2.25, they can be the source of large numerical errors if not implemented carefully. Computing the numerator in the exponent $\pi i c l^2(L+1)$ leads to a very large number, while only the part $cl^2(L+1) \bmod 2L$ matters in the final computation of the exponential. This modulo operation should be present in any implementation, as otherwise the results will contain numerical errors in an unacceptable range.

The combination of the two results above immediately yields the following theorem.

Theorem 2.27. *For any matrix $A \in \mathbb{Z}_L^2$ with $\det(A) = 1$, there exists a metaplectic operator U_A , such that for all $\lambda \in \mathbb{Z}_L^2$*

$$U_A \pi(\lambda) = \phi_A(\lambda) \pi(A\lambda) U_A. \quad (2.33)$$

The definition of a discrete symplectic lattice has to be done slightly different from the continuous case.

Definition 2.28. Let $A \in \mathbb{Z}_L^{2 \times 2}$. Then the generated lattice $\Lambda = A\mathbb{Z}_L^2$ is called *symplectic* if there exist matrices V, P with $\det V = \det P = 1$, a diagonal matrix D and a symplectic matrix P , such that

$$A = PDV. \quad (2.34)$$

Remark 2.29. The difference between the symplecticity in the discrete and continuous case stems from the different structure of the phase space. While the $V\mathbb{Z}_L^2 = \mathbb{Z}_L^2$ for any matrix V with determinant 1, the equivalent statement is not true in the continuous case, i.e. $\mathbb{Z}^2 \neq V\mathbb{Z}^2$ for $\det V = 1$ in general. The multiplication of $\mathbb{Z}_L^2$ with a determinant 1 matrix is just a reshuffling of the lattice elements. In the discrete case one applies the matrix to the whole phase space, while in the continuous setting the lattice matrix is applied to some standard lattice given by $\mathbb{Z}^2$ and that is already a subset of the TF-plane.

Analogous to the continuous case we find that every lattice in $\mathbb{Z}_L^2$ is symplectic via the *Smith normal form* for matrices with integer entries. This normal form has first been used in this context in [38].

Theorem 2.30 (Smith's normal form). *Let $A \in \mathbb{Z}^{m \times m}$. Then there exist matrices $P \in \mathbb{Z}^{m \times m}$ and $V \in \mathbb{Z}^{m \times m}$ with determinant 1 and a unique diagonal matrix D , such that*

- *the diagonal entries of D satisfy $d_i \geq 0$ for $i = 1, \dots, m$ and $d_i | d_{i+1}$ for $i = 1, \dots, m-1$,*
- *and $A = PDV$.*

The matrix D is called Smith normal form of A .

Proposition 2.31. *Let $\Lambda = A\mathbb{Z}_L^2$ be a lattice and $A = \tilde{P}\tilde{D}\tilde{V}$ the Smith decomposition of A . Then*

$$\Lambda = P\tilde{\Lambda}, \quad (2.35)$$

where $P = (\tilde{P} \bmod L)$, $D = (\tilde{D} \bmod L)$ and $\tilde{\Lambda} = D\mathbb{Z}_L^2$.

The above proposition uses the Smith normal form to show that any given lattice can be viewed as symplectic image of a rectangular grid. Before we are able to generalize the Janssen representation and the Wexler-Raz relations we need to explain adjoint lattices in the discrete setting. The definition is completely analogous to the continuous case.

Definition 2.32. Given a lattice $\Lambda \leq \mathbb{Z}_L^2$. Then we denote the *adjoint lattice* as

$$\Lambda^0 = \{(x, \xi) : \pi(x, \xi)\pi(\lambda) = \pi(\lambda)\pi(x, \xi) \text{ for all } \lambda \in \Lambda\}. \quad (2.36)$$

Using the lattice normal form (2.3) we can explicitly compute the adjoint lattice.

Proposition 2.33. *Given a lattice in normal form $\Lambda = \begin{pmatrix} a & 0 \\ s & b \end{pmatrix} \cdot \mathbb{Z}_L^2$. Then the adjoint lattice is given by*

$$\Lambda^0 = \begin{pmatrix} \frac{L}{b} & 0 \\ \frac{Ls}{ab} & \frac{L}{a} \end{pmatrix} \cdot \mathbb{Z}_L^2 \quad (2.37)$$

Proof. With the notation $\lambda = (x, \xi)^T$ and $\lambda^0 = (x^0, \xi^0)^T$ we find that

$$\pi(\lambda)\pi(\lambda^0) = e^{2\pi i(x\xi^0 - x^0\xi)/L} \pi(\lambda^0)\pi(\lambda). \quad (2.38)$$

Therefore, we need to determine all λ^0 , such that

$$x\xi^0 - x^0\xi = 0 \pmod{L} \quad \text{for all } \lambda \in \Lambda. \quad (2.39)$$

Due to the lattice normal form it is sufficient to find all the points in $\mathbb{Z}_L^2$ that satisfy the above equation for the generators $(a, s)^T$ and $(0, b)^T$. Plugging the second generator into (2.39) we immediately see that $x^0 = k \frac{L}{b}$ for some $k \in \mathbb{Z}$. Using this and plugging the second generator into (2.39) we find that for some $\tilde{k} \in \mathbb{Z}$

$$\xi^0 = k \frac{Ls}{ab} + \tilde{k} \frac{L}{a}, \quad (2.40)$$

since $(ab)|(Ls)$. The assertion follows because we just proved that

$$(x^0, \xi^0)^T = k \cdot \left(\frac{L}{b}, \frac{Ls}{ab} \right)^T + \tilde{k} \cdot \left(0, \frac{L}{a} \right)^T. \quad (2.41)$$

□

This allows to generalize the discrete Janssen representation as well as the discrete Wexler-Raz relations to non-separable lattices.

Corollary 2.34. *Given a lattice in normal form $\Lambda = \begin{pmatrix} a & 0 \\ s & b \end{pmatrix} \cdot \mathbb{Z}_L^2$ and a window function $g \in \mathbb{C}^L$. Then the frame operator assumes the form*

$$S_{g,\Lambda} = \frac{L}{ab} \sum_{\lambda^0 \in \Lambda^0} \langle g, \pi(\lambda^0) g \rangle \pi(\lambda^0). \quad (2.42)$$

Corollary 2.35. *Given a lattice in normal form $\Lambda = \begin{pmatrix} a & 0 \\ s & b \end{pmatrix} \cdot \mathbb{Z}_L^2$ and window functions $g, \gamma \in \mathbb{C}^L$. Then the following are equivalent*

1. $\sum_{\lambda \in \Lambda} \langle f, \pi(\lambda) g \rangle \pi(\lambda) \gamma = f$ for all $f \in \mathbb{C}^L$.
2. $\frac{L}{ab} \langle \gamma, \pi(\lambda^0) g \rangle = \delta(\lambda^0)$ for $\lambda^0 \in \Lambda^0$.

The proofs for Corollaries 2.34 and 2.35 follow immediately from their separable counterparts (Theorem 1.57, Theorem 1.58) using the theory of metaplectic operators.

2.3 Computation on non-separable lattices

This section is fully devoted to algorithmic considerations, therefore we restrict ourselves to the finite discrete setting. The material follows largely the recently submitted paper [108], where we present new and fast algorithms for the finite discrete Gabor transform. Nonseparable lattices in $\mathbb{Z}_L^2$ can be interpreted in a variety of ways. Several different approaches relate Gabor expansions on general lattices to one or several equivalent expansions on separable (or rectangular) lattices. From an algorithmic viewpoint, these are of particular interest, since a wealth of research [2, 5, 9, 82, 94, 110] has investigated efficient algorithms for analysis and synthesis using Gabor dictionaries on separable lattices. Each of the three approaches described in this section yield a simple relation between arbitrary given Gabor systems and Gabor systems on separable sampling sets that can be harnessed for efficient analysis and synthesis.

In some cases we switch to a description of a subgroup of $\mathbb{Z}_L^2$ that is equivalent to the normal form introduced in ((2.3)), as it comes up more natural in some

settings. Instead of the shear parameter s , one can also use the shear relative to b , given by

$$\lambda = \frac{s}{b} = \frac{\lambda_1}{\lambda_2}, \text{ with } \lambda_1 = \frac{s}{\gcd(b, s)}, \lambda_2 = \frac{b}{\gcd(b, s)}, \quad (2.43)$$

This easily explains how to convert s into λ_1 and λ_2 and vice versa. Unlike in the case of separable lattices, there is no immediate natural way of indexing the Gabor coefficients. However, it seems sensible to index by the position in time and counting the sampling points in frequency from the lowest nonnegative frequency upwards. Therefore we fix

$$c(m, n) = \sum_{l=0}^{L-1} f(l) \overline{g(l - an + 1)} e^{-2\pi i l b(m + w(n))/M}, \quad (2.44)$$

for the rest of this section, where the additional offset w is given by $w(n) = \text{mod}(n\lambda_1, \lambda_2)/\lambda_2$. This format is also implemented in the open source MATLAB/OCTAVE Toolbox *LTFAT* [1], used for the experiments in Section 2.4.

2.3.1 Correspondence via multiwindow Gabor

We decompose a given lattice into a union of λ_2 cosets of a sparser separable lattice, which allows us to use multiwindow methods for the computation. Using multiwindow methods for computation of Gabor transforms on nonseparable lattices has been proposed in [40, 112] and implementation has been discussed in [103, 104]. However, the latter only briefly mention the computation of dual Gabor windows, not discussing efficient implementations in detail.

Proposition 2.36. *Given the lattice Λ in normal form specified by the parameters a, b and s , then*

$$\Lambda = \cup_{r=0}^{\lambda_2-1} \left((ar, sr \bmod b)^T + \tilde{\Lambda} \right), \quad (2.45)$$

where $\lambda_2 = b/\gcd(b, s)$ and $\tilde{\Lambda}$ is the separable lattice generated by $(\lambda_2 a, 0)^T$ and $(0, b)$.

Proof. Let the matrix generating Λ be denoted by A and let us define $M_x = \{sx + b\xi : \xi \in \mathbb{Z}_L\}$, for $0 \leq x < L/a$. We note here, that M_x is the second coordinate of the set $A \cdot (x, \mathbb{Z}_L)^T$. Furthermore, $0 \in M_x$ if and only if x is a multiple of λ_2 . To see that, we first note that λ_1 and λ_2 are relatively prime. Then the following equation has a solution if and only if x is a multiple of λ_2

$$sx + b\xi = b \left(\frac{\lambda_1}{\lambda_2} x + \xi \right) = 0. \quad (2.46)$$

This yields

$$\begin{aligned} M_x &= M_{x+\lambda_2}, \quad \text{for } x \in \mathbb{Z}_{L/a} \\ M_x &= sx \bmod b + M_0 \end{aligned}$$

This observation yields the following decomposition of the original lattice

$$\begin{aligned}\Lambda &= \bigcup_{x=0}^{L/a-1} \{ax\} \times M_x \\ &= \bigcup_{r=0}^{\lambda_2-1} \left((ar, sr \bmod b)^T + \bigcup_{j=0}^{L/(a\lambda_2)-1} \{aj\lambda_2\} \times M_0 \right),\end{aligned}\tag{2.47}$$

which finishes the proof by observing

$$\tilde{\Lambda} = \bigcup_{j=0}^{L/(a\lambda_2)-1} \{aj\lambda_2\} \times M_0.\tag{2.48}$$

□

We can now describe a Gabor system $\mathcal{G}(g, \Lambda)$, with Λ in the normal form, and the related operators completely in terms of a multiwindow Gabor system $\mathcal{G}(\mathbf{g}, \tilde{\Lambda})$, where $\mathbf{g} = \{g_r\}_{r=0, \dots, \lambda_2-1}$ on the separable lattice $\tilde{\Lambda}$.

Proposition 2.37. *Let $\mathcal{G}(g, \Lambda)$, $\mathcal{G}(g_r, \tilde{\Lambda})$, with $\Lambda, \tilde{\Lambda}$ as in Proposition 2.36 and $g \in \mathbb{C}^L$, $g_r = M_{rs \bmod b} T_{ra} g$, for $0 \leq r < \lambda_2$, be Gabor systems, then*

$$S_{g, \Lambda} f = \sum_{r=0}^{\lambda_2-1} S_{g_r, \tilde{\Lambda}} f.\tag{2.49}$$

Moreover, the Gabor transform can be computed using the identity

$$\begin{aligned}c(m, n) &= \langle f, M_{mb+(ns \bmod b)} T_{na} g \rangle \\ &= e^{-2\pi i \tilde{n} a \lambda_2 (rs \bmod b)/L} \langle f, M_{mb} T_{\tilde{n} a \lambda_2} g_m \rangle,\end{aligned}\tag{2.50}$$

where $\tilde{n} = \lfloor n/\lambda_2 \rfloor$ and $m = n - \tilde{n} \lambda_2$.

Proof. Using the fact that $ns \bmod b = rs \bmod b$ we find that

$$\begin{aligned}M_{mb+(ns \bmod b)} T_{na} g &= M_{mb} M_{rs \bmod b} T_{\tilde{n} a \lambda_2} T_{ra} g \\ &= e^{2\pi i \tilde{n} a \lambda_2 (rs \bmod b)/L} M_{mb} T_{\tilde{n} a \lambda_2} M_{rs \bmod b} T_{ra} g \\ &= e^{2\pi i \tilde{n} a \lambda_2 (rs \bmod b)/L} M_{mb} T_{\tilde{n} a \lambda_2} g_r,\end{aligned}\tag{2.51}$$

yielding (2.50). Using this we find

$$\begin{aligned}S_{g, \Lambda} &= \sum_{n=0}^{L/a-1} \sum_{m=0}^{L/b-1} \langle f, M_{mb+ns \bmod b} T_{na} g \rangle M_{mb+ns \bmod b} T_{na} g \\ &= \sum_{r=0}^{\lambda_2-1} \sum_{\tilde{n}=0}^{L/\tilde{a}-1} \sum_{m=0}^{L/b-1} \langle f, M_{mb} T_{\tilde{n} \tilde{a}} g_r \rangle M_{mb} T_{\tilde{n} \tilde{a}} g_r,\end{aligned}\tag{2.52}$$

where $\tilde{a} = \lambda_2 a$.

□

2.3.2 Correspondence via Smith normal form

In this and the following section we will use the theory of metaplectic operators developed in Section 2.2.2, i.e. we will determine a determinant 1 matrix P , such that $\Lambda = P\tilde{\Lambda}$ for a general lattice Λ and a separable lattice $\tilde{\Lambda}$. This is precisely how we proved the symplecticity of any lattice in the TF plane in Proposition 2.31. Using this correspondence and Proposition 2.24 and Lemma 2.25 we can find an operator associated to P , which leads to the final computational procedure described in the following Corollary.

Corollary 2.38. *Let Λ be given. Furthermore determine some P using Smith's decomposition to find a separable lattice $\tilde{\Lambda}$ satisfying $\Lambda = P\tilde{\Lambda}$ and denote the metaplectic operator associated to P by U_P . Setting $\tilde{g} = U_P^{-1}g$ the frame operator satisfies*

$$S_{g,\Lambda} = U_P S_{g,\tilde{\Lambda}} U_P^{-1}. \quad (2.53)$$

Moreover, the Gabor coefficients can be computed using the identity

$$\langle f, \pi(z)g \rangle = \phi_P(z) \langle U_P^{-1}f, \pi(P^{-1}z)\tilde{g} \rangle, \quad (2.54)$$

for all $z = (x, \xi)^T \in \Lambda$.

The proof of the corollary above is straightforward using the tools developed in Section 2.2.2 on discrete metaplectic operators.

2.3.3 Correspondence via shearing

The method detailed in the previous section uses in general 6 elementary symplectic matrices to deform a general lattice into a separable lattice. In the following we will show that this can be reduced to 4 or less. Reducing computations on nonseparable lattices to well known computation on product lattices using shearing has been proposed earlier [103, 104], however the authors were able to describe only a subset of all lattices over $\mathbb{Z}_L^2$ as shears of rectangular lattices. In [103] the author speculates that it might be possible to describe every lattice as the image of a product lattice under a horizontal (frequency side) and a vertical (time side) shear. The following theorem shows that this is indeed true.

Theorem 2.39. *let $A \in \mathbb{Z}_L^{2 \times 2}$. There exist $s_0, s_1 \in \mathbb{Z}_L$ and $V \in \mathbb{Z}_L^{2 \times 2}$ with $|\det(V)| = 1$, such that*

$$A = P_{s_0, s_1} D V, \quad (2.55)$$

where $D \in \mathbb{Z}_L^{2 \times 2}$ is diagonal and

$$P_{s_0, s_1} = S_{-s_1} F^{-1} S_{s_0} F. \quad (2.56)$$

For better readability of the material we will proof the above result after stating the computational procedure connected to this particular reduction of nonseparable lattices. We can now rewrite Gabor transforms on nonseparable lattices in the vein of Proposition 2.16 using the metaplectic operator associated to P_{s_0, s_1} . Subsequently, the metaplectic operator associated with P_{s_0, s_1} will be denoted by

$$U_{s_0, s_1} = U_{S_{-s_1}} \mathcal{F}^{-1} U_{S_{s_0}} \mathcal{F}. \quad (2.57)$$

Proposition 2.40. *Let $\Lambda = A\mathbb{Z}_L^2$ be a lattice, D, P_{s_0, s_1} as in the previous theorem and $\tilde{\Lambda} = D\mathbb{Z}_L^2$. Furthermore let $g \in \mathbb{C}^L$ and $\tilde{g} = U_{s_0, s_1}^{-1}g$. Then*

$$S_{g, \Lambda} f = U_{s_0, s_1} S_{\tilde{g}, \tilde{\Lambda}} U_{s_0, s_1}^{-1} f \quad (2.58)$$

and

$$\langle f, M_\xi T_x g \rangle = \phi_{U_{s_0, s_1}}(z) \langle U_{s_0, s_1}^{-1} f, M_{\xi - s_1(x - s_0 \xi)} T_{x - s_0 \xi} \tilde{g} \rangle, \quad (2.59)$$

for all $z = (x, \xi)^T$. Moreover,

$$\phi_{U_{s_0, s_1}}(z) = e^{\pi i (s_0 \xi^2 - s_1 (x - s_0 \xi)^2) (L+1)/L}. \quad (2.60)$$

Proof. Everything but the explicit form of the phase factor $\phi_{U_{s_0, s_1}}$ is a direct consequence of Lemma 2.25 and Theorem 2.39, note

$$P_{s_0, s_1} = S_{-s_1} F^{-1} S_{s_0} F = \begin{pmatrix} 1 & -s_0 \\ -s_1 & s_0 s_1 + 1 \end{pmatrix}. \quad (2.61)$$

To complete the proof, set $y = (x - s_0 \xi)$ and determine the phase factor explicitly

$$\begin{aligned} & U_{S_{-s_1}} \mathcal{F}^{-1} U_{S_{s_0}} \mathcal{F} M_\xi T_x f \\ &= e^{\pi i s_0 \xi^2 (L+1)/L} U_{S_{-s_1}} \mathcal{F}^{-1} T_\xi M_{s_0 \xi - x} U_{S_{s_0}} \mathcal{F} f \\ &= e^{\pi i s_0 \xi^2 (L+1)/L} U_{S_{-s_1}} M_\xi T_y \mathcal{F}^{-1} U_{S_{s_0}} \mathcal{F} f \\ &= e^{\pi i (s_0 \xi^2 - s_1 y^2) (L+1)/L} M_{\xi - s_1 y} T_y U_{S_{-s_1}} \mathcal{F}^{-1} U_{S_{s_0}} \mathcal{F} f, \end{aligned} \quad (2.62)$$

where we used Lemma 2.25 and $\exp(2\pi i m(L+1)/L) = \exp(2\pi i m/L)$ for all $m \in \mathbb{Z}$. $\square$

For Proposition 2.40 to be valid, it remains to prove Theorem 2.39, establishing the representation of A through U_{s_0, s_1} .

Proof of Theorem 2.39. We can assume without loss of generality that A is in lattice normal form, i.e.

$$A = \begin{pmatrix} a & 0 \\ s & b \end{pmatrix}. \quad (2.63)$$

To prove equation (2.55), we rewrite $U_{s_0, s_1}^{-1} A = DV$ with a diagonal matrix D and a unitary matrix V . It can be seen that

$$U_{s_0, s_1}^{-1} = \begin{pmatrix} s_0 s_1 + 1 & s_0 \\ s_1 & 1 \end{pmatrix} = \begin{pmatrix} 1 & s_0 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & 0 \\ s_1 & 1 \end{pmatrix}. \quad (2.64)$$

Now, using Proposition 2.6 in the step from (2.65) to (2.66) below, we can write

$$\begin{aligned} U_{s_0, s_1}^{-1} A &= \begin{pmatrix} 1 & s_0 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & 0 \\ s_1 & 1 \end{pmatrix} \begin{pmatrix} a & 0 \\ s & b \end{pmatrix} \\ &= \begin{pmatrix} 1 & s_0 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} a & 0 \\ s_1 a + s & b \end{pmatrix} \end{aligned} \quad (2.65)$$

$$= \begin{pmatrix} 1 & s_0 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} \frac{ab}{X} & ak_1 \\ 0 & X \end{pmatrix} \begin{pmatrix} k_2 & -k_1 \\ Y & b/X \end{pmatrix} \quad (2.66)$$

$$= \begin{pmatrix} \frac{ab}{X} & s_0 X + ak_1 \\ 0 & X \end{pmatrix} \begin{pmatrix} k_2 & -k_1 \\ Y & b/X \end{pmatrix}. \quad (2.67)$$

Here $X = \gcd(s_1a + s, b)$, $Y = X^{-1}(s_1a + s)$ and k_1, k_2 stem from Bézout's identity when representing $\gcd(s_1a + s, b) = k_1(s_1a + s) + k_2b$. It is important to note that the second matrix in the last line has determinant one. This shows that the lattice $U_{s_0, s_1}A$ is separable if and only if $\tilde{D} = \begin{pmatrix} ab/X & s_0X + ak_1 \\ 0 & X \end{pmatrix}$ is equivalent to a diagonal matrix, i.e.

$$\text{mod } (s_0X + ak_1, ab/X) = 0. \quad (2.68)$$

We will now deduce numbers s_0 and s_1 satisfying our needs from the prime factor decomposition of the involved quantities. Therefore we represent $L = \prod_{j=1}^J p_j^{n_j}$ for a fixed set of prime numbers. Since a and b are divisors of L we find their prime factor decompositions to have exponents $\{\alpha_j\}_{j=1}^J$ and $\{\beta_j\}_{j=1}^J$, where $\alpha_j, \beta_j \leq n_j$. The shearing parameter has the decomposition $s = l \prod_{j=1}^J p_j^{\sigma_j}$, where $\gcd(l, L) = 1$.

We choose

$$s_1 = \prod_{j=1}^J p_j^{\mu_j}, \text{ where } \mu_j = \begin{cases} 1 & \text{for } \alpha_j = \sigma_j \\ 0 & \text{else.} \end{cases} \quad (2.69)$$

With this choice of s_1 we investigate

$$X = \gcd(s_1a + s, b) = \prod_{j=1}^J \gcd(s_1a + s, p_j^{\beta_j}). \quad (2.70)$$

To do so, we have to individually treat three cases:

1. $\alpha_j < \sigma_j$: Since s_1 and p_j are coprime we find $\gcd(s_1a + s, p_j^{\beta_j}) = p_j^{\min(\alpha_j, \beta_j)}$.
2. $\alpha_j > \sigma_j$: $\gcd(s_1a + s, p_j^{\beta_j}) = p_j^{\min(\sigma_j, \beta_j)}$, and $\min(\sigma_j, \beta_j) < \alpha_j$
3. $\alpha_j = \sigma_j$: Use Eq. (2.69) to determine that $\gcd(s_1a + s, p_j^{\beta_j}) = p_j^{\min(\alpha_j, \beta_j)}$

The above arguments show that with the choice of s_1 , we find that $X = \prod_{j=1}^J p_j^{\gamma_j}$, where $\gamma_j \leq \alpha_j$.

Now we turn to the choice of s_0 . To do so we first decompose $k_1 = l \prod_{j=1}^J p_j^{\kappa_j}$, where l and L are coprime. Let us explain how to choose the shear via the positive part of a vector

$$s_0 = \left(\prod_{j=1}^J p_j^{(\beta_j - \gamma_j - \kappa_j)_+} - l \right) \prod_{j=1}^J p_j^{\alpha_j + \kappa_j - \gamma_j}, \quad (2.71)$$

where $(x_+)_j = \max(x_j, 0)$. A straightforward calculation shows then that

$$s_0X + ak_1 = \prod_{j=1}^J p_j^{(\beta_j - \gamma_j - \kappa_j)_+ + \alpha_j + \kappa_j}, \quad (2.72)$$

and we furthermore see that

$$(\beta_j - \gamma_j - \kappa_j)_+ + \alpha_j + \kappa_j \geq \beta_j + \alpha_j - \gamma_j. \quad (2.73)$$

This proves that (2.68) is satisfied, completing the proof. $\square$

Remark 2.41. It follows directly from the proof above that $X|a$ and $X|b$, therefore ab/X is a multiple of X . As a consequence, the normal form described in Theorem 2.39 is indeed Smith's normal form.

2.3.4 Further optimization

In this section we will first determine which signal lengths are feasible for some given sampling strategy. The treatment here is somewhat different from the question of which subgroups of the TF plane we can find for a fixed signal length L , which we addressed in depth in Section 2.1. This point of view is essential for computations, as in practice padding the signal with zeros to a feasible signal length is a standard procedure.

Particularly when using the shear method described in Subsection 2.3.3 it is interesting to know for which signal lengths one of the two shears s_0 and s_1 , preferably the frequency side shear s_0 , can be chosen to be zero. This saves additional computation time.

We will introduce here some important constants related to the *time shift* a , the *frequency shift* b , the *number of channels* $M = L/b$ and the *number of time shifts* $N = L/a$. The number of channels is denoted by M , as well as the modulation operator. This is a slight abuse of notation but the operator is always used with a subscript. We define $c, d, p, q \in \mathbb{N}$ by

$$c = \gcd(a, M) \quad , \quad d = \gcd(b, N) \quad , \quad (2.74)$$

$$p = \frac{a}{c} = \frac{b}{d} \quad , \quad q = \frac{M}{c} = \frac{N}{d} \quad . \quad (2.75)$$

With these numbers, the redundancy of a Gabor system can be written as $L/(ab) = q/p$, where q/p is an irreducible fraction. It holds that $L = cdpq$. Some of the introduced notation will be important in the next section.

Proposition 2.42. *Given the parameters $\lambda = \lambda_1/\lambda_2$, a and M . Then the minimal signal length, for which these parameters are feasible is given by $L_{\min} = \lambda_2 \operatorname{lcm}(a, M)$. All the feasible signal lengths are multiples of $L_{\min}$.*

Proof. For the parameters in combination with a given signal length L to form a lattice we require the following conditions

$$\begin{aligned} a|L, \quad M|L, \\ \frac{L}{a}\lambda \in \mathbb{Z}, \\ \frac{L}{M}\lambda \in \mathbb{Z}, \end{aligned} \quad (2.76)$$

where the first conditions immediately yield $\operatorname{lcm}(a, M)|L$. From the other two conditions we can derive

$$a\lambda_2/\gcd(a, \lambda_1)|L \quad \text{and} \quad M\lambda_2/\gcd(M, \lambda_1)|L. \quad (2.77)$$

Therefore, the signal length has to be a multiple of

$$L_{\min} = \operatorname{lcm}\left(\frac{a\lambda_2}{\gcd(a, \lambda_1)}, \frac{M\lambda_2}{\gcd(M, \lambda_1)}, a, M\right). \quad (2.78)$$

We proceed to show that

$$\operatorname{lcm}\left(\frac{a\lambda_2}{\gcd(a, \lambda_1)}, a\right) = \lambda_2 a. \quad (2.79)$$

For this purpose we look at the prime factor decomposition of the involved quantities, where we denote by $\alpha_j, \gamma_j, \delta_j$ the exponents of the prime number p_j of a, λ_1 and λ_2 respectively. Then we find, since λ_1 and λ_2 are co-prime that the exponent of p_j of $\text{lcm}(a\lambda_2/\text{gcd}(a, \lambda_1), a)$ is given by

$$\max(\alpha_j - \min(\alpha_j, \gamma_j) + \delta_j, \alpha_j) = \alpha_j + \delta_j, \quad (2.80)$$

proving (2.79). The proof that $\text{lcm}(M\lambda_2/\text{gcd}(M, \lambda_1), M) = \lambda_2 M$ is completely analogous. Combine these to find

$$\begin{aligned} & \text{lcm}\left(\frac{a\lambda_2}{\text{gcd}(a, \lambda_1)}, \frac{M\lambda_2}{\text{gcd}(M, \lambda_1)}, a, M\right) \\ &= \text{lcm}(\lambda_2 a, \lambda_2 M) = \lambda_2 \text{lcm}(a, M). \end{aligned} \quad (2.81)$$

□

Now we shall investigate, which multiples of the just derived minimal signal length allow for computation without the frequency shear. To do so, it is instructive to compute the set of factors l , for which $L = lL_{\min}$ needs only the time shear.

Proposition 2.43. *Given λ, a and M . Let the prime factor decomposition of $c = \text{gcd}(a, M)$ be given by*

$$c = \prod_{j=1}^J p_j^{\gamma_j} \quad (2.82)$$

for some set of prime factors and corresponding exponents. Let

$$c_1 = \prod_{j=1}^J p_j^{\sigma_j}, \quad \sigma_j = \begin{cases} \gamma_j & \text{if } \text{gcd}(\lambda_2, p_j) = 1 \\ 0 & \text{else,} \end{cases} \quad (2.83)$$

then the frequency shear can be chosen to be 0 if the signal length satisfies

$$L = nL_{\min} \frac{c}{c_1}, \quad (2.84)$$

for some $n \in \mathbb{N}$. In words, c_1 are factors of c that are relatively prime to λ_2

Proof. With the standard notation we easily see that the time shear is sufficient if and only if $(s + kb)/a \in \mathbb{Z}$ for some $k \in \{0, \dots, M-1\}$. Rewriting this leads to

$$L = \tilde{l} \frac{Ma\lambda_2}{\lambda_1 + k\lambda_2} = l \frac{Ma\lambda_2}{\text{gcd}(\lambda_1 + k\lambda_2, Ma\lambda_2)}, \quad (2.85)$$

for some $l \in \mathbb{Z}$. However, these signal lengths might not be compatible with the feasibility condition from Proposition 2.42. Therefore we compute the ratio

$$\frac{L}{L_{\min}} = l \frac{\text{gcd}(M, a)}{\text{gcd}(\lambda_1 + k\lambda_2, Ma\lambda_2)}. \quad (2.86)$$

Since this fraction should be an integer number, we have to choose

$$l = n \frac{\text{gcd}(\lambda_1 + k\lambda_2, Ma\lambda_2)}{\text{gcd}(M, a, \lambda_1 + k\lambda_2, Ma\lambda_2)}, \quad (2.87)$$

for some $n \in \mathbb{N}$. Therefore, we can compute

$$L = nL_{\min} \frac{\gcd(M, a)}{\gcd(M, a, \lambda_1 + k\lambda_2)}. \quad (2.88)$$

With the notation introduced above we are now interested in computing

$$\max_{k \in \mathbb{N}} (\gcd(c, \lambda_1 + k\lambda_2)). \quad (2.89)$$

Firstly, we rewrite

$$\gcd(c, \lambda_1 + k\lambda_2) = \prod_{j=1}^J \gcd(p_j^{\gamma_j}, \lambda_1 + k\lambda_2). \quad (2.90)$$

Now we will individually investigate the factors in the product above.

Case 1. $\gcd(p_j, \lambda_2) = 1$, in which case we can find numbers $k_{1,j}, k_{2,j}$, such that

$$\lambda_1 + k_{1,j}\lambda_2 = k_{2,j}p_j^{\gamma_j}. \quad (2.91)$$

Furthermore, the full set of coefficients of λ_2 , for which the above equation can be satisfied is given by $K_j = \{k_{1,j} + mp_j^{\gamma_j} : m \in \mathbb{Z}\}$. Therefore, for any $k \in K_j$ we find

$$\gcd(p_j^{\gamma_j}, \lambda_1 + k\lambda_2) = p_j^{\gamma_j}. \quad (2.92)$$

Case 2. $\gcd(p_j, \lambda_2) \neq 1$, which implies directly that λ_2 is a multiple of p_j . In this case we have to argue that $\lambda_1 + k\lambda_2$ can never be a multiple of p_j . Indeed, any linear combination $k_{1,j}\lambda_2 + k_{2,j}p_j$ is a multiple of p_j and therefore not equal to λ_1 , which is assumed to be relatively prime to λ_2 . Consequently, for any choice of $k \in \mathbb{Z}$

$$\gcd(p_j^{\gamma_j}, \lambda_1 + k\lambda_2) = 1. \quad (2.93)$$

For all the indices j in case 1, it is easy to see that the intersection of the corresponding sets K_j is not empty. This is an immediate consequence from the fact that powers of two different prime numbers have no common divisors. Using the notation introduced above, we can conclude that there exists some $k \in \mathbb{Z}$, such that

$$\gcd(c, \lambda_1 + \lambda_2) = c_1. \quad (2.94)$$

The last argument needed is to show that $k \in \{0, \dots, M-1\}$. By construction $s + kb = \tilde{k}a$, for some $\tilde{k} \in \mathbb{Z}$. Therefore, for any $m \in \mathbb{Z}$

$$s + \left(k + m\frac{L}{b}\right)b = \left(\tilde{k} + m\frac{L}{a}\right)a, \quad (2.95)$$

and for an appropriate choice of m , the expression in brackets on the left hand side will equal some number in the desired range. $\square$

Remark 2.44. There are possibly other feasible signal lengths than those determined by (2.84). The full set of feasible lengths is determined by

$$\left\{ nL_{\min} \frac{\gcd(M, a)}{\gcd(M, a, \lambda_1 + k\lambda_2)} : n \in \mathbb{N}, k \in \{0, \dots, M-1\} \right\}. \quad (2.96)$$

This can be easily seen from (2.88) in the proof above. For simplicity we only construct the minimal factor, that $L_{\min}$ has to be multiplied with, as stated in (2.89).

Remark 2.45. Looking at (2.84) we see that if we are given a certain redundancy (2.75) q/p and a lattice type λ_1/λ_2 and want to get a low value of $L_{\min}c/c_1$ we must choose c such that it is relatively prime to λ_2 . As an example, consider a common choice of $a = 32$, $M = 64$ and $\lambda_1/\lambda_2 = 1/2$ (the quincunx lattice). In this case $L_{\min} = 128$ and $c = \gcd(a, M) = 32$ which is the worst possible case, as c is a power of $\lambda_2 = 2$ giving a value of $L_{\min}c/c_1 = 128 \cdot 32 = 4096$. If we instead choose $a = 27$, $M = 54$ (which is the same redundancy) we get $L_{\min}c/c_1 = 108 \cdot 1 = 108$. This illustrates that it is possible to work efficiently with the quincunx lattice by not choosing the rectangular lattice parameters to be powers of 2.

2.3.5 Extension to higher dimensions

It is well known [16, 39, 80] that multidimensional Gabor transforms and dual windows can be computed using algorithms designed for the 1D case, if both the Gabor window and the lattice used can be written as a tensor product. That is, we assume that with $l = (l_1, \dots, l_n)^T \in \mathbb{C}^{L_1} \times \dots \times \mathbb{C}^{L_n}$,

$$g(l) = g_1(l_1) \otimes \dots \otimes g_n(l_n)$$

and

$$\Lambda = \Lambda_1 \times \dots \times \Lambda_n = A_1 \mathbb{Z}_{L_1}^2 \times \dots \times A_n \mathbb{Z}_{L_n}^2$$

for some $A_j \in \mathbb{Z}_{L_j}^2 \times \mathbb{Z}_{L_j}^2$ for $j = 1, \dots, n$.

Equivalently, we can say that Λ can be described by a block matrix

$$A = \begin{pmatrix} D & E \\ F & G \end{pmatrix}. \quad (2.97)$$

with diagonal blocks $D, E, F, G \in \mathbb{Z}^{n \times n}$. In this case, the multidimensional transform and dual window can be computed by subsequently applying the algorithms presented in the previous sections in every dimension. A matrix describing the lower dimensional lattice corresponding to dimension j is simply given by

$$A_j = \begin{pmatrix} D_{j,j} & E_{j,j} \\ F_{j,j} & G_{j,j} \end{pmatrix}$$

and can be transformed into lattice normal form, allowing straightforward application of the presented algorithms.

However, we are not aware of a constructive method to determine whether a lattice, given by an arbitrary matrix, can be described by a banded matrix of the form (2.97).

In contrast to the methods based on metaplectic operators, it is easier to extend the multiwindow approach from 2.3.1 to higher dimensions. For reasons of readability we will not give details here.

2.4 Implementation and timing

In this section we discuss the implementation and speed of the proposed algorithms. It is important for this section to recall the definition of the constants c, d in (2.74) and p, q in (2.75).

Methodology for computing the computational complexity. To compute the Discrete Fourier transform, the familiar FFT algorithm is used. When computing the flop (floating point operations) count of the algorithm, we will assume that a complex FFT of length M can be computed using $4M \log_2 M$ flops. A review of flop counts for FFT algorithms is presented in [68]. When computing the flop count, we assume that both the window and signal are complex valued.

The cost of performing the computation of a DGT with a full length window on a rectangular lattice using the algorithm first reported in [91] is given by

$$8Lq + 4L \log_2 d + 4MN \log_2 d + 4MN \log_2 (M) \quad (2.98)$$

$$= L(8q + 4 \log_2 d) + 4MN (\log_2 L/p) \quad (2.99)$$

where the first terms in (2.98) come from the multiplication of the matrices in the factorization, the two middle terms come from creating the factorization of the signal and inverting the factorization of the coefficients, and the last term comes from the final application of FFTs. The terms can be collected as in (2.99), where the first term grows as the length of the signal L , and the second term grows as the total number of coefficients MN . In the following, we refer to this as the *full window* algorithm.

If the window is an SIR window supported on an index set with width L_g which is much smaller than the length of the signal L , the *weighted-overlap-add algorithm*, first reported in [82], can be used instead. It has a computational complexity of

$$8L \frac{L_g}{a} + 4NM \log_2 M. \quad (2.100)$$

In the following, we refer to this as the *SIR window* algorithm.

A third approach to computing a DGT is a hybrid approach, where a DGT using an SIR window can be computing using a full window algorithm on blocks of the input signal. The blocks are then combined using the classical *overlap-add* (OLA) algorithm, cf. [60, 93].

The OLA algorithm works by partitioning a system of length L into blocks of length L_b such that $L = L_b N_b$, where N_b is the number of blocks. The block length must be longer than the support of the window, $L_b > L_g$. To perform the computation we take a block of the input signal of length L_b and zero-extend it to length $L_x = L_b + L_g$, and compute the convolution with the extended signal using the similarly extended window. Because of the zero-extension of the window and signal, the computed coefficients will not be affected by the periodic boundary conditions, and it is therefore possible to overlay and add the computed convolutions of length L_x together to form the complete convolution of length L .

The equations (2.99), (2.100) are used to express the efficiency of the algorithms for the DGT on nonseparable lattices.

2.4.1 Implementation of the shear algorithm

The shear algorithm proposed in Proposition 2.40 computes the DGT on a nonseparable lattice using a DGT on a separable lattice with some suitable pre- and postprocessing steps. The computational complexity of the pre- and postprocessing steps is significant compared to the separable DGT, so we wish

to minimize the cost of these steps. An implementation of the shear algorithm is presented as Algorithm 1. Note that we assume the existence of several underlying routines: An implementation DGT of the separable Gabor transform, the periodic chirp $\text{PCHIRP}(L, s) = \exp(\pi i s \cdot^2 (L+1)/L)$ and SHEARFIND, a program that determines the shear parameters s_0, s_1 and the correct separable lattice to do the DGT on, following the constructive proof of Theorem 2.39.

Algorithm 1 The shear algorithm: $c = \text{DGTNS}(f, g, a, M, \lambda)$

```

1:  $[s_0, s_1, b_r] = \text{SHEARFIND}(L, a, M, \lambda)$ 
2: if  $s_1 \neq 0$  then
3:    $p \leftarrow \text{PCHIRP}(L, s_1)$ 
4:    $g(\cdot) \leftarrow p(\cdot)g(\cdot)$ 
5:    $f(\cdot) \leftarrow p(\cdot)f(\cdot)$ 
6: end if
7: if  $s_0 = 0$  then
8:    $c_r \leftarrow \text{DGT}(f, g, a, M)$ 
9:    $C_1 \leftarrow s_1 a (L+1) \pmod{2N}$ 
10:  for  $k = 0 \rightarrow N-1$  do
11:     $E \leftarrow e^{\pi i (C_1 k^2 \pmod{2N})/N}$ 
12:    for  $m = 0 \rightarrow M-1$  do
13:       $c(\lfloor \frac{-s_1 k a + m b \pmod{L}}{b} \rfloor, k) \leftarrow E c_r(m, k)$ 
14:    end for
15:  end for
16: else
17:    $a_r \leftarrow \frac{ab}{b_r}, \quad M_r \leftarrow \frac{L}{b_r}, \quad N_r \leftarrow \frac{L}{a_r}$ 
18:    $C_1 \leftarrow \frac{a_r}{a}, \quad C_2 \leftarrow -s_0 b_r / a$ 
19:    $C_3 \leftarrow a s_1 (L+1), \quad C_4 \leftarrow C_2 b_r (L+1)$ 
20:    $C_5 \leftarrow 2C_1 b_r, \quad C_6 \leftarrow (s_0 s_1 + 1) b_r$ 
21:    $p \leftarrow \text{PCHIRP}(L, -s_0)$ 
22:    $g(\cdot) \leftarrow p(\cdot) \text{FFT}(g(\cdot))/L$ 
23:    $f(\cdot) \leftarrow p(\cdot) \text{FFT}(f(\cdot))$ 
24:    $c_r \leftarrow \text{DGT}(f, g, b_r, N_r)$ 
25:   for  $k = 0 \rightarrow N_r - 1$  do
26:     for  $m = 0 \rightarrow M_r - 1$  do
27:        $s_{q1} \leftarrow C_1 k + C_2 m \pmod{2N}$ 
28:        $E \leftarrow e^{\pi i (C_3 s_{q1}^2 - m(C_4 m + C_5 k) \pmod{2N})/N}$ 
29:        $\tilde{m} \leftarrow C_1 k + C_2 m \pmod{N}$ 
30:        $\tilde{k} \leftarrow \lfloor \frac{-s_1 a_r k + C_6 m \pmod{L}}{b} \rfloor$ 
31:        $c(\tilde{k}, \tilde{m}) \leftarrow E c_r(-k \pmod{N_r}, m)$ 
32:     end for
33:   end for
34: end if

```

A simple trick is to notice that when a frequency-side shear is needed, the DFT of the signal f and the window g are multiplied by a chirp on the frequency side, $\tilde{f} = U_{s_0, s_1}^{-1} f$ and $\tilde{g} = U_{s_0, s_1}^{-1} g$. The total cost of this is 4 FFTs and two pointwise multiplications. However, instead of transforming the chirped signal and window back to the time domain, we can compute the nonseparable DGT

directly in the frequency domain using the well-known commutation relation of the DFT and the translation and modulation operators:

$$\langle f, M_m T_n g \rangle = e^{-\pi i m n / L} \langle \mathcal{F} f, M_{-n} T_m \mathcal{F} g \rangle \quad (2.101)$$

This trick saves the two inverse FFTs at the expense of the multiplication of the coefficients by a complex exponential and reshuffling. As we already need these operations to realize (2.59) and (2.60), they can be combined with no additional computational complexity.

The overlap-add algorithm can be used in conjunction with the shear algorithm in the following case: we wish to compute the DGT with an SIR window for a nonseparable lattice using the shear algorithm. Because of the frequency-side shearing, the window is converted from an SIR window into a full length window, making it impossible to perform real-time or block-wise processing. However, if the shear algorithm is used inside an OLA algorithm, this is no longer a concern, as the shearing will only convert the window into a window of length $L_g + L_b$, restoring the ability to perform block-wise processing.

In total, the shear-OLA algorithm for the DGT is calculated in three steps using the three algorithms:

1. Split the input signal into blocks using the overlap-add algorithm
2. Apply the shears to the blocks of the input signal as in the shear algorithm
3. Use the full-window rectangular lattice DGT on the sheared signal blocks.

The downside of the shear-OLA algorithm is that the total length of the DGTs is longer than the original DGT by

$$\rho = \frac{L_g + L_b}{L_b}, \quad (2.102)$$

where L_b is the block length. Therefore, a trade-off between the block length and the window length must be found, so that the block length is long enough for (2.102) to be close to one, but at the same time small enough to not impose a too long processing delay.

2.4.2 Dual and tight windows

The shear method in Proposition 2.40 can also be used to compute the canonical dual and canonical tight windows, using the factorization of the frame operator given in (2.58). The complete algorithm for the canonical dual window is shown in Algorithm 2, and uses the same trick as the Gabor transform algorithm to compute the canonical dual when a frequency side shear is needed: do it in the Fourier domain without transforming back. Again, we assume the existence of an implementation GABDUAL for the computation of Gabor dual windows on separable lattices.

To compute the canonical dual and tight windows on a separable lattice, the matrices are first factorized as in [91, 94] and then the factorized matrices are transformed as in [67].

Algorithm 2 Dual window via shearing: $\tilde{g} = \text{GABDUALNS}(g, a, M, \lambda)$

```

1:  $[s_0, s_1, b_r] = \text{SHEARFIND}(L, a, M, \lambda)$ 
2: if  $s_1 \neq 0$  then
3:    $p \leftarrow \text{PCHIRP}(L, s_1)$ 
4:    $g(\cdot) \leftarrow p(\cdot)g(\cdot)$ 
5: end if
6:  $b \leftarrow \frac{L}{M}, \quad M_r \leftarrow \frac{L}{b_r}, \quad a_r \leftarrow \frac{ab}{b_r}$ 
7: if  $s_0 = 0$  then
8:    $g_d \leftarrow \text{GABDUAL}(g, a_r, M_r)$ 
9: else
10:   $p_0 \leftarrow \text{PCHIRP}(L, -s_0)$ 
11:   $g(\cdot) \leftarrow p_0(\cdot)\text{FFT}(g)(\cdot)$ 
12:   $g_d \leftarrow L \cdot \text{GABDUAL}(g, L/M_r, L/a_r)$ 
13:   $g_d \leftarrow \text{IFFT}(\overline{p_0}(\cdot)g_d(\cdot))$ 
14: end if

```

Alg.:	Flop count
Multi-window	
SIR.	$8L \frac{L_g}{a} + 4NM \log_2 M$
Full.	$L\lambda_2 (8q_{mw} + 4 \log_2 d_{mw})$ $+ MN (4 \log_2 L/p_{mw} + 6)$
SNF	$L (8q + 4 \log_2 d_{sm} + 8 \log_2 L + 18)$ $+ MN (4 \log_2 L/p + 6)$
Shear alg.	
No freq. shear	$L (8q + 4 \log_2 d + 6k_{time})$ $+ MN (4 \log_2 L/p + 6k_{time})$
Freq. shear	$L (8q + 4 \log_2 Lc_{sh} + 6 + 6k_{time})$ $+ MN (4 \log_2 L/p + 6)$
Shear OLA	
No freq. shear	$\rho L (8q + 4 \log_2 \rho d_{shola} + 6k_{time})$ $+ \rho MN (4 \log_2 \rho L_b/p + 6k_{time})$
Freq. shear	$\rho L (8q + 4 \log_2 \rho Lc_{shola} + 6k_{time} + 6)$ $+ \rho MN (4 \log_2 \rho L_b/p + 6)$

Table 2.1: Flop counts for different ways of computing the DGT on a nonseparable lattice. First column list the algorithm, second column the flop count for the particular algorithm. Listed from the top, the algorithms are: The multiwindow algorithm using the full window rectangular lattice algorithm, the multiwindow algorithm using the SIR window rectangular lattice algorithm, the Smith normal form algorithm using the full window rectangular lattice algorithm, the shear algorithm when no frequency shear is needed, the shear algorithm including the frequency shear and finally the overlap-add versions of the shear algorithms. The term L_g denotes the length of the window used so L_g/a is the overlapping factor of the window.

2.4.3 Analysis of the computational complexity

The flop counts of the various algorithms used for computing the DGT on a nonseparable lattice is listed in Table 2.1. The additional parameters used are

defined as follows. For the multiwindow algorithms we define

$$\begin{aligned} c_{mw} &= \gcd(a_{mw}, M), & d_{mw} &= \gcd(b, N_{mw}), \\ p_{mw} &= N_{mw}/d_{mw}, & q_{mw} &= M/c_{mw}, \end{aligned} \quad (2.103)$$

where $a_{mw} = a\lambda_2$ and $N_{mw} = N/\lambda_2$. For the shear and shear OLA algorithms we define

$$\begin{aligned} c_{sh} &= \frac{ca_{sh}}{a}, & d_{sh} &= \frac{dM}{M_{sh}}, \\ c_{shola} &= \frac{ca_{shola}}{a}, & d_{shola} &= \frac{dM}{M_{shola}}, \end{aligned} \quad (2.104)$$

$$k_{time} = \begin{cases} 1 & \text{if a time shear is used} \\ 0 & \text{otherwise,} \end{cases} \quad (2.105)$$

where a_{sh}, M_{sh} and a_{shola}, M_{shola} are the parameters of the rectangular lattice the problem is reduced to in the respective algorithm.

Based on the computational complexity presented in the table, any of the algorithms may for some specific problem setup be the fastest, except for the Smith-normal form algorithm which is always slower than the shear algorithm:

- The multiwindow algorithm for SIR windows is the fastest for very short windows.
- The multiwindow-OLA algorithm is the fastest for simple lattices (λ_2 small) and medium length windows.
- The shear-OLA algorithm is the fastest for more complex lattices (λ_2 large) and medium length windows.
- The multiwindow algorithm is the fastest for simple lattices (λ_2 small) and very long windows.
- The shear algorithm is the fastest for more complex lattices (λ_2 large) and very long windows.

2.4.4 Numerical experiments

Implementations of the algorithms described in this paper can be found in the Large Time Frequency Analysis Toolbox (LTFAT), cf. [92], [1]. An appropriate algorithm will be automatically invoked when calling the DGT or DGTREAL functions. The implementations are done in both the MATLAB / OCTAVE scripting language and in C. All tests were performed on an Intel i7 CPU operating at 3.6 GHz.

As the speed of the algorithms depends on a large number of parameters $a, M, L, L_g, c, d, s_0, s_1$ and similar parameters relating to the multiwindow and shear transforms, we cannot provide an exhaustive illustration of the running times. Instead we will present some figures that illustrates the crossover point of when the the shear algorithm becomes faster than the multiwindow algorithm as the lattice complexity λ_2 increases. The behavior of the algorithms as the window length L_g increases is completely determined by the algorithms for the rectangular lattice, so we refer to [91] for illustrations.

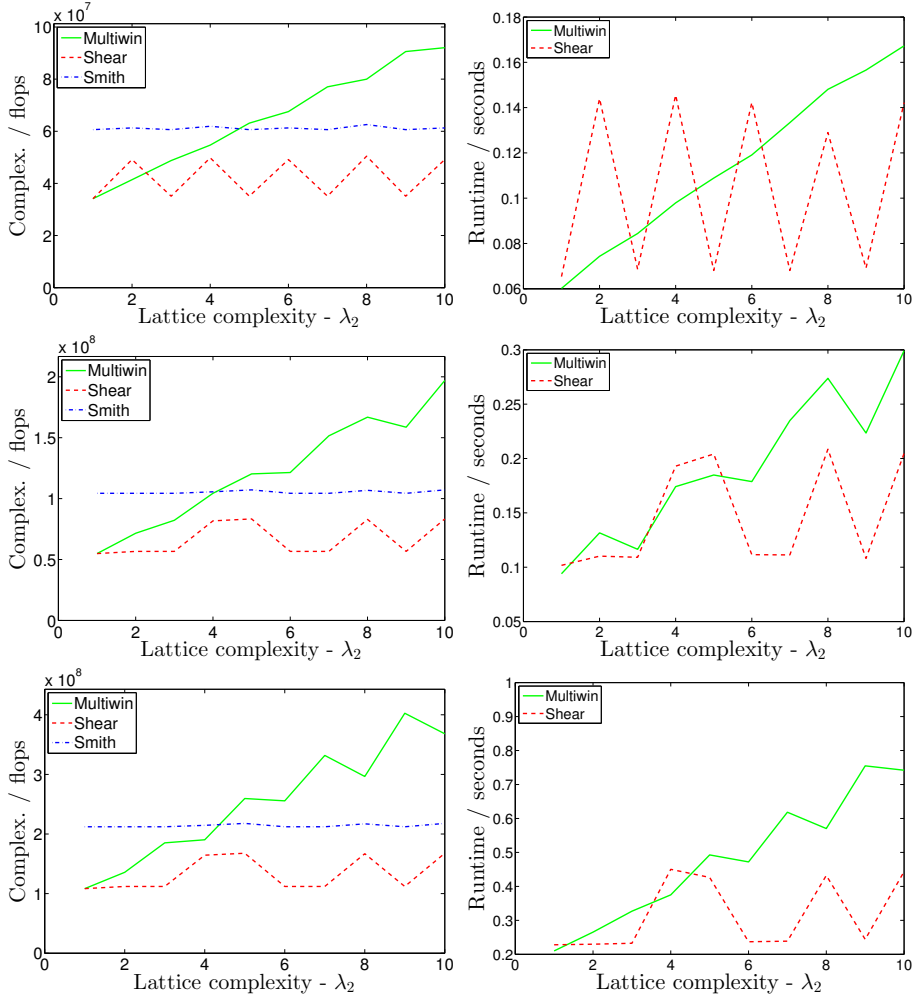


Figure 2.2: Computation of the DGT for nonseparable lattices with increasing lattice complexities, λ_2 . The length is kept fixed at $L = \text{lcm}(a, M) \cdot 2520$ which is the minimal legal transform length for all the tested lattices. (left) Accurate flop counts and (right) the actual running time. The Gabor system parameters are $a = 32$, $M = 64$ ($p/q = 1/2$) (1st row), $a = 40$, $M = 60$ ($p/q = 2/3$) (2nd row) and $a = 60$, $M = 80$ ($p/q = 3/4$) (3rd row).

The experiments shown in Figure 2.2 illustrate how the computational complexity of the running time of the algorithm depends on the lattice complexity λ_2 : The complexity of the shear algorithm is independent of λ_2 , while the complexity of the multiwindow algorithm grows linearly.

The bumps in the curves for the multiwindow algorithm are due to variations in q_{mw} : The multiwindow algorithm transforms the problem into λ_2 different DGTs that should be computed on a lattice with redundancy $q/(p\lambda_2)$. The number q_{mw} is the nominator of this written as an irreducible fraction, and depending on $p\lambda_2$ it may be smaller than q .

The bumps in the curves for the shear algorithm are caused by whether or not a frequency side shear is required for that particular lattice configuration, and to a lesser extent whether a time-side shear is needed. As the multiwindow algorithm is faster for simple lattices, there is a cross-over point where the shear algorithm becomes faster, but the cross-over point depends strongly on the exact lattice configuration. Just considering the flop counts would predict that the cross-over happens for a smaller value of λ_2 than what is really the case. This is due to the fact there are more complicated indexing operations and memory reshuffling for the shear algorithm than for the multiwindow algorithm, and this is not properly reflected in the flop count.

The cross-over point where one algorithm is faster than the other is highly dependent on the interplay between the algorithm and the computer architecture. Experience from the ATLAS [107], FFTW [49] and SPIRAL [29] projects shows that in order to have the highest performance, it is necessary to select the algorithm for a given problem size based on previous tests done on the very same machine. Performing such an optimization is beyond the scope of this paper, and we therefore cannot make statements about how to choose the most efficient cross-over points.

Chapter 3

Optimal Gabor frames

In the context of time-frequency analysis, it is commonly agreed that best results are obtained when one uses window functions that are well localized in the time-frequency domain. In Section 1.2.3 we have already seen that a window cannot be concentrated in time and frequency arbitrarily well at the same time. There we found the Gaussian to optimize the standard Heisenberg uncertainty principle. Loosely speaking, a waveform that is well concentrated will allow to extract time-frequency information from a signal at a fairly precise point. When talking about Gabor transforms it is intuitive to try to work with a window, that is concentrated on the fundamental domain of the lattice since this will allow to extract the desired information optimally. We will measure the concentration of a window in terms of its ambiguity function. The ambiguity function plays an important role in many different areas, e.g. for operator approximation [34] or for RADAR and coding applications [3, 21, 61]. Moreover, we have seen in Section 1.5, that this function also serves as the reproducing kernel for the STFT.

In the following we will give some arguments how to adapt the window and lattice to each other. Since in Chapter 2 we developed methods to compute the Gabor transform on general lattices we will consider non-separable lattices in our optimizations also. We will judge the quality of a Gabor system by the ratio of its frame bounds as they give an estimate how close the frame operator is to the identity, i.e. a tight system. Systems which are close to tight are desirable because the canonical dual system will consist of similar windows, which ensures that they do not lose the concentration properties. It is generally hard to make statements about the exact condition of Gabor systems, we will provide some results that give an intuition how to design systems by estimating the frame bounds from above and below by easily computable numbers.

3.1 The Janssen test and tilings of the TF-plane

Most of this section will be devoted to a detailed study of the role of the ambiguity function $V_g g$ for some given window function g . We will show that the ambiguity function plays an important role and that we can employ some geometrical arguments to obtain good sampling patterns.

The first criterion we describe is the so-called Janssen test, first described

in [101] and also used in [28]. It is an easily computable sufficient condition for a Gabor family to establish a frame, it is an almost immediate consequence from the Janssen representation.

Corollary 3.1 (Janssen's test). *Let $\mathcal{G}(g, \Lambda)$ be a Gabor family with $g \in \mathcal{S}_0(\mathbb{R}^d)(\in \mathbb{C}^L)$ and $\|g\|_2 = 1$. Furthermore, assume that*

$$\sum_{\lambda^0 \in \Lambda^0} |V_g g(\lambda^0)| \leq 1 + D < 2. \quad (3.1)$$

Then $\mathcal{G}(g, \Lambda)$ is a frame, where the frame bounds A, B satisfy

$$\text{vol}(\Lambda)^{-1}(1 - D) \leq A \leq B \leq \text{vol}(\Lambda)^{-1}(1 + D). \quad (3.2)$$

Proof. The conditions ensure that the frame operator $S_{g, \Lambda}$ has a Janssen representation. With (3.1) this implies

$$\|\text{vol}(\Lambda)S - I\| = \left\| \sum_{\lambda^0 \in \Lambda^0 \setminus \{0\}} V_g g(\lambda^0) \pi(\lambda^0) \right\| \leq \sum_{\lambda^0 \in \Lambda^0 \setminus \{0\}} |V_g g(\lambda^0)| \leq D. \quad (3.3)$$

Therefore we can use the Neumann series to calculate

$$\|(\text{vol}(\Lambda)S)^{-1}\| \leq \sum_{k=0}^{\infty} D^k = \frac{1}{1 - D}, \quad (3.4)$$

which shows that $\|S^{-1}\| \leq \text{vol}(\Lambda)/(1 - D)$. Using the Janssen representation directly we find that

$$\|\text{vol}(\Lambda)S\| \leq 1 + D. \quad (3.5)$$

Theorem 1.33 relates the norm of the inverse frame operator with the lower frame bound and yields the assertion. $\square$

It is reasonable to ask the question which sampling strategy yields the best estimates in the Janssen test instead of computing the analytic frame bounds, i.e. for some given window function we are interested in solving the problem

$$\Lambda_{\text{best}} = \arg \min_{\Lambda \text{ lattice}} \sum_{\lambda^0 \in \Lambda^0} |V_g g(\lambda^0)|. \quad (3.6)$$

While the Janssen test gives an easy sufficient condition for frames, it does not immediately show a connection to the quality of the tiling with the TF plane. The next results will use the periodization of the ambiguity function on the chosen lattice given by

$$F_{g, \Lambda}(x) = \sum_{\lambda \in \Lambda} |V_g g(x - \lambda)|^2. \quad (3.7)$$

If the ambiguity functions pack the TF-plane well, then this function will be almost constant, more precisely, its maximum and minimum lie close to each other. It is easy to see that the function $F_{g, \Lambda}(x) = F_{g, \Lambda}(x - \lambda)$ for any $\lambda \in \Lambda$, i.e. the function F is periodic on the lattice.

Proposition 3.2. *Let $g \in L^2(\mathbb{R})(\in \mathbb{C}^L)$ and a lattice in phase space be given. Then the following are equivalent*

1. The function $F_{g,\Lambda}(x) = c$ for some constant $c > 0$.
2. The Gabor family $\mathcal{G}(g, \Lambda)$ is a tight frame.

Furthermore, the frame is tight with $A = B = 1$ if $c = \text{vol}(\Lambda)^{-1}$.

Proof. We start by using Janssen's representation to compute

$$\begin{aligned}
 \sum_{\lambda \in \Lambda} |V_g g(x - \lambda)|^2 &= \sum_{\lambda \in \Lambda} \langle g, \pi(x - \lambda)g \rangle \langle \pi(x - \lambda)g, g \rangle \\
 &= \langle S_{g,\Lambda} \pi(-x)g, \pi(-x)g \rangle \\
 &= \sum_{\lambda^0 \in \Lambda^0} \langle g, \pi(\lambda^0)g \rangle \langle \pi(\lambda^0)\pi(-x)g, \pi(-x)g \rangle.
 \end{aligned} \tag{3.8}$$

If $\mathcal{G}(g, \Lambda^0)$ is an orthogonal system (OS), so is $\mathcal{G}(\pi(-x)g, \Lambda^0)$. Therefore we see from the above equation that $\langle \pi(\lambda^0)\pi(-x)g, \pi(-x)g \rangle$ is independent of x if $\mathcal{G}(g, \Lambda^0)$ is an OS. The converse is also true, therefore $\mathcal{G}(g, \Lambda^0)$ is an OS is equivalent to $F_{g,\Lambda}$ being constant. The family $\mathcal{G}(g, \Lambda^0)$ is an orthogonal family if and only if $\mathcal{G}(g, \Lambda)$ is a tight frame by the Wexler-Raz relations.

From (3.8) it is easy to see that $F_{g,\Lambda}(x) = \text{vol}(\Lambda)^{-1}$ is equivalent to $\mathcal{G}(g, \Lambda^0)$ being an orthonormal family, which is in turn equivalent to $\mathcal{G}(g, \Lambda)$ being a tight frame with frame constants equal to 1. $\square$

The above result allows to exactly characterize tight frames by looking at the periodization of the ambiguity function. Now we derive a necessary condition on $F_{g,\Lambda}$ for the corresponding Gabor family to form a frame.

Proposition 3.3. *Let $\mathcal{G}(g, \Lambda)$ be a frame for $L^2(\mathbb{R})(\mathbb{C}^L)$ with frame bounds A, B and $\|g\| = 1$. Then*

$$A \leq \min_x F_{g,\Lambda}(x) \leq \max_x F_{g,\Lambda}(x) \leq B. \tag{3.9}$$

Proof. We use the frame inequality and the unitarity of the operator $\pi(-x)$ for some x in phase space to see

$$A \leq \sum_{\lambda \in \Lambda} |\langle \pi(-x)g, \pi(\lambda)g \rangle|^2 \leq B. \tag{3.10}$$

Moving $\pi(-x)$ to the other side of the inner product we find that

$$A \leq \sum_{\lambda \in \Lambda} |V_g g(x - \lambda)|^2 \leq B. \tag{3.11}$$

$\square$

In order to fulfill Janssen's criterion and also get a good estimate from Proposition 3.3 we want the ambiguity function to be adapted to the fundamental domain, i.e. the fundamental tiles of the lattice.

Remark 3.4. In (3.11) of Proposition 3.3 we find a necessary condition for the family $\mathcal{G}(g, \Lambda)$ to be a frame. Proposition 3.2 on the other hand contains the converse statement, i.e. a sufficient condition. However, this sufficient condition holds true only in the case where the periodized ambiguity functions sum up to a constant. Nevertheless, this shows that trying to sample the TF plane in a way, such that $F_{g,\Lambda}$ is bounded from above and below is necessary when constructing frames. In Figure 3.1 examples for the frame bound estimators can be found for two different windows.

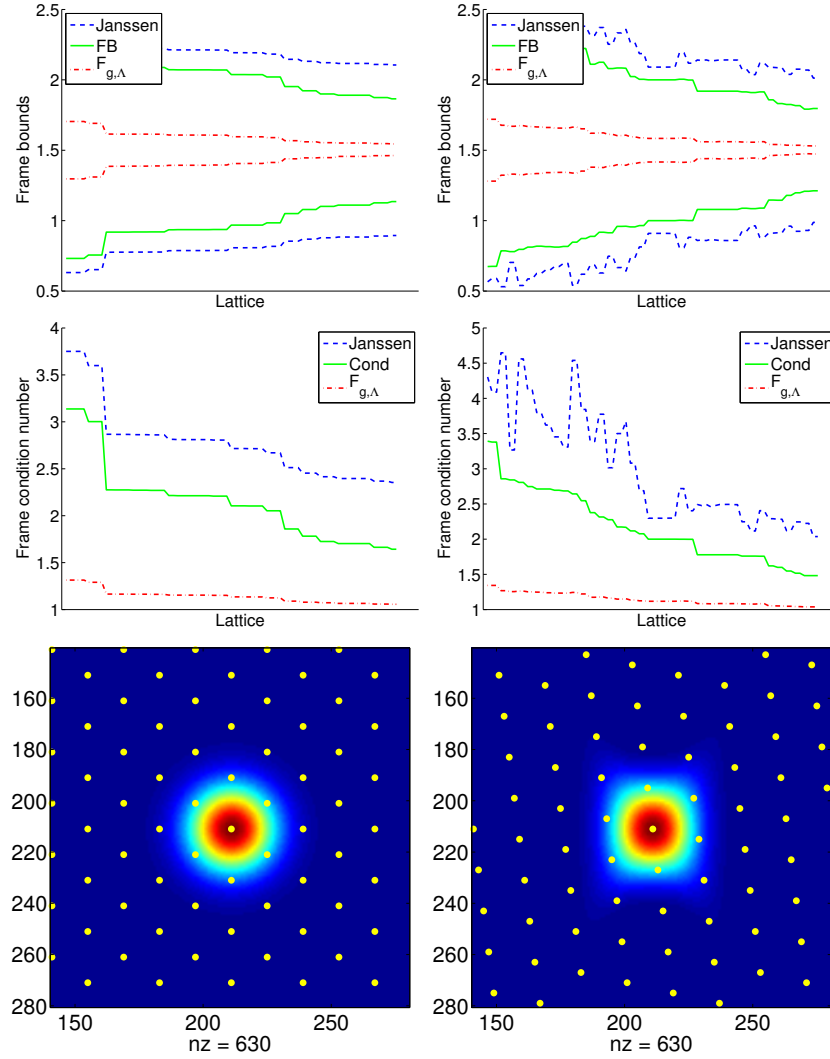


Figure 3.1: Left column: sampled periodized Gaussian window, Right column: Hanning window. The top two plots show the exact frame bounds, as well as the inner and outer estimates for the frame bounds computed by the periodization of the ambiguity function and the Janssen criterion respectively. The signal length is $L = 420$ and the redundancy is fixed to 1.5, all the lattices for that signal length and redundancy are sorted by condition number and the best 70 are shown. The two plots at the bottom show the Gaussian and the Hanning window in the ambiguity plane and the respective optimal lattices.

3.2 Optimal sampling for Gaussian windows

We can use the statements from the Propositions 3.2 and 3.3 to argue that we want to optimally pack the TF-plane with ambiguity functions of the standard Gaussian window g_0 . Since this ambiguity function is again a Gaussian (i.e. radial symmetric) these arguments lead to the conjecture that hexagonal sampling is the best pattern in this case, as the best packing of the plane with circles is hexagonal [18]. The idea that a hexagonal sampling pattern is optimal for the standard Gaussian window is not new and has been conjectured in [30]. Geometric arguments have lead to the same conclusion in [95].

In the following we show some experiments based on the Janssen test that support this conjecture. First of all, we note that hexagonal sampling is not uniquely determined.

Definition 3.5. The *standard hexagonal lattice* of redundancy R is given by

$$\Lambda_0 = \sqrt{\frac{1}{R}} \mathcal{A}_0 \mathbb{Z}^2, \quad (3.12)$$

where

$$\mathcal{A}_0 = \begin{pmatrix} \frac{\sqrt{3}}{\sqrt{2}} & 0 \\ \frac{1}{\sqrt{3}\sqrt{2}} & \frac{\sqrt{2}}{\sqrt{3}} \end{pmatrix}. \quad (3.13)$$

Other hexagonal lattices are generated by $Q\mathcal{A}_0$ for some orthogonal matrix Q .

Now we can use the Janssen test to find numerically which sampling optimizes the frame bound estimate given in (3.2). We only search for lattices with lower triangular generator matrix, all the other lattices can be obtained by rotation from this, as it is easy to see by decomposing a given matrix $A = QL$, where L is a lower triangular matrix and Q is an orthogonal matrix. A numerical evaluation shows that the standard hexagonal lattice is optimal for the standard Gaussian window function, c.f. Figure 3.2.

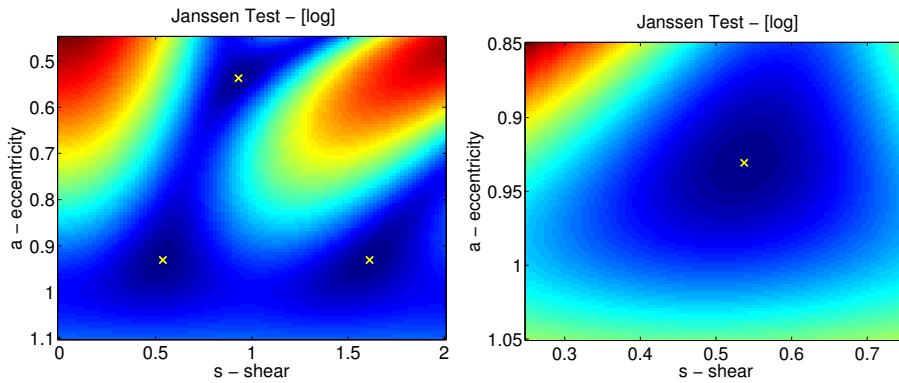


Figure 3.2: Janssen sum for the standard Gaussian function on different lattices; the yellow crosses mark the minima, which all correspond to hexagonal lattices; the right image shows a zoomed version to the region with the standard hexagonal sampling pattern

Proposition 3.6. *The Gabor system $\mathcal{G}(g_0, \Lambda_0)$ has the same frame bounds as $\mathcal{G}(g_0, Q\Lambda_0)$ for any rotation matrix Q .*

Proof. The standard Gaussian function is invariant under the fractional Fourier Transform. The metaplectic operator corresponding to the rotation matrix Q with some rotation angle α is given by $\mathcal{F}^\alpha$. Therefore we find by Corollary 2.16 that

$$S_{g_0, \Lambda_0} = \mathcal{F}^\alpha S_{g_0, Q\Lambda_0} \mathcal{F}^{-\alpha}, \quad (3.14)$$

and therefore by the unitarity of the fractional Fourier Transform the assertion follows. $\square$

Having established the best sampling patterns for the standard Gaussian function, we can use the theory of metaplectic operators to extend this to *generalized Gaussians*, much in the same fashion as presented in [95].

Definition 3.7. For some $a \in \mathbb{R}^+$ and $c \in \mathbb{R}$ the corresponding *generalized Gaussian* function is given by

$$g_{a,c}(t) = e^{i\pi ct^2} D_a g_0(t). \quad (3.15)$$

Proposition 3.8. *An optimal sampling pattern for the generalized Gaussian $g_{a,c}$ with respect to the Janssen criterion is given by*

$$\begin{pmatrix} 1 & 0 \\ c & 1 \end{pmatrix} \begin{pmatrix} a^{-1} & 0 \\ 0 & a \end{pmatrix} \mathcal{A}_0 \cdot \mathbb{Z}^2. \quad (3.16)$$

Proof. The statement easily follows from the fact that the operations applied to g_0 to obtain $g_{a,c}$ are the metaplectic operators to the two symplectic matrices acting on the standard hexagonal sampling matrix in (3.16). $\square$

3.3 Adaption of the window - optimal ambiguity functions

In Section 3.1 we derived some statements that suggest that it is desirable to have ambiguity functions that are concentrated on the fundamental domain of the underlying lattice. In this section we present some concentration measures and an algorithm to optimize with respect to such constraints. Putting constraints on the ambiguity function is not new, see e.g. [96], [63]. However, in the following we present a novel approach that allows for very flexible constraints, which we reported in [35].

3.3.1 Concentration measures

There are different concentration measures that have been used in the literature. Classical examples are Rényi entropies and the Shannon entropy measure [8].

Definition 3.9. Given some window function $g \in L^2(\mathbb{R})$ and $p \geq 0$. Then the *Rényi entropy* is defined as

$$I_g(p) = \|V_g g\|_p^p = \int_{\mathbb{R}^2} |\langle g, \pi(\lambda)g \rangle|^p d\lambda. \quad (3.17)$$

The *Shannon entropy* measure is defined as

$$S_g = - \int_{\mathbb{R}^2} |V_g g(\lambda)|^2 \log(|V_g g(\lambda)|^2) d\lambda. \quad (3.18)$$

For the discrete case both measures are defined analogously by replacing the integrals with sums over the finite discrete phase space.

Proposition 3.10. *The Rényi entropy and the Shannon entropy relate by*

$$S_g = -2I'_g(2). \quad (3.19)$$

Proof. Differentiating (3.17) we find

$$I'(p) = \int_{\mathbb{R}^2} |V_g g(\lambda)|^p \log(|V_g g(\lambda)|) d\lambda. \quad (3.20)$$

Setting $p = 2$ yields the result. $\square$

A bound on $I_g(p)$ can be found in [73], this is called *Lieb's uncertainty principle*. The original proof is for $\mathbb{R}$, however it can be generalized to locally compact abelian groups since the necessary inequalities for the proof, i.e. *Hausdorff-Young inequality* and *Young's inequality* hold true in this more general setting [53]. In the following we state the results for the two locally compact abelian groups that are of interest to us, i.e. $\mathbb{R}$ and $\mathbb{Z}_L$. It is interesting to note that in the discrete setting the minimizing waveforms are fundamentally different. While in the continuous domain we find the same optimizers as for the Heisenberg uncertainty principle, i.e. Gaussian waveforms (c.f. Definition 3.7), in the finite discrete setting it is delta trains optimizing the inequalities.

Theorem 3.11 (Lieb's uncertainty principle). *Let $g \in L^2(\mathbb{R})$. Then*

$$I_g(p) \leq \frac{2}{p} \|g\|_2^{2p} \quad 2 \leq p < \infty \quad (3.21)$$

$$I_g(p) \geq \frac{2}{p} \|g\|_2^{2p} \quad 1 \leq p \leq 2. \quad (3.22)$$

The inequalities saturate if and only if g is a (complex) multiple of a generalized Gaussian.

Let $g \in \mathbb{C}^L$. Then

$$I_g(p) \leq L \|g\|_2^{2p} \quad 2 \leq p < \infty \quad (3.23)$$

$$I_g(p) \geq L \|g\|_2^{2p} \quad 1 \leq p \leq 2. \quad (3.24)$$

The inequalities saturate if and only if g is a translated modulated copy of a delta train of the form $\sum_{l=0}^{L/a-1} \delta_{la}$ for some $a|L$.

Over the last years the l^1 -norm has become a main tool to measure sparsity [47]. Therefore, for $1 \leq p < 2$ Lieb's uncertainty principle can be interpreted as a bound on the sparsity of the ambiguity function. Indeed, the optimizing waveforms in the finite discrete setting are delta trains in the TF-plane (i.e. as sparse as possible) since for $g = \sum_{l=0}^{L/a-1} \delta_{la}$ we find

$$V_g g(x, \xi) = \sum_{m=0}^{L/a-1} \sum_{n=0}^{a-1} \delta_{ma}(x) \delta_{nL/a}(\xi). \quad (3.25)$$

Corollary 3.12. *Let $g \in L^2(\mathbb{R})$ with $\|g\|_2 = 1$. Then*

$$S_g \geq 1. \quad (3.26)$$

Let $g \in \mathbb{C}^L$ with $\|g\|_2 = 1$. Then

$$S_g \geq 0. \quad (3.27)$$

Equality in the continuous and discrete setting is attained for same class of functions as in Theorem 3.11.

Proof. We only proof the continuous case, the discrete result follows completely analogous. We first note that for $p = 2$ the orthogonality relations yield $I_g(2) = 1$ for a normalized function g . Explicitly writing the differential quotient and using Lieb's uncertainty principle we find

$$I'_g(2) = \lim_{\varepsilon \rightarrow 0} \frac{I_g(2) - I_g(2 - \varepsilon)}{\varepsilon} \leq \lim_{\varepsilon \rightarrow 0} \frac{1}{\varepsilon} \left(1 - \frac{2}{2 - \varepsilon} \right). \quad (3.28)$$

It is then easy to see that the limit on the right hand side is the differential quotient evaluating the derivative of the function $f(p) = 2/p$ at 2, i.e. $f'(2) = -1/2$. This yields (3.26) using the relation (3.19). $\square$

3.3.2 Continuous optimization principle

We describe a general scheme for optimizing prescribed time-frequency localization properties. Later we will show that in the finite discrete setting we are not only able to recover well known waveforms but also obtain interesting new window functions satisfying some optimality criterion. The rationale is to set the problem as a variational problem, and optimize a weighted L^2 -norm of the ambiguity function. The problem to solve reads in the continuous case

$$g_{\text{opt}} = \arg \max_{\|g\|=1} \int_{\mathbb{R}^2} m(|V_g g(\lambda)|, \lambda) |V_g g(\lambda)|^2 d\lambda, \quad (3.29)$$

where the *density function* $m : \mathbb{R}^+ \times \mathbb{R}^2 \rightarrow \mathbb{R}^+$ is chosen so as to enforce some specific localization properties. We shall consider here the two following specific cases, but the approach is not limited to these.

1. For finding a g which optimizes the sparsity of $V_g g$, we choose m of the form

$$m(|V_g g|, \lambda) = |V_g g(\lambda)|^{p-2}, \quad \text{for some } p > 2. \quad (3.30)$$

Then (3.29) becomes equivalent to finding the function which maximizes $\|V_g g\|_p^p$. This leads to the Gaussian here (c.f. Theorem 3.11) but our method clearly extends to more general settings than $L^2(\mathbb{R})$. We shall see in the next section that this also leads to discrete versions of the Gaussian function.

2. The other case we will consider is to choose a weight function of the form

$$m(u, \lambda) = m(0, \lambda), \quad (3.31)$$

i.e. a fixed weight in the ambiguity plane not dependent on the ambiguity function itself. In this case, the optimization (3.29) is expected to produce

waveforms g whose ambiguity function is concentrated in TF regions where m takes large values. Given the constraints on ambiguity functions, we limit ourselves to weight functions concentrated around the origin and satisfying the symmetry property $m(-\lambda) = m(\lambda)$. When m is an indicator of a ball the solution of the problem is the Gaussian function [56]. A similar result but for non-quadratic evaluation of the ambiguity function in the optimality criterion (3.29) can be found in [44].

Remark 3.13. The existence of a solution of (3.29) seems a hard problem and presently open for a general concentration function $m(u, \lambda)$. We maximize over the ball in $L^2(\mathbb{R})$, which is not compact. Therefore we cannot immediately guarantee the existence of a solution even though the objective function we seek to maximize given by

$$\Gamma(g) = \int_{\mathbb{R}^2} m(|V_g g(\lambda)|, \lambda) |V_g g(\lambda)|^2 d\lambda \quad (3.32)$$

is bounded on the unit ball of $L^2(\mathbb{R})$ if the density function m is bounded on its domain.

The solution of Equation (3.29) is highly non-linearly dependent on the window function. Therefore, we define the cost functional

$$\Gamma_\mu(g, h) = \int_{\mathbb{R}^2} m(|V_h h(\lambda)|, \lambda) |V_h g(\lambda)|^2 d\lambda - \mu(\|g\| - 1), \quad (3.33)$$

which we seek to maximize for some fixed $h \in L^2(\mathbb{R})$. The constraint $\|g\| = 1$ is added via a Lagrange multiplier, we will also use the notation $\Gamma(g, h) = \Gamma_0(g, h)$. In the following we relate solutions of (3.33) to the optimization problem (3.29). We also note that the weight function used above in the following be used as a weight function for a multiplier, this is a slight abuse of the notation that we introduced in Section 1.5.1.

Proposition 3.14. *Let m be a bounded, positive density function and $h \in L^2(\mathbb{R})$. Then $g \in L^2(\mathbb{R})$ maximizes $\Gamma_\mu(g, h)$ if and only if g satisfies the eigenvalue equation*

$$M_{m,h}g = \int_{\mathbb{R}^2} m(|V_h h(\lambda)|, \lambda) V_h g(\lambda) \pi(\lambda) h d\lambda = \mu g, \quad (3.34)$$

where μ is the maximal value of the spectrum of the multiplier $M_{m,h}$.

Proof. We use the standard perturbation argument from the calculus of variations to conclude that if g maximizes $\Gamma_\mu(g, h)$ for some fixed $h \in L^2(\mathbb{R})$, then

$$\left. \frac{d}{d\varepsilon} \Gamma_\mu(g + \varepsilon \tilde{g}, h) \right|_{\varepsilon=0} = 0, \quad (3.35)$$

for any given $\tilde{g} \in L^2(\mathbb{R})$. Computing the differential and setting ε to zero yields

$$\begin{aligned} \left. \frac{d}{d\varepsilon} \Gamma_\mu(g + \varepsilon \tilde{g}, h) \right|_{\varepsilon=0} &= \int_{\mathbb{R}^2} m(|V_h h(\lambda)|, \lambda) \left(V_h g(\lambda) \overline{V_h \tilde{g}} + \overline{V_h g(\lambda)} V_h \tilde{g}(\lambda) \right) d\lambda \\ &\quad - \mu(\langle \tilde{g}, g \rangle + \langle g, \tilde{g} \rangle). \end{aligned} \quad (3.36)$$

Equating this to 0, we find

$$\begin{aligned} 0 &= \langle M_{m,h}g - \mu g, \tilde{g} \rangle + \langle \tilde{g}, M_{m,h}g - \mu g \rangle \\ &= 2\Re \left(\int_{\mathbb{R}^2} (M_{m,h}g - \mu g) \bar{\tilde{g}} \right), \end{aligned} \quad (3.37)$$

where $\Re f$ denotes the real part of the function f . It is easy to see that this yields (3.34) in the weak sense.

Since the operator $M_{m,h}$ is compact it has a largest eigenvalue (with possible multiplicity), which leads to the desired maximization. $\square$

Remark 3.15. It is not at all trivial to show that the multiplier $M_{m,g}$ has simple eigenvalues. Nevertheless, in the following we will always assume that this is the case, even though there exist constructions of pairs of weight function and window, such that the largest eigenvalue of the corresponding multiplier has degeneracy greater than 1.

The above proposition shows that for some fixed $h \in L^2(\mathbb{R})$ the solution of

$$\arg \max_{\|g\|=1} \Gamma(g, h), \quad (3.38)$$

can be obtained by diagonalizing the operator $M_{m,h}$. This suggests Algorithm 3, which iteratively solves problems of the form (3.38). Finding sufficient conditions for the convergence of this algorithm in $L^2(\mathbb{R})$ is currently an unsolved problem. However, we can show convergence of the cost function in the case of $m(u, \lambda) = m(0, \lambda)$.

Algorithm 3 Algorithm for recursive quadratization of problem (3.29)

Input: m - weight function, h - initial window

Individual Steps:

- 1: **while** 1 **do**
 - 2: $h_{k+1} = \arg \max_{\|g\|=1} \Gamma(g, h_k)$
 - 3: **end while**
-

Proposition 3.16. *Let $m(u, \lambda) = m(\lambda)$ be bounded, positive and even with respect to λ . Then the sequence of numbers $\Gamma^{(k)} = \Gamma(h_{k+1}, h_k)$ is monotonically increasing and convergent.*

Proof. Algorithm 3 computes in the k -th step

$$h_{k+1} = \arg \sup_{\|g\|=1} \langle M_{F,h_k} g, g \rangle. \quad (3.39)$$

Therefore

$$\int_{\mathbb{R}} m(\lambda) |V_{h_k} f(\lambda)|^2 d\lambda \leq \int_{\mathbb{R}} m(\lambda) |V_{h_k} h_{k+1}(\lambda)|^2 d\lambda, \quad (3.40)$$

for any $\|f\| = 1$. We further note that $\Gamma(g, h) = \Gamma(h, g)$ since $|V_g h(\lambda)| = |V_h g(-\lambda)|$ and the weight function satisfies $m(\lambda) = m(-\lambda)$. From (3.40) we conclude that

$$\Gamma(h_{k-1}, h_k) = \Gamma(h_k, h_{k-1}) \leq \Gamma(h_{k+1}, h_k), \quad (3.41)$$

and therefore $\Gamma^{(k-1)} \leq \Gamma^{(k)}$. Furthermore, using the orthogonality relations for the STFT and the fact that all the iterates h_k are normalized we find

$$\Gamma^{(k)} \leq \|m\|_\infty. \quad (3.42)$$

Concludingly, $\Gamma^{(k)}$ is a bounded monotonically increasing sequence and therefore convergent. $\square$

Remark 3.17. For $m(u, \lambda) = m(\lambda) = m(|\lambda|)$ being a radial symmetric decreasing function on phase space Daubechies showed in [24] that

$$\arg \max_{\|g\|=1} \Gamma(g, g_0), \quad (3.43)$$

where g_0 denotes the standard Gaussian is again g_0 . Furthermore, all the eigenvectors of the multiplier M_{m, g_0} are given by the Hermite functions.

In the next step we relate the limiting function (in case of convergence) of Algorithm 3 to the solution of the original problem (3.29).

Proposition 3.18. *Assume that Algorithm 3 converges for any starting point $f \in B$, where B is a subset of the unit ball in $L^2(\mathbb{R})$. Then the limiting function g_{opt} is the unique solution of the problem*

$$g_{\text{opt}} = \arg \max_{g \in B} \int_{\mathbb{R}^2} m(|V_g g(\lambda)|, \lambda) |V_g g(\lambda)|^2 d\lambda. \quad (3.44)$$

Proof. We start by noting that for some $f \in B$

$$\Gamma(f, f) = \int_{\mathbb{R}^2} m(|V_f f(\lambda)|, \lambda) |V_f f(\lambda)|^2 d\lambda. \quad (3.45)$$

Since for this starting value the algorithm converges by assumption we can conclude denoting the iterates by h_k

$$\Gamma(f, f) \leq \Gamma(f, h_1) \leq \Gamma(h_1, h_2) \leq \cdots \leq \Gamma(g_{\text{opt}}, g_{\text{opt}}). \quad (3.46)$$

Uniqueness follows from the fact that for any $g_{\text{opt}} \neq f \in B$ we find that

$$\Gamma(f, f) < \Gamma(f, h_1), \quad (3.47)$$

otherwise the algorithm would converge to f contradicting the assumption. $\square$

Remark 3.19. As mentioned above, the convergence of the procedure is an open problem. However, for the case $m(u, \lambda) = m(\lambda)$ being an even function with compact support in the TF-plane it seems reasonable to conjecture that the set $\{h_k : k \in \mathbb{N}\}$ is compact in $L^2(\mathbb{R})$. Characterizations of compact sets in $L^2(\mathbb{R})$, that look promising for this problem can be found in [31].

We have discussed an algorithm to obtain some optimally localized function g_{opt} . Given that the algorithm converges we can diagonalize the operator $M_{m, g_{\text{opt}}}$ and consider its eigenfunctions. The eigenfunction to the largest eigenvalue will recover g_{opt} .

Definition 3.20. Let $g = \mu M_{m,g}$ for some weight function on the TF-plane, where μ is the largest eigenvalue of the multiplier. Then we call the eigenfunction to the k -th largest eigenvalue of $M_{m,g}$ the k -th *generalized Hermite function*.

Remark 3.21. This is indeed a generalization of the classical Hermite functions because for m radial symmetric and decreasing on the TF-plane and $g = g_0$ the standard Gaussian one recovers the Hermite functions. See also Remark 3.17 or [24].

Remark 3.22. STFT multipliers with radial symmetric weight commute with the Hamiltonian operator [24], therefore they have the same eigendecomposition. For more general weight functions it is an open question if there exists a differential operator commuting with the corresponding multiplier. This would allow to use the rich theory of partial differential equations to tackle the problem of existence of solutions of (3.29).

3.3.3 Numerical experiments

The above described algorithm can also be formulated for the finite discrete setting. All the convergence results carry over verbatim from $L^2(\mathbb{R})$ to $\mathbb{C}^L$. A restatement from Algorithm 3 can be found in Algorithm 4.

Algorithm 4 compute the optimal window ψ_{opt}

Input: ϕ - initial window, F - weight, eps .

Individual Steps:

- 1: $\psi_0 = 0, \psi_1 = \phi$
 - 2: **while** $\|\psi_1 - \psi_0\| > \text{eps}$ **do**
 - 3: $\psi_0 \leftarrow \psi_1$
 - 4: Compute the STFT-multiplier with weight F and window ψ_0
 - 5: Compute ψ_1 as the eigenvector to the largest eigenvalue
 - 6: Update F if necessary
 - 7: **end while**
 - 8: $\psi_{\text{opt}} = \psi_1$
-

In the following we show and discuss some numerical experiments for the following two choices of the weight function m . First we consider the case of a weight, that is fixed over the iterations, i.e. $m(u, \lambda) = m(\lambda)$. Secondly we consider a weight that depends on the ambiguity function of the current iterate, i.e. $m = m(|V_{h_k} h_k|)$. We look at two specific examples (c.f. Figure 3.3):

1. $m(\lambda) = f(\|\lambda\|_2)$: This is a radial symmetric weight and one obtains a Gaussian-like function. However, this is not a perfect match with the sampled periodized Gaussian function suggested in Section 1.6.1 (c.f. Remark 3.23). The generalized Hermite functions with respect to this weight form rings in the TF plane.
2. $m(\lambda) = f(\|\lambda\|_\infty)$: This weight function has its contour lines on quadratic boxes around the origin. Also in this case the algorithm converges, one obtains still a Fourier invariant function. This weight produces a completely different set of generalized Hermite functions, that satisfy a different concentration constraint. The generalized Hermite functions for this choice form a square box in the TF plane.

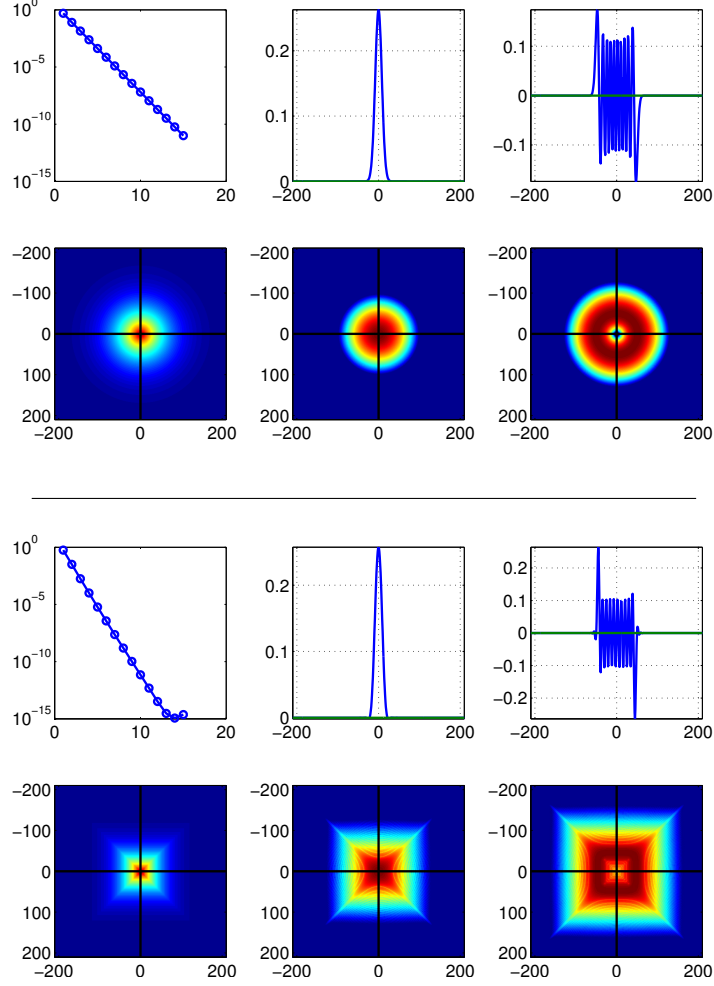


Figure 3.3: Results in the case of a TF mask of the form $m(|V_g g(\lambda)|, \lambda) = m(0, \lambda)$. Top 6 plots: radial symmetric mask. Bottom 6 plots: $m(\lambda) = f(\|\lambda\|_\infty)$. The following explanation holds for both blocks separately. Top left: $\|h_{n+1} - h_n\|_2$, log scale. Bottom left: TF mask m . Middle: optimal function in time (top) and in the TF plane through its ambiguity function (bottom). Right: an eigenvector of the last STFT-multiplier, the “generalized” Hermite function H_{20} in time (top) and in the TF plane through its ambiguity function (bottom). The signal length is $L = 420$.

Remark 3.23. Running the algorithm for a radial symmetric weight imitates the situation of classical Hermite functions. It is interesting to remark here that these functions obtained as eigenfunctions of an STFT multiplier do not exactly give the same results as the functions introduced in Section 1.3.2. Both of these sets of functions fulfill the main properties one would expect from a sensible definition of Hermite functions. However, the Hermite functions from the first chapter can be computed as eigenfunctions of a tridiagonal, i.e. a very sparse matrix. Therefore it is preferable to choose the earlier definition for computation, as it is possible to speed up the diagonalization algorithms utilizing the sparsity of the matrix [4]. In Figure 3.4 the correlation between the different versions of Hermite functions is shown, as well as an expansion of the sampled periodized Gaussian in terms of the discrete Hermite functions. The Gaussian function obtained as limit of the iterative procedure and the sampled periodized Gaussian differ in norm by 10^{-10} .

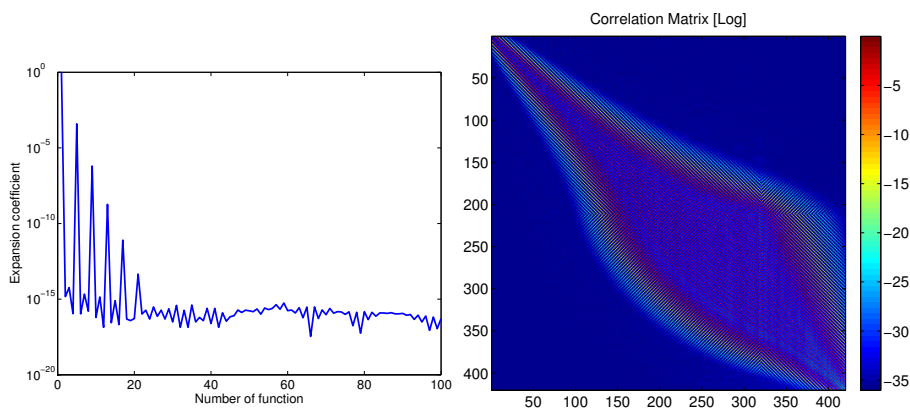


Figure 3.4: Left: expansion coefficients of the sampled periodized Gaussian with respect to the discrete Hermite functions from Section 1.3.2, Right: Correlation coefficients of the generalized Hermite functions from this section with the discrete Hermite functions from Section 1.3.2

Now we turn our attention to the optimization of the p -norm of the ambiguity function through the choice $m(|V_g g|, \lambda) = |V_g g(\lambda)|^{p-2}$ in the optimization (3.29). For this case we have less theoretical results describing the behavior of Algorithm 4, nevertheless it is interesting to discuss the outcome of some numerical experiments. We find that the result of the algorithm depends on the initial window. The maximum of the objective function Γ_μ is very degenerate, because there are numerous solutions of the optimization problem (3.29). Liebs uncertainty principle for the discrete case, i.e. Theorem 3.11 allows to explicitly determine all the maximizing waveforms. The algorithm seems to converge to the maximum closest to the initialization. For example, starting with a widespread signal (e.g. white noise) generally leads to a picket fence, or more simply a constant function (a Kronecker delta in the Fourier domain). The solution for a mildly localized starting window is a Gaussian. In Figure 3.5 we show the results for different starting windows, which suggests that in this case the Gaussian window functions are local maxima of the cost function. Also, Gaussian window functions of different widths are shown to be the limits

of Algorithm 4.

Remark 3.24. A very similar approach to the p -norm maximization shown in this section, that allows to optimize the window for some given signal is presented in [84]. The main goal of this publication is to optimize problems of the form

$$g_{\text{opt}} = \arg \max_{\|g\|=1} \|V_g f\|_p^p. \quad (3.48)$$

This is strongly related to the maximization presented in Equation (3.29), the motivation to solve such a problem is slightly different. For a given signal f this optimization determines the window function g , such that the time-frequency representation becomes maximally sparse. However, the numerical realization chosen there is more sophisticated and is suitable for larger signal lengths than we have discussed here.

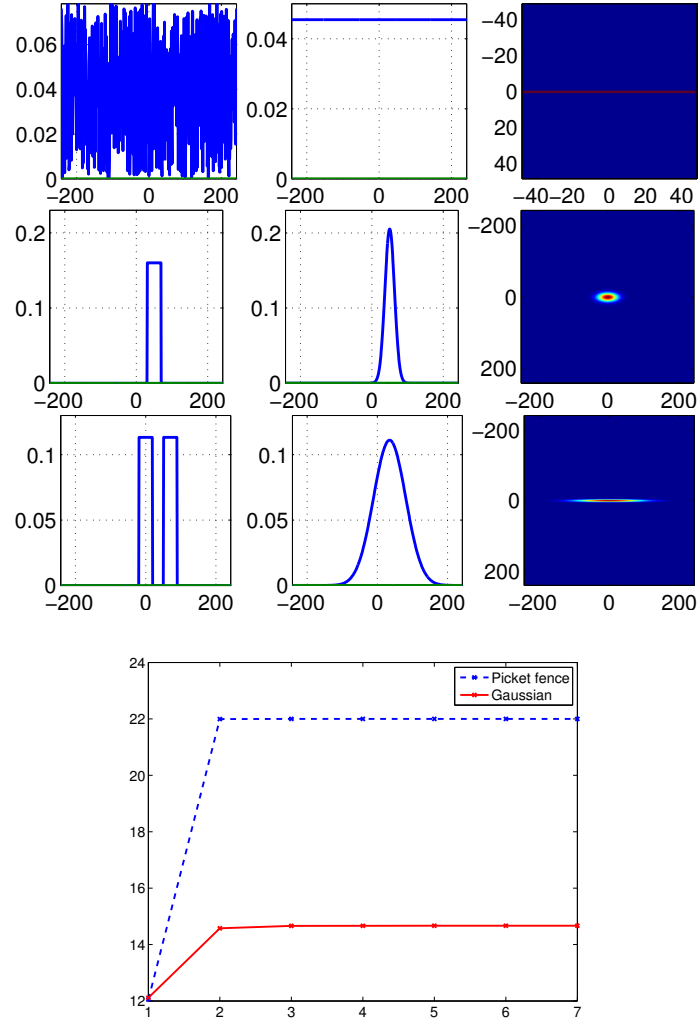


Figure 3.5: Top 3×3 array of plots: On the left the initial window is plotted, in the middle the limit and on the right the ambiguity function of the limit (the top right plot is obtained by zooming into the ambiguity plane to make the line better visible). The bottom plot shows the p -norm for $p = 3$ of the iterates for the first two starting windows above. The signal length is $L = 484$.

Chapter 4

Beyond Gabor - Wavelets and ERB-lets

In the previous chapters we discussed Gabor frames in detail, the classical tool of time-frequency analysis. This chapter is devoted to a flexible method to construct well-conditioned or even tight time-frequency frames. The systems under investigation are so-called *non-stationary Gabor* systems, first introduced in [7]. While in standard Gabor theory shifted copies of a prototype window function are used to localize the signal f , in this generalized setting one can use different windows at each time position. For each of these windows one specifies a modulation step parameter, meaning that the spacing in frequency stays regular to keep some of the structure of ordinary Gabor systems. The same principle can be applied on the frequency side to obtain *frequency side non-stationary Gabor* systems. These systems, adaptive in frequency, are also called *filter banks*.

In this chapter we use a *warping* of the frequency axis to obtain a transform adapted to a desired frequency scaling. Wavelets for example can be realized using a logarithmic warping of the frequency axis. This fact has first been used implicitly by mapping a Gabor system to a wavelet system to proof the so-called *painless conditions* for wavelets in [27]. The same idea has recently been further developed in [17]. While the mentioned papers only consider a logarithmic transformation of the frequency axis, we will relax this idea to allow for more general transformations. Especially promising for applications is the *ERB-scale*, as it describes filters adapted to the human perception of sound. A corresponding transform has potential applications in audio signal processing.

Recent efforts have shown that it is possible to extend the principles of time-frequency analysis and filter banks to graphs, see [88]. The main tool is the so-called graph Laplacian operator, the diagonalization of which yields the spectrum and the corresponding pure frequencies on the graph. The spectrum is purely discrete and the points are in general non-uniformly distributed between the lowest and the largest eigenvalue. Using the idea of warping it is possible to adapt a filter bank to the specific distribution of the eigenvalues of the graph. The exact details of our construction are beyond the scope of this thesis, the interested reader can find the details online in [89].

4.1 Filter banks via warping

Filter banks are a widely known and used tool in the signal processing community [23, 102, 105], we are mostly interested in the frame properties of such systems.

Definition 4.1. Let $\mathbf{g} = \{g_m\}_{m \in \mathbb{Z}} \subset L^2(\mathbb{R})$ and $\mathbf{a} = \{a_m\}_{m \in \mathbb{Z}}$ a sequence of corresponding downsampling factors. Then we call the system

$$\mathcal{G}(\mathbf{g}, \mathbf{a}) = \{g_{m,n}\}_{m,n \in \mathbb{Z}}, \quad g_{m,n} := T_{na_m} \widetilde{g}_m, \quad (4.1)$$

a *non-uniform filterbank*. The windows $\mathbf{g}$ are called *filter transfer functions*.

Filterbanks have received considerable attention in literature, also under different names. The frame properties of filterbanks have been studied e.g. in [10], as generalized shift-invariant systems in [62, 86] and as (frequency-side) nonstationary Gabor systems in [7, 32, 33, 64].

Remark 4.2. In Definition 4.1 we defined filterbanks on $L^2(\mathbb{R})$, which is the standard setting. In the following, we consider filterbanks of $\mathcal{F}^{-1}L^2(D)$ instead, where $D \subseteq \mathbb{R}$ is an interval, in our examples we only consider $D = \mathbb{R}$ and $D = \mathbb{R}^+$. This generalization is easily explained by just requiring that $g_m \in L^2(D)$ and making the appropriate changes in the definition above.

Having the definition of filter banks at hand, we turn our attention to a special method of constructing them. The method is based on the premise that the filter transfer functions should be uniformly spaced in the desired frequency scale. Later, we show that the suggested construction encompasses regular Gabor transforms, in which the filter transfer functions are just translates of a prototype transfer function, and wavelets, where the filter transfer functions are dilates of a prototype. However, the method provides enough flexibility to construct a large class of transforms adapted to a variety of frequency scales. The frequency scale is determined by a *warping function* $F : D \rightarrow \mathbb{R}$, which should be increasing for some given (possibly unbounded) interval D . To warp a regular filter bank given by $T_m \theta$ we consider the windows

$$\{\theta_{F,m}\}_{m \in \mathbb{Z}}, \quad \text{where } (T_m \theta) \circ F, \quad (4.2)$$

for some warping function determining the frequency scale of interest. In Figure 4.1 the warping for $F = \log$ is visualized, as we will see later, this constructs wavelets. In order to make the theory of warped filterbanks precise we need to assume some more conditions on the warping function F .

We start by introducing weight functions for the later treatment of the appearing function spaces. For a more in depth analysis of weight functions and their role in harmonic analysis we refer to [55].

Definition 4.3. • A weight function $v : \mathbb{R} \rightarrow \mathbb{R}^+$ is called *submultiplicative* if

$$v(x+y) \leq v(x) + v(y). \quad (4.3)$$

• A weight function $v : \mathbb{R} \rightarrow \mathbb{R}^+$ is called *v-moderate* if

$$w(x+y) \leq C v(x) w(y), \quad (4.4)$$

for some submultiplicative weight function v and constant C .

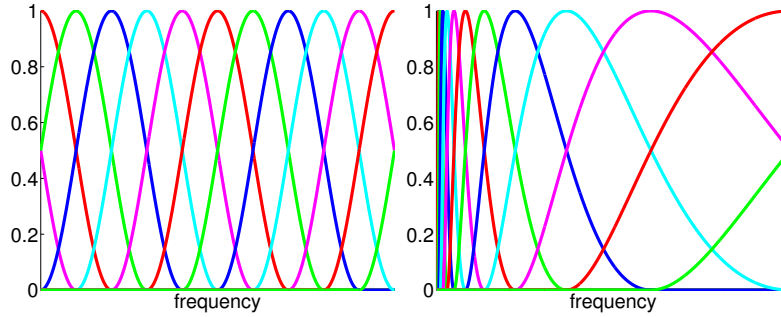


Figure 4.1: Left: unwarped system of translates $T_m \theta$, Right: warped system $(T_m \theta) \circ \log$.

Submultiplicative and moderate weight functions are an important concept in the theory of function spaces, as they are closely related to the translation invariance of the corresponding weighted spaces [36, 54]. This definition also allows us to explicitly state our conditions on the function F .

Definition 4.4. A bijective function $F : D \rightarrow \mathbb{R}$ is called *warping function*, if $F \in \mathcal{C}^1(D)$ with $F' > 0$ and the associated weight function given by

$$w(t) = (F^{-1})'(t) = \frac{1}{F'(F^{-1}(t))}, \quad (4.5)$$

is v -moderate for some submultiplicative weight v .

Proposition 4.5. Let $F : D \rightarrow \mathbb{R}$ be a warping function and $\theta \in L_w^2$. Then

- the space L_w^2 is translation invariant.
- $\|(T_m \theta) \circ F\|_{L^2(D)} = \|T_m \theta\|_{L_w^2(\mathbb{R})} \leq C v(m) \|\theta\|_{L_w^2(\mathbb{R})}$, for some submultiplicative v .

Proof. The first item is an easy consequence from the condition that the weight function w is moderate by investigating the integral

$$\|T_m f\|_{L_w^2} = \int_{\mathbb{R}} |f(t)|^2 w(t+m) dt \leq C v(m) \|f\|_{L_w^2}, \quad (4.6)$$

for some submultiplicative function v .

The second statement follows by a substitution

$$\begin{aligned} \|\theta \circ F\|_{L^2} &= \int_D |\theta(F(t))|^2 dt \\ &= \int_{\mathbb{R}} |\theta(u)|^2 (F^{-1})'(u) du \end{aligned} \quad (4.7)$$

□

Definition 4.6. Let $F : D \rightarrow \mathbb{R}$ be a warping function and $\theta \in L_w^2(\mathbb{R}) \cup L^\infty(\mathbb{R})$. Furthermore, let $\mathbf{a} = \{a_m\}_{m \in \mathbb{Z}}$ be a set of downsampling factors. Then the *warped filterbank* associated to θ and $\mathbf{a}$ is given by

$$\mathcal{G}(\theta, F, \mathbf{a}) = \{T_{na_m} \widetilde{g_m}\}_{m, n \in \mathbb{Z}}, \quad (4.8)$$

with

$$g_m = \sqrt{a_m} \theta_{F, m}, \quad (4.9)$$

where the functions $\theta_{F, m}$ are defined in Equation (4.2).

We further call $a_m = \tilde{a}/(Cv(m))$ (C and v stem from the moderateness of w) for some $\tilde{a}$ *natural downsampling factors*, as they lead to bounded L^2 -norms of the filter transfer functions.

Before we start to investigate the frame properties of warped filter banks we discuss some prototypical warping functions. There are two essentially different types of warping functions we consider, one where the domain is $D = \mathbb{R}^+$ and another one, where $D = \mathbb{R}$. In the latter case we always assume point symmetry around 0 for simplicity.

Proposition 4.7. *The following functions are warping functions:*

1. $F : \mathbb{R}^+ \rightarrow \mathbb{R}$, $t \mapsto c \log(t/d)$ for some $c, d > 0$.
2. $F : \mathbb{R} \rightarrow \mathbb{R}$, $t \mapsto \text{sign}(t) c \log(1 + |t|/d)$ for some $c, d > 0$.
3. $F : \mathbb{R} \rightarrow \mathbb{R}$, $t \mapsto \text{sign}(t) (|t| + 1)^l - 1$ for some $l \in (0, 1)$.

Proof. In the first case we find that

$$w(x) = (F^{-1})'(x) = \frac{d}{c} e^{x/c}. \quad (4.10)$$

For this weight function we find $w(x + y) = e^{x/c} w(y)$, it is easy to see that the weight function $v(x) = e^{x/c}$ is submultiplicative. Therefore w is indeed v -moderate.

The second function is easily identified as being once continuously differentiable and point symmetric around 0. The corresponding weight function is given by

$$w(x) = (F^{-1})'(x) = d e^{|x|/c}. \quad (4.11)$$

Analogously to the weight function from the first case we find that $w(x + y) = e^{|x|/c} w(y)$, which shows that w is indeed moderate with respect to the submultiplicative weight function $v(x) = e^{|x|/c}$.

For the last warping function we also start by noting that it is in $\mathcal{C}^1(\mathbb{R})$, further we find

$$w(x) = (F^{-1})'(x) = \frac{1}{l} (|x| + 1)^{1/l-1}. \quad (4.12)$$

Moreover, we define the weight function $v(x) = (|x| + 1)^{1/l-1}$, which is submultiplicative since

$$(1 + |x + y|)^{1/l-1} \leq (1 + |x| + |y| + |xy|)^{1/l-1} = (|x| + 1)^{1/l-1} (|y| + 1)^{1/l-1} \quad (4.13)$$

Using the same argument as in the equation above for the weight function w shows that $w(x + y) \leq v(x) w(y)$. $\square$

The first two warping functions discussed in Proposition 4.7 are of special interest, as they encompass two important frequency scales. Therefore they are subject to more detailed discussion in the following examples.

Example 4.8 (Wavelets). Choosing $F = \log$ leads to a system of the form

$$\theta_{F,m}(t) = \theta(\log(t) - m) = \theta_{F,0}(\log(te^{-m})). \quad (4.14)$$

This warping function therefore leads to $\theta_{F,m}$ being a dilated version of $\theta_{F,0}$ up to normalization. An F of this form has been discussed in the first case of Proposition 4.7 and a good choice for the subsampling factors is $a_m = a_0/v(m) = a_0e^{-m}$. This shows that $\mathcal{G}(\theta, \log, a_0e^{-m})$ is indeed a wavelet system on the frequency side with the minor modification that our scales are reciprocal to the usual definition of wavelets.

Example 4.9 (ERB-lets). In psychoacoustics the investigation of filter banks adapted to the spectral resolution of the human ear has been subject to a wealth of research, see [74] for an overview. We mention here the Equivalent Rectangular Bandwidth scale (ERB-scale) described in [51], which introduces a set of bandpass filters following the human perception. In [76] the authors construct a filter bank that is designed to be adapted to the ERB-scale, in the following we show how to construct a similar filter bank but with better control over the frame properties. In our terminology the ERB warping function is given by

$$F_{\text{ERB}}(t) = \text{sign}(t) c_1 \log \left(1 + \frac{|t|}{c_2} \right), \quad (4.15)$$

where the constants are given by $c_1 = 9.265$ and $c_2 = 228.8$. Proposition 4.7 shows that F_{ERB} is a warping function in the sense of Definition 4.4. The transform based on an ERB filter bank has potential applications in sound signal processing, as it provides a perfectly invertible transform adapted to the human perception of sound.

4.2 Warped frames

The purpose of this section is the derivation of necessary and sufficient conditions for warped systems as introduced in Definition 4.6 to be frames. We also show that for reasonable choices of θ and F with a natural set of downsampling factors the set $\mathcal{G}(\theta, F, \mathbf{a})$ is a Bessel sequence.

We start with a collection of results for general non-stationary Gabor systems. The following proposition combines [7, Cor 1] with [64, Pro 2]. The lower bound in Proposition 4.10(i), not considered in [64] can be proven in the same way as the upper bound.

Proposition 4.10. *Let $\mathcal{G}(\mathbf{g}, \mathbf{a})$ be a filterbank.*

(i) *If $\mathcal{G}(\mathbf{g}, \mathbf{a})$ is a frame with frame bounds $0 < A \leq B < \infty$, then*

$$0 < A \leq \sum_m \frac{1}{a_m} |g_m|^2 \leq B < \infty, \text{ almost everywhere.} \quad (4.16)$$

(ii) Assume that, there are some constants $c_m, d_m \in \mathbb{R}$, such that $\text{supp}(g_m) \subseteq [c_m, d_m]$ and a_m satisfies $a_m^{-1} \geq d_m - c_m$, for all $m \in \mathbb{Z}$. Then $\mathcal{G}(\mathbf{g}, \mathbf{a})$ forms a frame if and only if (4.16) holds. Furthermore, $\mathcal{G}(\tilde{\mathbf{g}}, \mathbf{a})$, where

$$\tilde{g}_m = \frac{g_m}{\sum_k \frac{1}{a_m} |g_m|^2}, \quad (4.17)$$

is a dual frame for $\mathcal{G}(\mathbf{g}, \mathbf{a})$.

We refer to the setup in Proposition 4.10(ii) as the *painless case*, a generalization of the classical painless non-orthogonal expansions is introduced in [27].

Corollary 4.11. Let $\mathcal{G}(\theta, F, \mathbf{a})$ such that $\text{supp}(\theta) \subseteq [c, d]$ for some constants $0 < c < d < \infty$. If $a_m^{-1} \geq F^{-1}(d + m) - F^{-1}(c + m)$ for all m , then $\mathcal{G}(\theta, F, \mathbf{a})$ forms a frame with frame bounds A and B , if and only if

$$0 < A \leq \sum_{m \in \mathbb{Z}} |\mathbf{T}_m \theta|^2 \leq B < \infty. \quad (4.18)$$

In that case, the canonical dual frame is of the form $\mathcal{G}(\tilde{\theta}, F, \mathbf{a})$, with

$$\tilde{\theta}(t) = \frac{\theta(t)}{\sum_m |\mathbf{T}_m \theta(t)|^2}, \quad \text{a.e.} \quad (4.19)$$

Proof. Insert $g_m = \sqrt{a_m} (\mathbf{T}_m \theta) (F(\cdot))$ into (4.16) and (4.17). $\square$

Similar to Proposition 4.10, the following theorem is a generalization of a sufficient frame condition for Gabor systems [54], first proven by Daubechies in [25]. In the following we need to extend functions defined on the interval D , possibly a subset of $\mathbb{R}$, to the whole real line. At any place necessary this is realized by extending functions with zeros.

Theorem 4.12. Let $\mathcal{G}(\theta, F, \mathbf{a})$ be a warped filterbank with $\mathbf{a}$ being a set of natural downsampling factors. If

$$A := \min_{t \in D} \left[\sum_{m \in \mathbb{Z}} \frac{1}{a_m} \left(|g_m(t)|^2 - \sum_{k \neq 0} |g_m(t) \overline{g_m}(t - k/a_m)| \right) \right] > 0 \quad (4.20)$$

$$B := \max_{t \in D} \left[\sum_{m \in \mathbb{Z}} \sum_{k \in \mathbb{Z}} \frac{1}{a_m} |g_m(t) \overline{g_m}(t - k/a_m)| \right] < \infty, \quad (4.21)$$

almost everywhere, then $\mathcal{G}(\theta, F, \mathbf{a})$ constitutes a frame for $\mathcal{F}^{-1}L^2(D)$. Furthermore, we have $A \leq A_0 \leq B_0 \leq B$, where A_0, B_0 are the optimal frame bounds of $\mathcal{G}(\theta, F, \mathbf{a})$.

Proof. We first note that $\mathcal{G}(\theta, F, \mathbf{a})$ constituting a frame for $\mathcal{F}^{-1}(L^2(D))$ with frame bounds A, B is equivalent to $\mathcal{G}(\theta, F, \mathbf{a}) = \{M_{na_m} g_m\}_{n, m \in \mathbb{Z}}$ constituting a frame for $L^2(D)$ with the same frame bounds by Parseval's formula. In the following we use the latter formulation of the problem to prove the theorem. Let

$$F_m := \sum_{n \in \mathbb{Z}} T_{ka_m}(f g_m). \quad (4.22)$$

Since $f, g_m \in L^2(\mathbb{R})$ we obtain $F_m \in L^1([0, a_m^{-1}])$, using Cauchy-Schwarz' inequality. We now assume that f has compact support and is bounded, which shows that $F_m \in L^2([0, a_m^{-1}])$ by observing that

$$|F_m(t)|^2 \leq \sum_{k \in \mathbb{Z}} |T_{ka_m^{-1}} f(t)|^2 \cdot \sum_{k \in \mathbb{Z}} |T_{ka_m^{-1}} g_m(t)|^2, \quad (4.23)$$

where the first sum is bounded because of the assumptions on f . Parseval's equality now yields

$$\sum_{n \in \mathbb{Z}} \left| \int_0^{1/a_m} F_m(t) e^{2\pi i n a_m t} dt \right|^2 = \frac{1}{a_m} \int_0^{1/a_m} |F_m(t)|^2 dt.$$

We now rewrite the norm of the coefficient sequence into a main term and remainder using Poisson's formula (after the computations we rectify the interchange of summations and integrals)

$$\begin{aligned} \sum_{m, n \in \mathbb{Z}} |\langle f, M_{na_m} g_m \rangle|^2 &= \sum_{m, n \in \mathbb{Z}} \left| \int_0^{1/a_m} F_m(t) e^{2\pi i n a_m t} dt \right|^2 \\ &= \sum_{m \in \mathbb{Z}} \frac{1}{a_m} \int_0^{1/a_m} |F_m(t)|^2 dt \\ &= \sum_{m \in \mathbb{Z}} \frac{1}{a_m} \int_0^{1/a_m} \left(\sum_{k \in \mathbb{Z}} T_{ka_m^{-1}} (fg_m)(t) \right) \left(\sum_{l \in \mathbb{Z}} T_{la_m^{-1}} \overline{(fg_m)}(t) \right) dt \\ &= \sum_{m \in \mathbb{Z}} \int_D \frac{1}{a_m} f(t) g_m(t) \left(\sum_{l \in \mathbb{Z}} T_{la_m^{-1}} \overline{(fg_m)}(t) \right) dt \\ &= \int_D \sum_{m \in \mathbb{Z}} \frac{1}{a_m} |f(t)|^2 |g_m(t)|^2 dt + \underbrace{\int_{\mathbb{R}} \sum_{m \in \mathbb{Z}} \sum_{l \neq 0} \frac{1}{a_m} (fg_m T_{la_m^{-1}} \overline{(fg_m)})(t) dt}_{=: R(t)}. \end{aligned} \quad (4.24)$$

In the second line of the equation above the summation over k is finite since we assumed f to have compact support, therefore we can interchange integration and summation. The last interchange in the second to last line is rectified by observing that

$$\begin{aligned} &\sum_{m \in \mathbb{Z}} \int_D \frac{1}{a_m} |f(t) g_m(t)| \left| \sum_{l \in \mathbb{Z}} T_{la_m^{-1}} \overline{(fg_m)}(t) \right| dt \\ &\leq \sum_{l \in \mathbb{Z}} \sum_{m \in \mathbb{Z}} \int_D |f(t)| |f(t - l/a_m)| |\theta(F(t) - m)| |\theta(F(t - l/a_m) - m)| dt < \infty, \end{aligned} \quad (4.25)$$

since θ is essentially bounded and therefore the integral can be estimated by $\|f\|_2^2$.

To obtain the desired result, we now estimate the rest term R . Begin by

applying Cauchy-Schwarz inequality twice:

$$\begin{aligned}
|R| &\leq \sum_{m \in \mathbb{Z}} \frac{1}{a_m} \sum_{l \neq 0} \int_D \left(|f(t)| |T_{la_m^{-1}} f(t)| |g_m(t) T_{la_m^{-1}} \overline{g_m}(t)| \right) (t) dt \\
&\leq \sum_{m \in \mathbb{Z}} \frac{1}{a_m} \left(\sum_{l \neq 0} \int_D |f(t)|^2 |g_m T_{la_m^{-1}} \overline{g_m}(t)| dt \right)^{1/2} \\
&\quad \left(\sum_{l \neq 0} \int_D |T_{la_m^{-1}} f(t)|^2 |g_m T_{la_m^{-1}} \overline{g_m}(t)| dt \right)^{1/2}.
\end{aligned} \tag{4.26}$$

By the substitution $u = t - la_m^{-1}$ and an index shift in the sums over l , we find the two factors in the estimate above to be the same. This leads to the final estimate of the rest term

$$|R| \leq \int_D |f(t)|^2 \sum_{l \neq 0} \sum_{m \in \mathbb{Z}} \frac{1}{a_m} |g_m T_{la_m^{-1}} \overline{g_m}(t)| dt.$$

The result now follows by applying the triangle equality and pointwise estimation of the sum in the integrand.

For a general function $f \in L^2(D)$ the result follows by a density argument. $\square$

The condition above is a mild generalization of the painless case condition in Proposition 4.10(ii) and easily satisfied whenever the g_m have sufficient overlap and are *almost* compactly supported on intervals of the length a_m^{-1} . This implies small values and good decay outside that interval, i.e. there exist c_m and $d_m = c_m + a_m^{-1}$ such that $\|g_m|_{\mathbb{R} \setminus [c_m, d_m]}\|_2 \ll \|g_m\|_2$ for all m . In the following we turn our attention to a sufficient condition on the function θ , such that $\mathcal{G}(\theta, F, \mathbf{a})$ for some natural downsampling factors is a Bessel sequence. For the investigation we need an auxiliary result on warping functions as introduced in Definition 4.4.

From now on we always restrict ourselves to two cases, one where $D = \mathbb{R}^+$ and in the other case we assume $D = \mathbb{R}$ and the warping function is point symmetric around 0. It would be possible to consider other types of warping functions also but in order to keep the technical effort as low as possible we restrict ourselves to the treatment of warping functions following largely the behavior of the interesting examples presented in Proposition 4.7.

Lemma 4.13. *Let $F : D \rightarrow \mathbb{R}$ be a warping function, such that $(F^{-1})'$ is v -moderate and either $D = \mathbb{R}^+$ or $D = \mathbb{R}$ and F point symmetric around 0. Then for any $x, y \geq 0$*

$$F(y) + F(x + F^{-1}(0)) \leq F(y + Cv(F(y))x), \tag{4.27}$$

where C denotes the constant from the moderateness.

Proof. By the moderateness of the derivative of the inverse function we find that

$$\int_0^{\tilde{x}} (F^{-1})'(u + \tilde{y}) du \leq Cv(\tilde{y}) \int_0^{\tilde{x}} (F^{-1})'(u) du. \tag{4.28}$$

Using the fundamental theorem of calculus one arrives at the inequality

$$\tilde{x} + \tilde{y} \leq F(Cv(\tilde{y})(F^{-1}(\tilde{x}) - F^{-1}(0)) + F^{-1}(\tilde{y})). \quad (4.29)$$

Setting $\tilde{x} = F(x + F^{-1}(0))$ and $\tilde{y} = F(y)$ yields the assertion

$$F(x + F^{-1}(0)) + F(y) \leq F(Cv(F(y))x + y) \quad (4.30)$$

□

Theorem 4.14. *Let $F : D \rightarrow \mathbb{R}$ be a warping function and either*

- $D = \mathbb{R}^+$ with non-increasing derivative, or
- $D = \mathbb{R}$, F is point symmetric around 0, F' is non-increasing on $\mathbb{R}^+$.

Furthermore, assume that $\theta \in \mathcal{O}((1 + |F^{-1}(x)|)^{-1-\epsilon})$, for some $\epsilon > 0$ and satisfies (4.18) for some $0 < A \leq B < \infty$. For any natural set of downsampling factors $\mathbf{a}$, the warped filterbank $\mathcal{G}(\theta, F, \mathbf{a})$ forms a Bessel sequence.

Proof. We first restrict ourselves to the case $D = \mathbb{R}$ with a point symmetric warping function, at the end of the proof we explain how to adapt the steps for the other case. Instead of proving the Bessel condition directly, we show that the sufficient condition (4.21) from Thm 4.12 is satisfied. First, we rewrite that condition as follows

$$\begin{aligned} \infty &> \operatorname{ess\,sup}_{t \in D} \left[\sum_{m \in \mathbb{Z}} \frac{1}{\sqrt{a_m}} |g_m(t)| \sum_{k \in \mathbb{Z}} \frac{1}{\sqrt{a_m}} \overline{g_m(t - k/a_m)} \right] \\ &= \operatorname{ess\,sup}_{t \in D} \left[\sum_{m \in \mathbb{Z}} \mathbf{T}_m |\theta(F(t))| \sum_{k \in \mathbb{Z}} \mathbf{T}_m |\theta(F(t - k/a_m))| \right]. \end{aligned} \quad (4.31)$$

In the next step, we use $\theta \in \mathcal{O}((1 + F^{-1}(|t|))^{-1-\epsilon})$ to estimate the inner sum.

$$\sum_{k \in \mathbb{Z}} \mathbf{T}_m |\theta(F(t - k/a_m))| \leq C \sum_{k \in \mathbb{Z}} (1 + |F^{-1}(F(t - k/a_m) - m)|)^{-1-\epsilon}. \quad (4.32)$$

For $m = 0$, we get the following estimate

$$C \sum_{k \in \mathbb{Z}} (1 + |t - k/a_0|)^{-1-\epsilon} \leq 2C \sum_{k \in \mathbb{N}_0} (1 + |k/a_0|)^{-1-\epsilon} < \infty. \quad (4.33)$$

We now use $a_m = a_0/(Cv(m)) = a_0/C_m$ to arrive at

$$\begin{aligned} &\sum_{k \in \mathbb{Z}} (1 + |F^{-1}(F(t + k/a_m) - m)|)^{-1-\epsilon} \\ &= \sum_{k \in \mathbb{Z}} (1 + |F^{-1}(F(F^{-1}(m) + C_m(\tilde{t} + k/a_0)) - m)|)^{-1-\epsilon}, \end{aligned} \quad (4.34)$$

where $t = F^{-1}(m) + C_m \tilde{t}$. Now let $m \geq 0$ and wlog assume $\tilde{t} \in]0, a_0^{-1}]$ allowing us to divide the sum into terms with $k < 0$ and $k \geq 0$. For the latter we apply

(4.27) with the setting $y = F^{-1}(m)$ and $x = \tilde{t} + k/a_0$

$$\begin{aligned}
& \sum_{k \in \mathbb{N}_0} (1 + |F^{-1}(F(\tilde{t} + k/a_0 + F^{-1}(0)))|)^{-1-\epsilon} \\
&= \sum_{k \in \mathbb{N}_0} (1 + |\tilde{t} + k/a_0 + F^{-1}(0)|)^{-1-\epsilon} \\
&\leq \sum_{k \in \mathbb{N}_0} (1 + |k/a_0|)^{-1-\epsilon}.
\end{aligned} \tag{4.35}$$

However, since F' is nonincreasing away from 0, we obtain

$$|F(F^{-1}(m) + C_m(\tilde{t} + k/a_0)) - m| \leq |F(F^{-1}(m) + C_m(\tilde{t} - k/a_0)) - m|, \tag{4.36}$$

for all $k \in \mathbb{N}$ and consequently

$$\begin{aligned}
& \sum_{k \in \mathbb{N}} (1 + |F^{-1}(F(F^{-1}(m) + C_m(\tilde{t} - k/a_0)) - m))|)^{-1-\epsilon} \\
&\leq 2 \sum_{k \in \mathbb{N}_0} (1 + k/a_0)^{-1-\epsilon}.
\end{aligned} \tag{4.37}$$

This leads to the final estimate

$$\sum_{k \in \mathbb{Z}} \mathbf{T}_m |\theta(F(t - k/a_m))| \leq 2C \sum_{k \in \mathbb{N}_0} (1 + k/a_0)^{-1-\epsilon} < \infty, \tag{4.38}$$

for all $t \in D$ and $m \geq 0$. The symmetry of F can be used to arrive to an analogous estimate for $m < 0$. To conclude the proof of the Bessel property, we note that $F^{-1} \geq cx$ for some $c > 0$ and therefore

$$\operatorname{ess\,sup}_{t \in D} \left[\sum_{m \in \mathbb{Z}} \mathbf{T}_m |\theta(F(t))| \sum_{k \in \mathbb{Z}} \mathbf{T}_m |\theta(F(t - k/a_m))| \right] \tag{4.39}$$

$$\leq 2C \sum_{k \in \mathbb{N}_0} (1 + k/a_0)^{-1-\epsilon} \operatorname{ess\,sup}_{t \in D} \sum_{m \in \mathbb{Z}} \mathbf{T}_m |\theta(F(t))| \tag{4.40}$$

$$\leq 4C \sum_{k \in \mathbb{N}_0} (1 + k/a_0)^{-1-\epsilon} \sum_{m \in \mathbb{Z}} (1 + c|m|)^{-1-\epsilon} < \infty. \tag{4.41}$$

The proof for $D = \mathbb{R}^+$ is almost analogous, however some of the sums can actually be simplified to be only over the natural numbers. That makes the estimate in Equation (4.36) superfluous, the rest follows along the same lines. $\square$

4.3 Construction of tight warped frames

This section is devoted to the construction of tight frames following a given frequency scale. Using the framework of warped filterbanks this is easily realized using the painless case result presented in Corollary 4.11, this is one of the major assets of the presented construction. Tight transforms are desirable, as they allow for analysis and synthesis with the same system. Therefore it is not necessary to compute and store a dual system. Furthermore, the primal system

can be constructed to have good localization properties. For tight systems the dual system automatically shares the same properties as opposed to non-tight systems where the application of the inverse frame operator, to compute the canonical dual system, may break the localization.

In the following we use the painless case result given in Corollary 4.11 to show that we can easily construct tight frames if the underlying system $\{T_m\theta\}_{m \in \mathbb{Z}}$ is a partition of unity.

Definition 4.15. We call a system of translates $\{T_m\theta\}_{m \in \mathbb{Z}}$ a *partition of unity* if it satisfies for some constant C

$$\sum_{m \in \mathbb{Z}} |T_m\theta(t)|^2 = C \quad (4.42)$$

This definition of a partition of unity is slightly different from the standard in literature, as it is usually defined without the square and for $C = 1$. However, from a standard partition of unity one can easily construct one in the sense of the above definition by just taking square roots. For our purposes this definition will prove to be useful, therefore it makes sense to slightly deviate from the established concept.

Corollary 4.16. Let $\mathcal{G}(\theta, F, \mathbf{a})$ be a warped filterbank.

- If $\mathcal{G}(\mathbf{g}, \mathbf{a})$ satisfies the conditions of Corollary 4.11, then $\mathcal{G}(\theta, F, \mathbf{a})$ is a tight frame if and only if $\{\mathbf{T}_m\theta\}_{m \in \mathbb{Z}}$ is a partition of unity.
- If $\mathcal{G}(\theta, F, \mathbf{a})$ forms a tight frame, then $\{\mathbf{T}_m\theta\}_{m \in \mathbb{Z}}$ is a partition of unity.

Proof. The first part is an immediate consequence from Corollary 4.11, the second part follows from Proposition 4.10. $\square$

We see that even if the condition $a_m^{-1} \geq F^{-1}(d+m) - F^{-1}(c+m)$ is violated, $\{\mathbf{T}_m\theta\}_{m \in \mathbb{Z}}$ forming a partition of unity is a necessary condition for tightness. Therefore, partitions of unity are the optimal starting point when aiming to construct well-conditioned warped frames.

In the following we present a method to construct partitions of unity for a wide range of window functions encompassing the well known Hanning, Blackman-Harris and Nuttall windows. The following theorem has been reported in [89].

Theorem 4.17. Let $K \in \mathbb{N}$ and $a_k \in \mathbb{R}$ for $k \in \{0, 1, \dots, K\}$, and define

$$q(t) := \sum_{k=0}^K a_k \cos\left(2\pi k\left(t - \frac{1}{2}\right)\right) 1_{\{0 \leq t < 1\}}. \quad (4.43)$$

Then for any $R > 2K$

$$\sum_{m \in \mathbb{Z}} \left| q\left(t - \frac{m}{R}\right) \right|^2 = Ra_0^2 + \frac{R}{2} \sum_{k=1}^K a_k^2, \quad \forall t \in \mathbb{R};$$

i.e. the squares of a system of regular translates sums up to a constant function.

Proof. Let $t \in \mathbb{R}$ be arbitrary. Then

$$\begin{aligned}
\sum_{m \in \mathbb{Z}} \left| q\left(t - \frac{m}{R}\right) \right|^2 &= \sum_{m \in \mathbb{Z}} \left| \sum_{k=0}^K a_k \cos\left(2\pi k\left(t - \frac{m}{R} - \frac{1}{2}\right)\right) 1_{\{\frac{m}{R} \leq t < 1 + \frac{m}{R}\}} \right|^2 \\
&= \sum_{m=\lfloor Rt-(R-1) \rfloor}^{\lfloor Rt \rfloor} \left| \sum_{k=0}^K a_k \cos\left(2\pi k\left(t - \frac{m}{R} - \frac{1}{2}\right)\right) \right|^2 \\
&= \sum_{k=0}^K \sum_{j=0}^K a_k a_j \sum_{m=\lfloor Rt-(R-1) \rfloor}^{\lfloor Rt \rfloor} \cos\left(2\pi k\left(t - \frac{m}{R} - \frac{1}{2}\right)\right) \cos\left(2\pi j\left(t - \frac{m}{R} - \frac{1}{2}\right)\right).
\end{aligned} \tag{4.44}$$

Defining $z := 2\pi k\left(t - \frac{1}{2}\right)$ and $y := 2\pi j\left(t - \frac{1}{2}\right)$, and the inner terms of (4.44) as

$$A_{k,j} := \sum_{m=\lfloor Rt-(R-1) \rfloor}^{\lfloor Rt \rfloor} \cos\left(z - \frac{2\pi km}{R}\right) \cos\left(y - \frac{2\pi jm}{R}\right), \quad 0 \leq k, j \leq K,$$

and expanding the cosines into complex exponentials, we have

$$\begin{aligned}
A_{k,j} &= \frac{1}{4} \sum_{m=\lfloor Rt-(R-1) \rfloor}^{\lfloor Rt \rfloor} \left\{ \begin{aligned} &\left[\exp\left(iz + 2\pi i \frac{km}{R}\right) + \exp\left(-iz - 2\pi i \frac{km}{R}\right) \right] \\ &\cdot \left[\exp\left(iy + 2\pi i \frac{jm}{R}\right) + \exp\left(-iy - 2\pi i \frac{jm}{R}\right) \right] \end{aligned} \right\} \\
&= \frac{1}{4} \sum_{m=0}^{R-1} \left\{ \begin{aligned} &\left[\exp\left(iz + 2\pi i \frac{km}{R}\right) + \exp\left(-iz - 2\pi i \frac{km}{R}\right) \right] \\ &\cdot \left[\exp\left(iy + 2\pi i \frac{jm}{R}\right) + \exp\left(-iy - 2\pi i \frac{jm}{R}\right) \right] \end{aligned} \right\} \\
&= \frac{1}{4} \sum_{m=0}^{R-1} \exp\left(iz + 2\pi i \frac{km}{R}\right) \exp\left(iy + 2\pi i \frac{jm}{R}\right)
\end{aligned} \tag{4.45}$$

$$+ \frac{1}{4} \sum_{m=0}^{R-1} \exp\left(-iz - 2\pi i \frac{km}{R}\right) \exp\left(iy + 2\pi i \frac{jm}{R}\right) \tag{4.46}$$

$$+ \frac{1}{4} \sum_{m=0}^{R-1} \exp\left(iz + 2\pi i \frac{km}{R}\right) \exp\left(-iy - 2\pi i \frac{jm}{R}\right) \tag{4.47}$$

$$+ \frac{1}{4} \sum_{m=0}^{R-1} \exp\left(-iz - 2\pi i \frac{km}{R}\right) \exp\left(-iy + 2\pi i \frac{jm}{R}\right). \tag{4.48}$$

We now use the fact that if $\xi \neq 1$ is any R^{th} root of unity, then

$$\sum_{m=0}^{R-1} \xi^m = 0.$$

Since $K < R/2$, $k + j < R$ for all $0 \leq k, j \leq K$, and therefore

$$(4.45) = \begin{cases} \frac{R}{4} \exp(iz + iy), & \text{if } k = j = 0 \\ 0, & \text{otherwise} \end{cases},$$

$$(4.46) = (4.47) = \begin{cases} \frac{R}{4}, & \text{if } k = j \\ 0, & \text{otherwise} \end{cases}, \quad (4.49)$$

$$\text{and } (4.48) = \begin{cases} \frac{R}{4} \exp(-iz - iy), & \text{if } k = j = 0 \\ 0, & \text{otherwise} \end{cases}.$$

Since for $k = j = 0$ we have $y = z = 0$, we see that

$$A_{k,j} = \begin{cases} R, & \text{if } k = j = 0 \\ \frac{R}{2}, & \text{if } k = j \neq 0 \\ 0, & \text{otherwise} \end{cases}. \quad (4.50)$$

Finally, substituting (4.50) back into (4.44) yields

$$\sum_{m \in \mathbb{Z}} \left| q\left(t - \frac{m}{R}\right) \right|^2 = \sum_{k=0}^K \sum_{j=0}^K a_k a_j A_{k,j} = R a_0^2 + \frac{R}{2} \sum_{k=1}^K a_k^2.$$

□

Remark 4.18. For any fixed number K of cosine terms in Equation 4.43, we can choose a nonzero coefficient sequence $\{a_k\}_{k=0,\dots,K}$, such that the resulting window $q \in \mathcal{C}^{2K-1}$. This can be seen by differentiating the function to find a linear system of equations for the coefficients [59, 77]. To satisfy this system of equations, the coefficient sequence must be in the kernel of a $K \times (K+1)$ matrix. Since this kernel is never trivial, it is always possible to find a sequence of non-zero coefficients yielding the desired smoothness.

From Theorem 4.17 we can easily construct a system of integer translates forming a partition of unity.

Corollary 4.19. *Let $K \in \mathbb{N}$ and q be defined as in Equation 4.43 for some coefficient sequence $\{a_k\}_{k=0,\dots,K}$. Define for some $R > 2K$*

$$\theta(t) = q(t/R + 1/2). \quad (4.51)$$

Then $\{T_m \theta\}_{m \in \mathbb{Z}}$ forms a partition of unity. Furthermore, $\text{supp}(\theta) = [-R/2, R/2]$.

Proof. The substitution in Equation (4.51) easily shows that

$$\sum_{m \in \mathbb{Z}} |\theta(t - m)|^2 = \sum_{m \in \mathbb{Z}} \left| q\left(\frac{t}{R} - \frac{m}{R} + \frac{1}{2}\right) \right|^2 = R a_0^2 + \frac{R}{2} \sum_{k=1}^K a_k^2. \quad (4.52)$$

□

This construction of partitions of unity in connection with Corollary 4.16 allows us to construct tight frames if we choose the downsampling factors to be large enough, such that they satisfy

$$a_m^{-1} \geq F^{-1}(-R/2 + m) - F^{-1}(R/2 + m). \quad (4.53)$$

4.4 A continuous warped transform

Up to now we have always considered standard filter banks yielding a discrete set of transform coefficients. We now take the idea one step further by defining a continuous version of such systems. The corresponding transformation is a generalization of the continuous wavelet and STFT transforms. The general setting of warping functions stays the same as presented in Definition 4.4.

For a warping function F , we recall the corresponding weight function given by

$$w(t) = (F^{-1})'(t), \quad (4.54)$$

which has to be v -moderate for some submultiplicative weight function v .

Definition 4.20. Let F be a warping function and $\theta \in L_w^2(\mathbb{R})$. Then we define the *continuous warped filter bank* based on θ and F as

$$g_{x,\xi} = T_\xi \check{g}_x, \quad (4.55)$$

where $x \in D$ and $\xi \in \mathbb{R}$ and

$$g_x = \sqrt{F'(x)} \theta(F(t) - F(x)). \quad (4.56)$$

The *phase space* associated to this family is given by $D \times \mathbb{R}$.

From now on we assume that θ and F always satisfy the assumptions of the definition above.

Proposition 4.21. Let $g_{x,\xi}$ be a continuous warped filter bank based on the warping function F and $\theta \in L_w^2(\mathbb{R})$. Then $g_{x,\xi} \in L^2(D)$. If w is not only moderate but submultiplicative even the following uniform bound on the norms holds true

$$\|g_{x,\xi}\|_{L^2(D)} \leq \|\theta\|_{L_w^2(\mathbb{R})}. \quad (4.57)$$

Proof. Using the substitution $u = F(t) - F(x)$ we find

$$\begin{aligned} \|g_{x,\xi}\|_{L^2(\mathbb{R})} &= \|M_{-\xi} g_x\|_{L^2(D)} = \int_D F'(x) |\theta(F(t) - F(x))|^2 dt \\ &= F'(x) \int_{\mathbb{R}} |\theta(u)|^2 w(u + F(x)). \end{aligned} \quad (4.58)$$

Using the fact that w is v -moderate we find

$$\|g_{x,\xi}\|_{L^2(\mathbb{R})} \leq C v(F(x)) F'(x) \|\theta\|_{L_w^2}. \quad (4.59)$$

If the weight function is even submultiplicative, $v(F(x)) = 1/F'(x)$ and $C = 1$. $\square$

Definition 4.22. Given a continuous warped filter bank based on θ and F . Then we define the continuous transform of a function $f \in \mathcal{F}^{-1}(L^2(D))$

$$V_{\theta,F} f : D \times \mathbb{R} \rightarrow \mathbb{C}, \quad (x, \xi) \mapsto \langle f, g_{x,\xi} \rangle. \quad (4.60)$$

Theorem 4.23. *Let F be a warping function and $\theta_1, \theta_2 \in L_w^2$. Furthermore, assume that θ_1 and θ_2 are given, such that the following admissibility condition is fulfilled*

$$\int_D F'(t) \theta_1(F(t)) \overline{\theta_2(F(t))} dt = K < \infty. \quad (4.61)$$

Then the two continuous warped filter banks satisfy the orthogonality relations

$$\int_D \int_{\mathbb{R}} V_{\theta_1, F} f_1(x, \xi) \overline{V_{\theta_2, F} f_2(x, \xi)} d\xi dx = K \langle f_1, f_2 \rangle, \quad (4.62)$$

for some $f_1, f_2 \in \mathcal{F}^{-1}(L^2(D))$.

Proof. The associated continuous warped filter banks to θ_1 and θ_2 will be denoted by $g_{x, \xi}^1$ and $g_{x, \xi}^2$, respectively. We use the fact that $V_{\theta_i, F} f(x, \xi) = \mathcal{F}(\hat{f} \cdot g_x^i)(\xi)$ for $i = 1, 2$ to calculate

$$\begin{aligned} & \int_D \int_{\mathbb{R}} V_{\theta_1, F} f_1(x, \xi) \overline{V_{\theta_2, F} f_2(x, \xi)} d\xi dx \\ &= \int_D \int_{\mathbb{R}} \mathcal{F}(\hat{f}_1 \cdot g_x^1)(\xi) \overline{\mathcal{F}(\hat{f}_2 \cdot g_x^2)(\xi)} d\xi dx \\ &= \int_D \hat{f}_1(t) \overline{\hat{f}_2(t)} \int_D g_x^1(t) \overline{g_x^2(t)} dx dt \end{aligned} \quad (4.63)$$

Using the substitution $F(u) = F(t) - F(x)$ we can simplify the inner integral

$$\begin{aligned} \int_D g_x^1(t) \overline{g_x^2(t)} dx &= \int_D F'(x) \theta_1(F(t) - F(x)) \overline{\theta_2(F(t) - F(x))} dx \\ &= \int_D F'(u) \theta_1(F(u)) \overline{\theta_2(F(u))} du = K. \end{aligned} \quad (4.64)$$

Using Parseval's formula we can conclude the assertion. $\square$

This theorem shows that general continuous warped filter banks satisfy orthogonality relations in the same sense as the STFT and wavelets do. This is an important property, as it immediately yields an inverse of the continuous filter bank transform.

Corollary 4.24. *Given a warping function F and some $\theta \in L_w^2$ satisfying*

$$\int_D F'(t) |\theta(F(t))|^2 dt = K < \infty. \quad (4.65)$$

Then any $f \in \mathcal{F}^{-1}(L^2(D))$ can be reconstructed from its transform coefficient by

$$f = \frac{1}{K} \int_D \int_{\mathbb{R}} V_{\theta, F} f(x, \xi) g_{x, \xi} d\xi dx. \quad (4.66)$$

The equation holds in the weak sense.

Proof. The assertion follows easily from the orthogonality relations by setting $\theta = \theta_1 = \theta_2$ since for any given $f_2 \in \mathcal{F}^{-1}(L^2(D))$ we have the relation

$$\langle f, f_2 \rangle = \frac{1}{K} \int_D \int_{\mathbb{R}} \langle f, g_{x, \xi} \rangle \langle g_{x, \xi}, f_2 \rangle d\xi dx. \quad (4.67)$$

$\square$

The orthogonality relations already describe the operator $V_{\theta,F}$ well. In the following we consider the continuity of the function $V_{\theta,F}f : D \times \mathbb{R} \rightarrow \mathbb{C}$ for some fixed $f \in \mathcal{F}^{-1}(L^2(D))$.

Proposition 4.25. *Let F and θ be given, furthermore let $f \in \mathcal{F}^{-1}(L^2(D))$. Then the function $V_{\theta,F}f$ is continuous as a function from $D \times \mathbb{R}$ into $\mathbb{C}$.*

Proof. We further compute the following estimate

$$\begin{aligned} |V_{\theta,F}f(x, \xi) - V_{\theta,F}f(\tilde{x}, \tilde{\xi})| &= |\langle \hat{f}, M_\xi g_x - M_{\tilde{\xi}} g_{\tilde{x}} \rangle| \\ &\leq \|\hat{f}\|_{L^2(D)} \left(\|M_\xi g_x - M_{\tilde{\xi}} g_{\tilde{x}}\| + \|M_{\tilde{\xi}}(g_x - g_{\tilde{x}})\| \right). \end{aligned} \quad (4.68)$$

Since modulations are continuous on $L^2(D)$ it suffices to show that $\|g_{x,0} - g_{\tilde{x},0}\| \rightarrow 0$, as $\tilde{x}$ tends to x . To see this we calculate

$$\begin{aligned} \|g_{x,0} - g_{\tilde{x},0}\|^2 &= \int_D \left| \sqrt{F'(x)}(T_{F(x)}\theta)(F(t)) - \sqrt{F'(\tilde{x})}(T_{F(\tilde{x})}\theta)(F(t)) \right|^2 dt \\ &= F'(x) \|T_{F(x)}\theta - \sqrt{F'(\tilde{x})/F'(x)} T_{F(\tilde{x})}\theta\|_{L_w^2}^2 \\ &= F'(x) \|T_{F(x)}\theta - T_{F(\tilde{x})}\theta + T_{F(\tilde{x})}\theta - \sqrt{F'(\tilde{x})/F'(x)} T_{F(\tilde{x})}\theta\|_{L_w^2}^2. \end{aligned} \quad (4.69)$$

Now an $\varepsilon/2$ argument finishes the proof since $\sqrt{F'(\tilde{x})/F'(x)} \rightarrow 1$, $F(\tilde{x}) \rightarrow F(x)$ as $\tilde{x} \rightarrow x$ and translations are continuous on the weighted space L_w^2 because of the moderateness of the weight function. $\square$

4.5 Discretization and numerical examples

The purpose of this section is the description of warped time-frequency systems for $\mathbb{C}^L$ by sampling a warped system of translates $\theta_{F,m}(t) = \theta(F(t) - m)$. The domain of the warping function F is always either $\mathbb{R}$ or $\mathbb{R}^+$, for both of the cases different issues arise in the discretization procedure. The painless case results carry over verbatim from the infinite dimensional setting presented in Section 4.3. For simplicity we assume $\text{supp}(\theta) \subseteq [-R/2, R/2]$, since fast algorithms only exist for the painless case [7]. In addition to the signal length we always assume that we know the sampling rate, therefore at a sampling rate of ξ_s Hz the samples on the Fourier side are equally spaced in $[-\xi_s/2, \xi_s/2]$.

Case 1 $D = \mathbb{R}$. We assume that F is point symmetric around 0. Let $m_{\max} = \max\{m \in \mathbb{Z} : F^{-1}(R/2 + m) \leq \xi_s/2\}$, the index of the last translate, that has its support completely inside the frequency range spanned by the discrete frequencies. Therefore for $-m_{\max} \leq m \leq m_{\max}$, the functions $\text{supp}(\theta_{F,m}) \subseteq [-\xi_s/2, \xi_s/2]$. That allows to define

$$g_m(l) = N_m^{-1/2} \theta(F(l/\xi_s) - m), \quad (4.70)$$

where $N_m = L/a_m$ and $a_m = L/|\text{supp } g_m|$. In order to span the whole frequency axis with filters we need to define a padding function around the Nyquist frequency $\xi_s/2$.

$$g_{m_{\max}+1}(l) := N_{m_{\max}+1}^{-1/2} \left(\sum_{m \in \mathbb{Z} \setminus [m_{\min}, m_{\max}]} |g_m^0(l/\xi_s)|^2 \right)^{1/2}, \quad (4.71)$$

where $N_{m_{\max}+1}$ and $a_{m_{\max}+1}$ are defined analogously to the other filters. This construction is tailored to yield a painless filterbank.

Case 2 $D = \mathbb{R}^+$. Here we obtain a covering of the whole discrete frequency range by mirroring the filters around the 0 frequency, which requires distinguishing more cases than before. We define $m_{\max}$ as above, in addition to this maximal index we fix a minimal index $m_{\min} < m_{\max}$. This minimal index is chosen by the user and is not inherent to the construction, as it determines how many scales one wants to consider, it can also be chosen to be negative. This leads to the definition of two sets of filters

$$g_m^r(l) = \begin{cases} N_m^{-1/2} \theta(F(l/\xi_s) - m) & l > 0, \\ 0 & \text{else,} \end{cases} \quad (4.72)$$

$$g_m^l(l) = \begin{cases} N_m^{-1/2} \theta(F(-l/\xi_s) - m) & l < 0, \\ 0 & \text{else,} \end{cases} \quad (4.73)$$

for all $m = m_{\min}, \dots, m_{\max}$ and $N_m = L/a_m$ and $a_m = L/|\text{supp } g_m|$. In addition to these filters we need to design padding functions around both, the 0 frequency and the Nyquist frequency. We define them respectively as follows

$$g^0(l) := \begin{cases} N_n^{-1/2} \left(\sum_{m=-\infty}^{m_{\min}} |\theta(F(|l|/\xi_s) - m)|^2 \right)^{1/2} & l \neq 0, \\ N_n^{-1/2} & l = 0, \end{cases}, \text{ and} \quad (4.74)$$

$$g^{\text{Ny}}(l) := N_{\text{Ny}}^{-1/2} \left(\sum_{m=m_{\max}+1}^{\infty} |g_m^0(|l|/\xi_s)|^2 \right)^{1/2} \quad (4.75)$$

Again, we define the normalization factors to be the length of the support. Also these filters are designed to fall into the painless case.

Remark 4.26. When only functions in $\mathbb{R}^L$ are of interest, e.g. in audio signal processing, the hermite symmetry of the Fourier transform of real-valued signals can be used to avoid the mirroring procedure described above. Instead of obtaining a symmetric filterbank, covering the whole frequency range, we obtain a filterbank covering the positive frequency range only, which is sufficient for handling real-valued signals. Due to the point symmetry of the warping function F in the case $D = \mathbb{R}$ we also obtain a symmetric filter bank if we assume that θ is symmetric around 0 and can therefore also use the Hermite symmetry.

The definition of padding functions in both cases are chosen for a good reason. Essentially, we want to keep the good properties of the original warped system intact in their discretization. But we also need to consider the finite, circular structure of time-frequency systems over $\mathbb{C}^L$. In the upper frequency range, this means that it is reasonable to only consider those filters, which are completely inside the frequency range $[-\xi_s/2, \xi_s/2]$. An illustration of ERB-let filter transfer functions and wavelet filter transfer functions constructed with the described mirroring procedure can be found in Figure 4.2.

In Case 2 ($D = \mathbb{R}$), an infinite number of filters is placed in any neighborhood of the zero frequency. At some point, the sampling rate ξ_s is not high enough to distinguish the different center frequencies anymore, hence they should be considered the same. Furthermore, the filters that are not sampled densely enough do not bear any resemblance to their continuous model and show completely

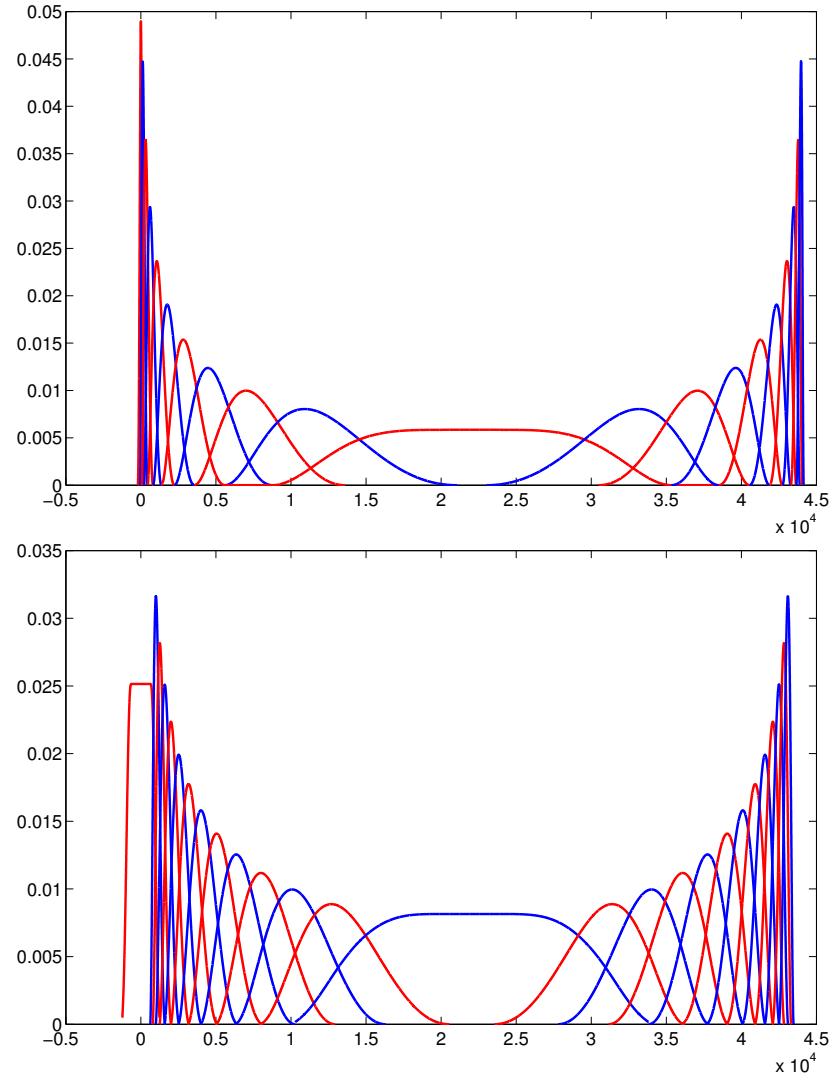


Figure 4.2: Top: ERB filter transfer functions with padding function around the Nyquist frequency, Bottom: Wavelet filter transfer functions, padding functions are necessary around 0 and the Nyquist frequency

different time-frequency concentration. Therefore, we propose a single filter, covering all these critical cases and preserving the summation properties again. In the wavelet community such a function is often referred to as *scaling function*, as it incorporates the 0 frequency and therefore has non-zero integral as opposed to the wavelets.

In practice, the resulting high- and low-pass filters can often be chosen such that the information contained in the high-pass and low-pass bands is irrelevant to the user.

Remark 4.27. Since the discretized filters inherit the summability properties from the continuous construction we can use functions constructed in Theorem 4.17 to build tight frames. For all the examples that follow we use the Hanning window defined as

$$h(t) = \frac{1}{2}(1 + \cos(2\pi t)), \quad (4.76)$$

which falls into the class described by this theorem with the setting $a_0 = a_1 = 1/2$. This particular choice of window can therefore be used to construct tight frames with $R > 2$ filters overlapping. With the subsampling factors defined above this leads to integer redundancies $R > 2$. In Figure 4.3 we show the transform coefficients of 3 different warping functions. All of the underlying function systems are tight frames, the frequency scales are quite different.

Remark 4.28. The construction of the padding functions can be simplified if we start with a partition of unity for the construction. In this case we can determine the appearing sums over the rest of the filters exactly, as we know they sum up to a constant.

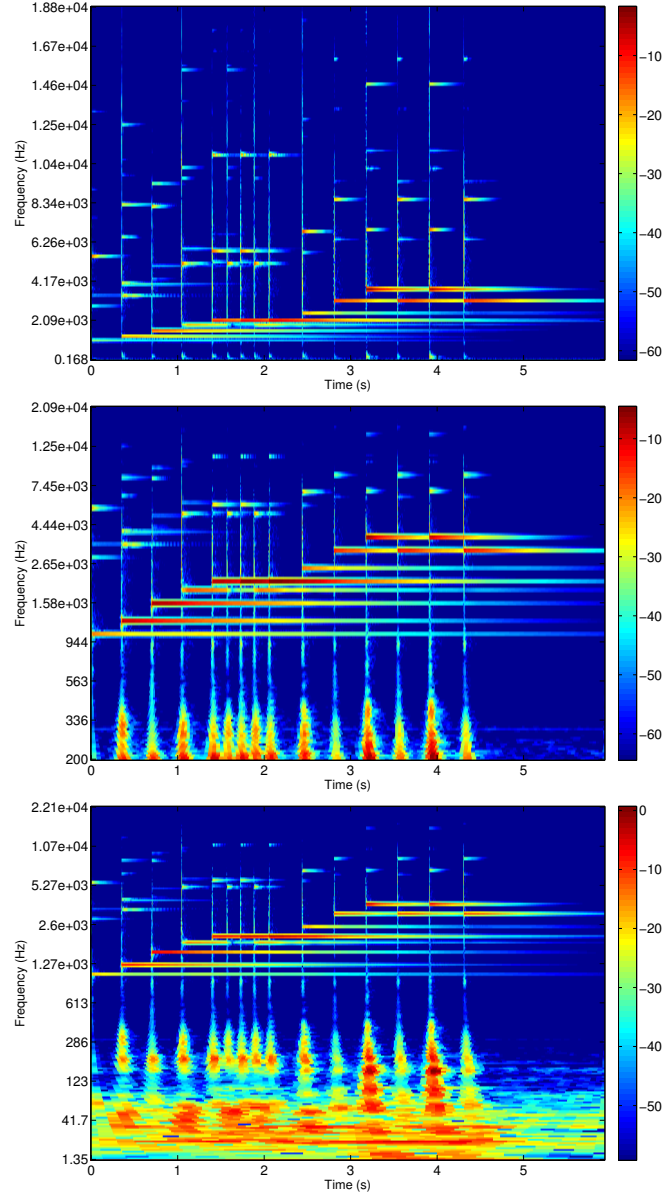


Figure 4.3: We show here three different transform coefficients of the Glockenspiel audio signal (Sampling rate 44100 Hz) for different warping functions with the same original window $\theta = h(t)$. From top to bottom: Gabor transform, wavelet transform, ERB-let transform. The minimal represented frequency for the wavelet transform is chosen to be 200. Since the input signal is real it suffices to show the positive frequencies only.

Bibliography

- [1] LTFAT - The Large Time-Frequency Analysis Toolbox. <http://ltfat.sourceforge.net/>.
- [2] J. Allen and L. Rabiner. A unified approach to short-time Fourier analysis and synthesis. *Proceedings of the IEEE*, 65(11):1558–1564, 1977.
- [3] W. Alltop. Complex sequences with low periodic correlations. *IEEE Trans. Info. Th.*, 26(3):350–354, May 1980.
- [4] E. Anderson, Z. Bai, C. Bischof, J. Demmel, J. Dongarra, A. Greenbaum, S. Hammarling, and D. Sorensen. *LAPACK Users' Guide*. Society for Industrial and Applied Mathematics, Philadelphia, PA, Third edition, 1999.
- [5] L. Auslander, I. Gertner, and R. Tolimieri. The discrete Zak transform application to time-frequency analysis and synthesis of nonstationary signals. *IEEE Trans. Signal Process.*, 39(4):825–835, 1991.
- [6] P. Balazs, D. Bayer, and A. Rahimi. Multipliers for continuous frames in Hilbert spaces. *J. Phys. A*, Special issue: Coherent states(45):244023, 2012.
- [7] P. Balazs, M. Dörfler, F. Jaillet, N. Holighaus, and G. A. Velasco. Theory, implementation and applications of nonstationary Gabor frames. *J. Comput. Appl. Math.*, 236(6):1481–1496, 2011.
- [8] R. Baraniuk, P. Flandrin, A. J. Janssen, and O. Michel. Measuring time frequency information content using the renyi entropies. *IEEE Trans. Inf. Theory*, 47(4):1391–1409, 2001.
- [9] M. J. Bastiaans and M. C. W. Geilen. On the discrete Gabor transform and the discrete Zak transform. *Signal Process.*, 49(3):151–166, 1996.
- [10] H. Bölcskei, F. Hlawatsch, and H. G. Feichtinger. Frame-theoretic analysis of oversampled filter banks. *IEEE Trans. Signal Process.*, 46(12):3256–3268, 1998.
- [11] A. Bultheel and H. Martinez. A shattered survey of the Fractional Fourier Transform. 2003.
- [12] G. Calugareanu. The total number of subgroups of a finite Abelian group. *Scientiae Mathematicae Japonicae*, 60(1):157–168, 2004.

- [13] C. Candan, M. Kutay, and H. Ozaktas. The discrete fractional Fourier transform. *IEEE Trans. Signal Process.*, 48(5):1329–1337, 2002.
- [14] P. G. Casazza and M. Fickus. Fourier transforms of finite chirps. *EURASIP J. Adv. Signal Process.*, 2006:1–7, 2006.
- [15] O. Christensen. *An Introduction to Frames and Riesz Bases*. Applied and Numerical Harmonic Analysis. Birkhäuser, Boston, 2003.
- [16] O. Christensen, H. G. Feichtinger, and S. Paukner. *Gabor Analysis for Imaging*, volume 3, pages 1271–1307. Springer Berlin, 2011.
- [17] O. Christensen and S. Goh. From dual pairs of Gabor frames to dual pairs of wavelet frames and vice versa. *Appl. Comput. Harmon. Anal.*, 2013.
- [18] J. Conway and N. J. A. Sloane. *Sphere Packings, Lattices and Groups*. Volume 290. Third Edition edition, 1999.
- [19] J. Cooley and J. Tukey. An algorithm for the machine calculation of complex Fourier series. *Math. Comp.*, 19:297–301, 1965.
- [20] E. Cordero and K. Gröchenig. Time-frequency analysis of localization operators. *J. Funct. Anal.*, 205(1):107–131, 2003.
- [21] J. P. Costas. A study of a class of detection waveforms having nearly ideal range-doppler ambiguity properties. *Proc. of the IEEE*, 72:996–1009, 1984.
- [22] M. G. Cowling and J. F. Price. Bandwidth versus time concentration: the Heisenberg-Pauli-Weyl inequality. *SIAM J. Math. Anal.*, 15(1):151–165, 1984.
- [23] Z. Cvetkovic and M. Vetterli. Oversampled filter banks. *IEEE Trans. Signal Process.*, 46(5):1245–1255, May 1998.
- [24] I. Daubechies. Time-frequency localization operators: a geometric phase space approach. *IEEE Trans. Inform. Theory*, 34(4):605–612, July 1988.
- [25] I. Daubechies. The wavelet transform, time-frequency localization and signal analysis. *IEEE Trans. Inform. Theory*, 36(5):961–1005, 1990.
- [26] I. Daubechies. *Ten lectures on wavelets*, volume 61 of *CBMS-NSF Regional Conference Series in Applied Mathematics*. SIAM, Philadelphia, PA, 1992.
- [27] I. Daubechies, A. Grossmann, and Y. Meyer. Painless nonorthogonal expansions. *J. Math. Phys.*, 27(5):1271–1283, May 1986.
- [28] M. A. de Gosson and D. Onchis. Multivariate symplectic Gabor frames with Gaussian windows. *preprint*, 2012.
- [29] F. de Mesmay, Y. Voronenko, and M. Püschel. Offline library adaptation using automatically generated heuristics. In *International Parallel and Distributed Processing Symposium (IPDPS)*, 2010.
- [30] M. Dörfler and L. D. Abreu. An inverse problem for localization operators. *Inverse Problems*, 28, 2012.

- [31] M. Dörfler, H. G. Feichtinger, and K. Gröchenig. Compactness criteria in function spaces. *Colloq. Math.*, 94(1):37–50, 2002.
- [32] M. Dörfler and E. Matusiak. Nonstationary Gabor Frames - Approximately Dual Frames and Reconstruction Errors. *ArXiv e-prints*, (arXiv:1301.1802), 2012.
- [33] M. Dörfler and E. Matusiak. Nonstationary Gabor Frames - Existence and Construction. *ArXiv e-prints*, (arXiv:1112.5262), 2012.
- [34] M. Dörfler and B. Torrésani. Representation of operators in the time-frequency domain and generalized Gabor multipliers. *J. Fourier Anal. Appl.*, 16(2):261–293, 2010.
- [35] H. Feichtinger, D. Onchis, B. Ricaud, B. Torrésani, and C. Wiesmeyr. A method for optimizing the ambiguity function concentration. In *Proceedings of the European Signal Processing Conference*, pages 804–808. IEEE, 2012.
- [36] H. G. Feichtinger. English translation of: Gewichtsfunktionen auf lokalkompakten Gruppen. *Sitzungsber.d.österr. Akad.Wiss.*, 188.
- [37] H. G. Feichtinger. On a new Segal algebra. *Monatsh. Math.*, 92:269–289, 1981.
- [38] H. G. Feichtinger, M. Hazewinkel, N. Kaiblinger, E. Matusiak, and M. Neuhauser. Metaplectic operators on C^n . *Quart. J. Math. Oxford Ser.*, 59(1):15–28, 2008.
- [39] H. G. Feichtinger and N. Kaiblinger. 2D-Gabor analysis based on 1D algorithms. In *Proc. OEAGM-97 (Hallstatt, Austria)*, 1997.
- [40] H. G. Feichtinger, W. Kozek, P. Prinz, and T. Strohmer. On multidimensional non-separable Gabor expansions. In *Proc. SPIE: Wavelet Applications in Signal and Image Processing IV*, August 1996.
- [41] H. G. Feichtinger, F. Luef, and T. Werther. A guided tour from linear algebra to the foundations of Gabor analysis. In *Gabor and Wavelet Frames*, volume 10 of *Lect. Notes Ser. Inst. Math. Sci. Natl. Univ. Singap.*, pages 1–49. World Sci. Publ., Hackensack, 2007.
- [42] H. G. Feichtinger and T. Strohmer. *Advances in Gabor Analysis*. Birkhäuser, Basel, 2003.
- [43] C. Fernández and A. Galbis. Some remarks on compact Weyl operators. *Integral Transforms Spec. Funct.*, 18(8):599–607, 2007.
- [44] P. Flandrin. Maximum signal energy concentration in a time-frequency domain. *ICASSP-88*, pages 2176–2179, 1988.
- [45] G. B. Folland. *Harmonic Analysis in Phase Space*. Princeton University Press, Princeton, N.J., 1989.
- [46] G. B. Folland and A. Sitaram. The uncertainty principle: A mathematical survey. *J. Fourier Anal. Appl.*, 3(3):207–238, 1997.

- [47] S. Foucart and H. Rauhut. *A mathematical introduction to compressive sensing*. Applied and Numerical Harmonic Analysis. Springer, 2013.
- [48] F. Friedlander and M. Joshi. *Introduction to the Theory of Distributions*. Cambridge Univ Pr, 1998.
- [49] M. Frigo and S. Johnson. The design and implementation of FFTW 3. *Proceedings of the IEEE*, 93(2):216–231, 2005.
- [50] D. Gabor. Theory of communication. *J. IEE*, 93(26):429–457, 1946.
- [51] B. Glasberg and B. Moore. Derivation of auditory filter shapes from notched-noise data. *Hearing Research*, 47:103–138, 1990.
- [52] M. Grady. A group theoretic approach to a famous partition formula. *The American Mathematical Monthly*, 112(7):645–651, 2005.
- [53] K. Gröchenig. Aspects of Gabor analysis on locally compact abelian groups. In H. G. Feichtinger and T. Strohmer, editors, *Gabor Analysis and Algorithms: Theory and Applications*, pages 211–231. Birkhäuser Boston, Boston, MA, 1998.
- [54] K. Gröchenig. *Foundations of Time-Frequency Analysis*. Appl. Numer. Harmon. Anal. Birkhäuser Boston, Boston, MA, 2001.
- [55] K. Gröchenig. Weight functions in time-frequency analysis. In L. Rodino and et al., editors, *Pseudodifferential Operators: Partial Differential Equations and Time-Frequency Analysis*, volume 52 of *Fields Inst. Commun.*, pages 343–366. Amer. Math. Soc., Providence, RI, 2007.
- [56] K. Gröchenig and M. Neuhauser. An uncertainty principle for the radar ambiguity function. *In preparation*.
- [57] M. Hampejs, N. Holighaus, L. Tóth, and C. Wiesmeyr. On the subgroups of the group $Z_m \times Z_n$. *ArXiv e-prints*, (arXiv:1211.1797), 2012.
- [58] G. H. Hardy. A theorem concerning Fourier transforms. *Journal L. M. S.*, 8:227–231, 1933.
- [59] F. J. Harris. On the use of windows for harmonic analysis with the discrete Fourier transform. *Proc. of the IEEE*, 66(1):51–83, 1978.
- [60] H. Helms. Fast Fourier transform method of computing difference equations and simulating filters. *IEEE Transactions on Audio and Electroacoustics*, 15(2):85–90, 1967.
- [61] M. A. Herman and T. Strohmer. High-resolution radar via compressed sensing. *IEEE Trans. Sig. Proc.*, 57(6):2275–2284, 2009.
- [62] E. Hernández, D. Labate, and G. Weiss. A unified characterization of reproducing systems generated by a finite family. II. *J. Geom. Anal.*, 12(4):615–662, 2002.
- [63] F. Hlawatsch and W. Krattenthaler. Bilinear signal synthesis. *IEEE Trans. on Sig. Proc.*, 40(2):352–363, 1992.

- [64] N. Holighaus. Structure of nonstationary Gabor frames and their dual systems. *ArXiv e-prints*, (arXiv:1306.5037), 2013.
- [65] N. Holighaus. *Theory and implementation of adaptive time-frequency transforms*. PhD thesis, University of Vienna, 2013.
- [66] A. J. E. M. Janssen. Positivity and spread of bilinear time-frequency distributions. In F. Hlawatsch and W. Mecklenbräuker, editors, *The Wigner Distribution. Theory and Applications in Signal Processing.*, pages 1–58. Elsevier Science, Amsterdam, 1997.
- [67] A. J. E. M. Janssen and P. L. Søndergaard. Iterative algorithms to approximate canonical Gabor windows: Computational aspects. *J. Fourier Anal. Appl.*, 13(2):211–241, 2007.
- [68] S. Johnson and M. Frigo. A Modified Split-Radix FFT With Fewer Arithmetic Operations. *IEEE Trans. Signal Process.*, 55(1):111, 2007.
- [69] N. Kaiblinger. *Metaplectic representation, eigenfunctions of phase space shifts, and Gelfand-Shilov spaces for LCA groups*. PhD thesis, Dept. Mathematics, Univ. Vienna, 1999.
- [70] N. Kaiblinger. Approximation of the Fourier transform and the dual Gabor window. *J. Fourier Anal. Appl.*, 11(1):25–42, 2005.
- [71] N. Kaiblinger and M. Neuhauser. Metaplectic operators for finite abelian groups and R^d . *Indag. Math.*, 20(2):233–246, 2009.
- [72] M. A. Kutay, H. M. Ozaktas, and Z. Zalevsky. *The Fractional Fourier Transform, with Applications in Optics and Signal Processing*. John Wiley and Sons, 2001.
- [73] E. H. Lieb. Integral bounds for radar ambiguity functions and Wigner distributions. *J. Math. Phys.*, 31(3):594–599, 1990.
- [74] B. Moore. *An introduction to the psychology of hearing*, volume 4. Academic press San Diego, 2003.
- [75] S. Nam. An Uncertainty Principle for Discrete Signals. *ArXiv e-prints*, (arXiv:1307.6321), 2013.
- [76] T. Necciari, P. Balazs, N. Holighaus, and P. Søndergaard. The ERBlet transform: An auditory-based time-frequency representation with perfect reconstruction. In *Proceedings of the 38th International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2013)*, pages 498–502, 2013.
- [77] A. Nuttall. Some windows with very good sidelobe behavior. *Acoustics, Speech and Signal Processing, IEEE Transactions on*, 29(1):84–91, 1981.
- [78] F. Olver. *NIST handbook of mathematical functions*. Cambridge University Press, 2010.
- [79] A. Paeth. A fast algorithm for general raster rotation. In *Graphics Interface*, volume 86, pages 77–81, 1986.

- [80] S. Paukner. Foundations of Gabor Analysis for Image Processing. Master's thesis, 2007.
- [81] J. Petrillo. Counting Subgroups in a Direct Product of Finite Cyclic Groups. *The College Mathematics Journal*, 42(3):215–222, 2011.
- [82] M. Portnoff. Implementation of the digital phase vocoder using the fast Fourier transform. *IEEE Trans. Acoust. Speech Signal Process.*, 24(3):243–248, 1976.
- [83] J. Ramanathan and T. Steger. Incompleteness of sparse coherent states. *Appl. Comput. Harmon. Anal.*, 2(2):148–153, 1995.
- [84] B. Ricaud, G. Stempfel, B. Torr sani, C. Wiesmeyer, H. Lachambre, and D. Onchis. An optimally concentrated Gabor transform for localized time-frequency components. *ArXiv e-prints*, (arXiv:1310.8573), 2013.
- [85] B. Ricaud and B. Torr sani. A survey of uncertainty principles and some signal processing applications. *ArXiv e-prints*, (arXiv:1211.5914), 2012.
- [86] A. Ron and Z. Shen. Generalized shift-invariant systems. *Constr. Approx.*, 22:1–45, 2005.
- [87] R. Schmidt. *Subgroup lattices of groups*, volume 14. Walter de Gruyter, 1994.
- [88] D. Shuman, S. Narang, P. Frossard, A. Ortega, and P. Vandergheynst. The emerging field of signal processing on graphs: Extending high-dimensional data analysis to networks and other irregular domains. *Signal Processing Magazine, IEEE*, 30(3):83–98, 2013.
- [89] D. Shuman, C. Wiesmeyer, N. Holighaus, and P. Vandergheynst. Spectrum-Adapted Tight Graph Wavelet and Vertex-Frequency Frames. *ArXiv e-prints*, (arXiv:1311.0897), 2013.
- [90] P. L. Sondergaard. Gabor frames by sampling and periodization. *Adv. Comput. Math.*, 27(4):355–373, 2007.
- [91] P. L. S ndergaard. Efficient Algorithms for the Discrete Gabor Transform with a long FIR window. *J. Fourier Anal. Appl.*, 18(3):456–470, 2012.
- [92] P. L. S ndergaard, B. Torr sani, and P. Balazs. The Linear Time Frequency Analysis Toolbox. *International Journal of Wavelets, Multiresolution Analysis and Information Processing*, 10(4), 2012.
- [93] T. Stockham Jr. High-speed convolution and correlation. In *Proceedings of the April 26-28, 1966, Spring joint computer conference*, pages 229–233. ACM, 1966.
- [94] T. Strohmer. Numerical algorithms for discrete Gabor expansions. In *Gabor Analysis and Algorithms*, chapter 8, pages 267–294. Birkh user, 1998.
- [95] T. Strohmer and S. Beaver. Optimal OFDM system design for time-frequency dispersive channels. *IEEE Trans. Comm.*, 51(7):1111–1122, July 2003.

- [96] S. M. Sussman. Least-squares synthesis of radar ambiguity functions. *IRE Transactions on Information Theory*, 8(3):246–254, 1962.
- [97] M. Suzuki. On the lattice of subgroups of finite groups. *Transactions of the American Mathematical Society*, 70(2):345–371, 1951.
- [98] M. Tarnauceanu. An arithmetic method of counting the subgroups of a finite abelian group. *Bull. Math. Soc. Sci. Math. Roumanie (NS)*, 53(101):373–386, 2010.
- [99] M. E. Taylor. *Measure Theory and Integration.*, volume 76 of *Graduate Studies in Mathematics*. American Mathematical Society (AMS), Providence, RI, 2006.
- [100] R. Tolimieri and R. Orr. Poisson summation, the ambiguity function, and the theory of Weyl–Heisenberg frames. *J. Fourier Anal. Appl.*, 1:233–247, 1995.
- [101] T. Tschurtschenthaler. The Gabor Frame Operator (its Structure and numerical Consequences). Master’s thesis, University of Vienna, 2000.
- [102] P. P. Vaidyanathan. *Multirate Systems and Filter Banks*. Prentice - Hall, 1993.
- [103] A. J. van Leest. *Non-separable Gabor schemes. Their Design and Implementation*. PhD thesis, Tech. Univ. Eindhoven, 2001.
- [104] A. J. van Leest and M. J. Bastiaans. Implementations of non-separable Gabor schemes. In *Proc. Eusipco 2004 ,12th European Signal Processing Conference, Vienna, Austria, 7-10 September 2004*, pages 1565–1568, 2004.
- [105] M. Vetterli. Filter banks allowing perfect reconstruction. *Signal Process.*, 10:219–244, 1986.
- [106] D. F. Walnut. Continuity properties of the Gabor frame operator. *J. Math. Anal. Appl.*, 165(2):479–504, 1992.
- [107] R. C. Whaley and A. Petitet. Minimizing development and maintenance costs in supporting persistently optimized BLAS. *Software: Practice and Experience*, 35(2):101–121, February 2005.
- [108] C. Wiesmeyr, N. Holighaus, and P. Sondergaard. Efficient algorithms for discrete Gabor transforms on a nonseparable lattice. *IEEE Trans. Signal Process.*, 61(20):5131 – 5142, 2013.
- [109] M. Wong. *Localization Operators*. Seoul National University, Seoul, 1999.
- [110] Y. Y. Zeevi and M. Zibulski. Oversampling in the Gabor scheme. *IEEE Trans. Signal Process.*, 41(8):2679–2687, 1993.
- [111] M. Zibulski and Y. Zeevi. Signal- and image-component separation by a multi-window Gabor-type scheme. In *Pattern Recognition, 1996., Proceedings of the 13th International Conference on., volume 2*, pages 835–839, Vienna , Austria, aug 1996. IEEE.

- [112] M. Zibulski and Y. Y. Zeevi. Analysis of multiwindow Gabor-type schemes by frame methods. *Appl. Comput. Harmon. Anal.*, 4(2):188–221, 1997.
- [113] M. Zibulski and Y. Y. Zeevi. Discrete multiwindow Gabor-type transforms. *IEEE Trans. Signal Process.*, 45(6):1428–1442, 1997.
- [114] M. Zibulski and Y. Y. Zeevi. The generalized Gabor scheme and its application in signal and image representation. In *Signal and Image Representation in Combined Spaces*, volume 7 of *Wavelet Anal. Appl.*, pages 121–164. Academic Press, San Diego, CA, 1998.
- [115] Y. Zou. Gaussian binomials and the number of sublattices. *Acta Crystallographica Section A: Foundations of Crystallography*, 62(5):409–410, 2006.

Curriculum Vitae

Personal Information

Name	Christoph Wiesmeyr
Date of Birth	October 12, 1984
Nationality	Austrian
E-mail	christoph.wiesmeyr@univie.ac.at

Current Position

- Duration: September 2010 - Present.
- Position: Research assistant.
- Organization: Numerical Harmonic Analysis Group (NuHAG), Faculty of Mathematics, University of Vienna.
- Research: Work topics include adaptive signal representations and uncertainty principles.

Academic Qualifications

- Dipl. Ing.(DI), in Industrial Mathematics from the University of Linz (2010).

Publications

Submitted

- M. Hampejs, N. Holighaus, L. Tóth, and C. Wiesmeyr. On the subgroups of the group $Z_m \times Z_n$. *ArXiv e-prints*, (arXiv:1211.1797), 2012.
- B. Ricaud, G. Stempfel, B. Torr sani, C. Wiesmeyr, H. Lachambre, and D. Onchis. An optimally concentrated Gabor transform for localized time-frequency components. *ArXiv e-prints*, (arXiv:1310.8573), 2013.
- D. Shuman, C. Wiesmeyr, N. Holighaus, and P. Vandergheynst. Spectrum-Adapted Tight Graph Wavelet and Vertex-Frequency Frames. *ArXiv e-prints*, (arXiv:1311.0897), 2013.

Published

- C. Wismeyr, N. Holighaus, and P. Sondergaard. Efficient algorithms for discrete Gabor transforms on a nonseparable lattice. *IEEE Trans. Signal Process.*, 61(20):5131 – 5142, 2013.
- H. Feichtinger, D. Onchis, and C. Wismeyr. Construction of approximate dual wavelet frames. *Advances in Computational Mathematics*, accepted, 2013.
- S. Bar lev, O. Boxma, W. Stadjé, F. Schouten, and C. Wismeyr. Two-stage queueing network models for quality control and testing. *European Journal of Operational Research*, 198:859–866, 2009.

Conference Papers

- H. Feichtinger, D. Onchis, B. Ricaud, B. Torr sani, and C. Wismeyr. A method for optimizing the ambiguity function concentration. In *Proceedings of the European Signal Processing Conference*, pages 804–808. IEEE, 2012.