

NUMERICAL COMPUTATIONS USING GABOR FRAMES FOR PERIODIC SPACES.

PETER LEMPEL SØNDERGAARD.

RESUME

This thesis presents Gabor frames for $L^2(\mathbb{T})$ and for the space of periodic sequences $\mathbb{C}_p^L$, as being analogue to Fourier series and the discrete Fourier transform. Results from the standard theory of Gabor frames for $L^2(\mathbb{R})$ are transferred to the periodic setting. This includes results about sampling Gabor systems and existence of Gabor frames. The interrelation of Gabor frames for the continuous space $L^2(\mathbb{T})$ and the discrete space $\mathbb{C}_p^L$ is afterwards used to construct a new method for numerical differentiation with properties close to the method of spectral differentiation. Finally, the various results are illustrated by numerical examples.

RESUMÉ

Rapporten præsenterer Gabor frames for $L^2(\mathbb{T})$ og for rummet af periodiske talfølger $\mathbb{C}_p^L$ som erstatning for fourierrækker og den diskrete fouriertransformation. Resultater fra Gabor frame teori for $L^2(\mathbb{R})$ overføres til $L^2(\mathbb{T})$. Dette inkluderer resultater omhandlende sampling af Gabor systemer og eksistens af Gabor frames. Sammenhængen mellem Gabor frames for det kontinuerte rum $L^2(\mathbb{T})$ og det diskrete rum $\mathbb{C}_p^L$ bliver herefter brugt til at konstruere en ny metode til numerisk differentiation, med egenskaber meget lig spectral differentiation. Til slut i rapporten illustreres de forskellige resultater med numeriske eksempler.

CONTENTS

Resume	1
Resumé	1
List of symbols	2
1. Introduction.	3
2. Periodic Spaces.	3
3. Fourier analysis.	4
3.1. The discrete Fourier transform.	4
3.2. Fourier series.	8
4. Gabor frames for $\mathbb{C}_p^L$	10
4.1. Frames	10
4.2. Basic properties.	12
4.3. Algorithms.	19
5. Gabor frames for $L^2(\mathbb{T})$.	25
5.1. Basic properties.	25
5.2. Sampling.	31
5.3. Bounds.	41
5.4. Existence.	45

6. Interpolation and Numerical differentiation	48
6.1. Polynomial differentiation	49
6.2. Spectral differentiation	50
6.3. The Gabor derivative.	53
7. Numerical experiments.	57
7.1. Dual windows.	57
7.2. Differentiation.	64
8. Concluding remarks.	68
References	69

LIST OF SYMBOLS¹

Symbol: Description:

a.e.	“almost everywhere”. Statements holds true except on a set of measure zero.
$\mathbb{R}_p^L, \mathbb{C}_p^L$	The spaces of real and complex periodic sequences of length L .
$\mathbb{T}$	The circle parameterized by $[0; 1[$.
$L^2(\mathbb{T})$	The space of square-integrable functions on the circle $\mathbb{T}$.
$C^p(\mathbb{T})$	The space of p times differentiable functions with continuous p 'th derivative on $\mathbb{T}$.
$\mathbb{R}^{m,n}, \mathbb{C}^{m,n}$	The spaces of real and complex matrices of size $m \times n$.
$\{f_n\}, \{g_n\}$	Sequences.
f, g, h	Vectors.
$\mathbf{A}, \mathbf{B}, \mathbf{C}$	Matrices.
$\mathbf{A}_{:,i}$	The i 'th column of a matrix.
$\mathbf{A}^*$	Conjugate transpose of the matrix $\mathbf{A}$.
$\mathbf{A} \otimes \mathbf{B}$	Kronecker product.
$\hat{f}$	The DFT or Fourier series for a vector f .
$\mathcal{F}$	The discrete Fourier transform or the Fourier series operator.
$\bar{f}$	Complex conjugation of f .
$f^{(p)}$	The p 'th derivative of f .
$\lfloor \frac{a}{b} \rfloor$	The result of a divided by b rounded down.
$\mathcal{O}(\cdot)$	Bounded from above, used only for computing-speed considerations.
ω_L	The first complex L -root of unity, $\omega_L = e^{\frac{2\pi i}{L}}$.
$\theta_{k,L}$	The k 'th mode of length L , $\theta_{k,L} = \left\{ \frac{1}{\sqrt{L}} e^{2\pi i \frac{kn}{L}} \right\}_n$.
$\mathbf{F}_L$	The Fourier-matrix of order L .
$\mathcal{I}, \mathbf{I}_L$	The identity operator and the identity matrix of size L .
δ_j	Kronecker's delta sequence, j being an integer.
$\mathcal{C}, (\mathcal{C}_\gamma)$	The analysis operator (with window γ).
$\mathcal{D}, (\mathcal{D}_g)$	The pre-frame operator (with window g).
$\mathcal{S}$	The frame operator.
$\mathcal{T}_k$	The translation operator.

¹The notation used in this thesis differ somewhat from the typical notation used in Gabor analysis literature. All operators are written using the `\mathcal` font: $\mathcal{D}, \mathcal{M}, \mathcal{T}$, matrices are written using bold font: $\mathbf{A}, \mathbf{B}, \mathbf{C}$ and Gabor systems are written with the `\mathfrak` font $\mathfrak{G}, \mathfrak{H}, \mathfrak{R}$. All sequences are indexed by subscript g_k instead of $g(k)$. This make Fourier series look a little bit out of the ordinary.

$\mathcal{M}_l$	The modulation operator.
$\mathfrak{G}, \mathfrak{H}, \mathfrak{J}, \mathfrak{K}$	Gabor systems.
g, γ	Gabor windows.
$\mathbf{D}_g, \mathbf{D}_\gamma$	Gabor matrices with window sequences $\{g_k\}$ and $\{\gamma_k\}$.
G_n	The n 'th correlation function.
$\mathcal{G}, \mathbf{G}$	The operator and matrix defined in theorem 5.35.
$\mathcal{P}\varphi_W$	The periodized Gaussian of width W .

1. INTRODUCTION.

When solving partial differential equations, one of the best methods is based on Fourier series. This works by approximating a function via a truncated Fourier series expansion, and differentiating this series termwise. This gives a remarkably good approximation to the derivative of the function, the so-called spectral derivative.

This only works well if the function is smooth. If the function is not smooth, or perhaps not even differentiable, then the spectral derivative has oscillations. Furthermore these oscillations affect the whole domain, because the complex exponentials in the Fourier series have global support. The complex exponentials have perfect frequency resolution but no time resolution.

Gabor frames have a good trade-off of time- and frequency resolution, and also a great flexibility in the choice of window function and parameters.

The idea behind this thesis is to use Gabor frames that resembles Fourier series and the discrete Fourier transform to construct a new method for numerical differentiation. Hopefully this new method will inherit some of the strengths of the spectral differentiation, while limiting some of the weaknesses.

Section 2 presents the relevant function spaces needed. Section 3 contains some necessary result from Fourier analysis. Section 4 presents some theory of Gabor frames for periodic sequences. These are the type of Gabor frames that will be used for the actual computation on a computer. The section therefore also contains some results to improve the computational efficiency.

Section 5 presents the theory of Gabor frames for the square integrable function on the circle, $L^2(\mathbb{T})$. These Gabor frames will be used for the error-analysis of the numerical differentiation. The results presented in this section have all been adapted from similar results for $L^2(\mathbb{R})$, sometimes with big simplifications in the complexity of the proofs. The results concern sampling of a Gabor frame, decay properties of the Gabor coefficients and existence results. The similar existence results for $L^2(\mathbb{R})$ are very complicated to check, but for $L^2(\mathbb{T})$ they simplify to something that can be checked numerically.

Section 6 presents the various methods for numerical differentiation and derives some error bounds, and section 7 presents numerical results for existence of Gabor frames and for the accuracy of the numerical differentiation methods.

2. PERIODIC SPACES.

Numerical computations on a computer can only be done on finite dimensional spaces. This could be a sequence of complex numbers, $\mathbb{C}^L$, enumerated by the numbers $0..L-1$. However, when doing calculations values enumerated by -1 and L are frequently needed. The simplest solution to this is to only consider *periodic* sequences. This means that the finite sequence of complex numbers will be enumerated by the ring $\mathbb{Z}_L = \mathbb{Z} \bmod L$. The value enumerated by

-1 is the value enumerated by $L - 1$ and so on. This solves any problems associated with the ends of the sequence, because there is no ends!

A complex, periodic sequence $\{f_n\}$ of length L will be denoted by $\mathbb{C}_p^L$, p for periodic. The space $\mathbb{C}_p^L$ will be equipped with the inner product

$$(2.1) \quad \langle \{f_n\}, \{g_n\} \rangle = \sum_{n=0}^{L-1} f_n \bar{g}_n.$$

This definition is very close to the the definition of the inner product for $\mathbb{C}^L$ and the arguments that show that $\mathbb{C}^L$ is a Hilbert space can also be used to show that $\mathbb{C}_p^L$ is a Hilbert space with inner product (2.1), and that (2.1) really is an inner product.

When there is no need to use the periodicity or the index of a sequence $\{f_n\}$, it will be referred to as a vector f .

The notation here presented, $\mathbb{C}_p^L$, is my own contribution.

The continuous space corresponding to $\mathbb{C}_p^L$ is the circle, $\mathbb{T}$, $\mathbb{T}$ for torus. In this thesis it will always be parameterized by the interval $[0; 1[$. A function on the circle $\mathbb{T}$ can be associated with a periodic function of period one on the real line $\mathbb{R}$. This identification will be used frequently.

The function space $L^2(\mathbb{T})$ is a Hilbert space with the inner product:

$$\langle f, g \rangle = \int_0^1 f(x) \bar{g}(x) dx.$$

The space will sometimes be written as $L^2([0; 1])$ if the mapping with periodic functions on $\mathbb{R}$ is used. Other spaces of this kind will also be used, for example written as $L^2([0; \frac{1}{N}])$. This corresponds to the L^2 space of periodic functions with period $\frac{1}{N}$ on $\mathbb{R}$.

Note that the relations between common function spaces are somewhat simpler than for the real line:

$$C^p(\mathbb{T}) \subset C(\mathbb{T}) \subset L^\infty(\mathbb{T}) \subset L^2(\mathbb{T}) \subset L^1(\mathbb{T})$$

Perspectives. Working with periodic spaces is not always convenient - not all problems are periodic of nature. In this thesis everything will be periodic for the sake of simplicity, and I will not make any considerations about non-periodic spaces. If results from this thesis should be applied to non-periodic problems, a way to do it would be to periodize the problems in some way and adapt the problem to the periodic space.

3. FOURIER ANALYSIS.

3.1. The discrete Fourier transform. This section is mostly inspired by [Briggs95] with the normalization of the DFT fitted to match [Stroh98], and the notation $\theta_{j,L}$ being my own addition.

Fourier analysis on finite, periodic sequences is done through the discrete Fourier transform:

Definition 3.1. The *discrete Fourier transform* of $\{f_n\} \in \mathbb{C}_p^L$ is a sequence $\{\hat{f}_k\}$ given by

$$\mathcal{F} \{f_n\} = \{\hat{f}_k\} = \left\{ \frac{1}{\sqrt{L}} \sum_{n=0}^{L-1} f_n e^{-2\pi i \frac{nk}{L}} \right\}_{k \in \mathbb{Z}}.$$

The discrete Fourier transform is abbreviated as the DFT.

Definition 3.1 only tells us that $\{\hat{f}_k\}$ is a sequence. The following proposition will show that it is also periodic.

Proposition 3.2. *The DFT of a periodic sequence $\{f_n\} \in \mathbb{C}_p^L$ is again a periodic sequence $\{\hat{f}_k\} \in \mathbb{C}_p^L$.*

Proof. Proving this amounts to proving that $\hat{f}_k = \hat{f}_{k \bmod L}$ for any k . Since the complex exponential function is periodic with period 2π then $e^{-2\pi i n \frac{k}{L}} = e^{-2\pi i n \frac{(k \bmod L)}{L}}$. $\square$

The DFT can be also be formulated as a matrix-vector product. For convenience define $\omega_L = e^{\frac{2\pi i}{L}}$. This is the first complex L -root of unity, i.e. $(\omega_L)^L = 1$. Furthermore, define the modes:

Definition 3.3. The k 'th mode $\theta_{k,L} \in \mathbb{C}_p^L$ of period L is defined by

$$\begin{aligned} \theta_{k,L} &= \left\{ \frac{1}{\sqrt{L}} e^{2\pi i \frac{kn}{L}} \right\}_{n \in \mathbb{Z}} \\ &= \left\{ \frac{1}{\sqrt{L}} \omega_L^{kn} \right\}_{n \in \mathbb{Z}}, \end{aligned}$$

for $k \in \mathbb{Z}$ and $L \in \mathbb{N}$.

The modes are periodic sequences of period L : Since the exponential function is periodic with period 2π , then the modes are periodic with at most period L .

Using ω_L and the modes, the definition of the DFT can be written as

$$\begin{aligned} (3.1) \quad \{\hat{f}_k\} &= \left\{ \frac{1}{\sqrt{L}} \sum_{n=0}^{L-1} f_n \omega_L^{-nk} \right\}_{k \in \mathbb{Z}} \\ &= \{\langle f, \theta_{k,L} \rangle\}_{k \in \mathbb{Z}}. \end{aligned}$$

This can be reformulated as a matrix product:

$$\hat{f} = \mathbf{F}_L^* f,$$

where the j, k 'th entry of the matrix $\mathbf{F}_L$ is given by

$$(\mathbf{F}_L)_{j,k} = \frac{1}{\sqrt{L}} \omega_L^{jk}.$$

This matrix is called the *Fourier matrix*²:

$$(3.2) \quad \mathbf{F}_L = \frac{1}{\sqrt{L}} \begin{pmatrix} 1 & 1 & 1 & \cdots & 1 \\ 1 & \omega_L^1 & \omega_L^2 & \cdots & \omega_L^{L-1} \\ 1 & \omega_L^2 & \omega_L^4 & \cdots & \omega_L^{2(L-1)} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & \omega_L^{(L-1)} & \omega_L^{2(L-1)} & \cdots & \omega_L^{(L-1)(L-1)} \end{pmatrix}.$$

The modes appears as the columns of the matrix.

To gain more insight in the structure of the Fourier matrix, the following important lemma about the modes is needed.

²The matrix defined this way is the transpose of the similar matrix defined in [Stroh98]. This makes the matrix methods for the DFT and IDFT resemble those for the Gabor frame methods more closely.

Lemma 3.4. Orthogonality relations. *Let $k, j \in \mathbb{Z}$ and $L \in \mathbb{N}$. Then*

$$\langle \theta_{j,L}, \theta_{k,L} \rangle = \delta_{|j-k| \bmod L},$$

where δ_j is Kronecker's delta function, defined by

$$\delta_j = \begin{cases} 1 & \text{if } j = 0 \\ 0 & \text{if } j \neq 0 \end{cases}.$$

Proof. All terms of the form ω_L^k are roots of the polynomial $z^L - 1$. This polynomial can be factored as

$$\begin{aligned} (3.3) \quad z^L - 1 &= (z - 1)(z^{L-1} + z^{L-2} \dots + z + 1) \\ &= (z - 1) \sum_{n=0}^{L-1} z^n. \end{aligned}$$

Now the inner product can be expanded as:

$$\begin{aligned} (3.4) \quad \langle \theta_{j,L}, \theta_{k,L} \rangle &= \sum_{n=0}^{L-1} \frac{1}{\sqrt{L}} \omega_L^{nj} \frac{1}{\sqrt{L}} \omega_L^{-nk} \\ &= \frac{1}{L} \sum_{n=0}^{L-1} \omega_L^{n(j-k)} = \frac{1}{L} \sum_{n=0}^{L-1} \left(\omega_L^{j-k} \right)^n. \end{aligned}$$

If $j = k$ modulo L then $z = \omega_L^{j-k} = 1$ and so $\langle \theta_{j,L}, \theta_{k,L} \rangle = 1$.

If $j \neq k$ modulo L then $z = \omega_L^{j-k} \neq 1$, and $z^L - 1 = 0$ because ω_L^{j-k} is a complex root of unity. Inserting this into (3.3) it is seen that $\sum_{n=0}^{L-1} z^n$ must be zero. This is the last term occurring in (3.4) so the conclusion is that $\langle \theta_{j,L}, \theta_{k,L} \rangle = 0$.

These two cases prove the lemma. $\square$

If this lemma is applied to the Fourier matrix the following proposition comes as the result.

Proposition 3.5. *The Fourier matrix $\mathbf{F}_L$ is a unitary matrix.*

Proof. This amounts to proving that $\langle \mathbf{F}_{:,j}, \mathbf{F}_{:,k} \rangle = \delta_{|j-k|}$, where $\mathbf{F}_{:,j}$ denotes the j 'th column of $\mathbf{F}$. Since the columns are the modes:

$$\begin{aligned} \langle \mathbf{F}_{:,j}, \mathbf{F}_{:,k} \rangle &= \langle \theta_{j,L}, \theta_{k,L} \rangle \\ &= \delta_{|j-k| \bmod L} \end{aligned}$$

the result follows from 3.4. The “mod L ” can be discarded since the indices only run in the range $0, \dots, L - 1$. $\square$

The inverse matrix of a complex unitary matrix is the conjugate transpose of the matrix, so this proves the existence of the *inverse discrete Fourier transform*, the IDFT, and provides a matrix expression:

$$f = (\mathbf{F}_L^*)^{-1} \hat{f} = \mathbf{F}_L \hat{f}.$$

From this a summation expression for the IDFT can be obtained:

Corollary 3.6. *The inverse discrete Fourier transform $\{f_n\}$ of a sequence $\{\hat{f}_k\}$ is a mapping from $\mathbb{C}_p^L \rightarrow \mathbb{C}_p^L$ and is given by*

$$\mathcal{F}^{-1} \left\{ \hat{f}_k \right\} = \{f_n\} = \left\{ \frac{1}{\sqrt{L}} \sum_{k=0}^{L-1} \hat{f}_k e^{2\pi i \frac{nk}{L}} \right\}_{n \in \mathbb{Z}}.$$

Proof. The summation expression is read directly from the matrix expression, and the periodicity is proven in the same manner as for the DFT. The only structural difference in the definition of the DFT and the IDFT is a minus-sign in the exponential function. $\square$

An important consideration for applications is how the DFT of a real signal $\{f_n\} \in \mathbb{R}_p^L$ behaves. In that case $f_n = f_n^*$ and so

$$\begin{aligned} \hat{f}_k^* &= \left(\frac{1}{\sqrt{L}} \sum_{n=0}^{L-1} f_n e^{-2\pi i \frac{nk}{L}} \right)^* \\ (3.5) \quad &= \frac{1}{\sqrt{L}} \sum_{n=0}^{L-1} f_n e^{-2\pi i \frac{n(-k)}{L}} \end{aligned}$$

$$(3.6) \quad = \hat{f}_{-k}.$$

This means that one does only need to compute and store half the complex DFT coefficients for a real signal, as the other half can be found by complex conjugation.

The DFT (and the IDFT) can be computed numerically using the matrix methods

$$\begin{aligned} \hat{f} &= \mathbf{F}_L^* f \\ f &= \mathbf{F}_L \hat{f}. \end{aligned}$$

This method requires $\mathcal{O}(L^2)$ floating point operations to compute either one of the transforms, where $\mathcal{O}(L^2)$ means that it is bounded from above by a constant times L^2 .

A more efficient way of computing the transforms is through the use of a *fast Fourier transform*, an FFT or for the inverse transform, an IFFT. These methods require $\mathcal{O}(L \log L)$ floating point operations.

Perspectives. So far I have chosen to enumerate finite periodic sequences with the integers running from 0 to $L - 1$. One could just as well use the indices $-\frac{L}{2} + 1$ to $\frac{L}{2}$. This gives the alternate definition of the DFT:

Definition 3.7. Using the indices $k \in -\frac{L}{2} + 1, \dots, \frac{L}{2}$ and $n \in -\frac{L}{2} + 1, \dots, \frac{L}{2}$ the DFT is defined by

$$\left\{ \hat{f}_k \right\} = \left\{ \frac{1}{\sqrt{L}} \sum_{n=-\frac{L}{2}+1}^{\frac{L}{2}} f_n e^{-2\pi i \frac{nk}{L}} \right\}_{k \in \mathbb{Z}}.$$

The coefficients obtained this way are exactly the same as for the original definition of the DFT, because everything work modulo L . This definition only works when L is even, but this is not a problem. It can easily be formulated for L being odd. Similarly, the IDFT

$$(3.7) \quad \left\{ f_n \right\} = \left\{ \frac{1}{\sqrt{L}} \sum_{k=-\frac{L}{2}+1}^{\frac{L}{2}} \hat{f}_k e^{2\pi i \frac{nk}{L}} \right\}_{n \in \mathbb{Z}}$$

At first these definitions look more troublesome because of the need to divide by two. There are however good reasons for using them, and these are the problems of *aliasing*.

The coefficient $\hat{f}_{L-1}$ is calculated by the inner product of f and $\theta_{L-1,L}$. The vector $\theta_{L-1,L}$ is formed from L samples of the complex function $e^{-2\pi i(L-1)x}$ in the interval $x \in [0; 1]$. If these samples are plotted one might expect to see some points forming $L - 1$ waves on the interval $[0; 1]$. This is done on figure 3.1.

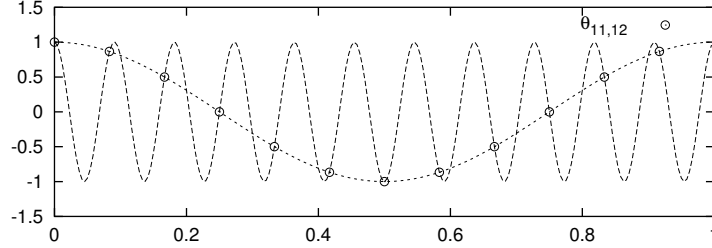


FIGURE 3.1. Plot of the mode $\theta_{11,12}$, with lines tracing out $e^{2\pi i(11)}$ and $e^{2\pi i(-1)}$. Only the real part is displayed.

The figure shows the 12 points of the mode $\theta_{11,12}$ tracing out the wave $e^{2\pi i(-1)}$ instead of the wave $e^{2\pi i(11)}$. This illustrates that the high-frequency mode $\theta_{11,12}$ is aliased to the low-frequency mode $\theta_{-1,12}$.

When using an already programmed DFT routine on a computer, one should check which definition of the DFT it uses, definition 3.1, definition 3.7 or a third definition. The software package MATLAB uses definition 3.1 and so does the software package OCTAVE, in which the numerical examples for this thesis are written.

The fast Fourier transform (FFT) methods provide a fast way of computing the DFT. They exist in many variants. Different algorithms for the computation of the FFT are presented in [Briggs95, chapter 10].

3.2. Fourier series. The section contains statements necessary to develop the theory of Gabor frames for $L^2(\mathbb{T})$ and to derive error estimates for the numerical differentiation methods that will be presented later.

Theorem 3.8. Plancherel's theorem for the circle. Let $f \in L^2(\mathbb{T})$ and let

$$(\mathcal{F}f)_k = \hat{f}_k = \int_0^1 f(x) e^{-2\pi i k x} dx$$

be the k 'th Fourier coefficient. Then f can be expanded as a Fourier series

$$f(x) = \sum_{k \in \mathbb{Z}} \hat{f}_k e^{2\pi i k x} \text{ a.e. } x \in \mathbb{T}.$$

and

$$\int_0^1 |f(x)|^2 dx = \|f\|_{L^2(\mathbb{T})}^2 = \sum_{k \in \mathbb{Z}} |\hat{f}_k|^2 = \|\hat{f}\|_{l^2(\mathbb{Z})}^2.$$

Proof. This is [Groch01, theorem 1.3.1] or [Rud66, theorem 9.13]. □

The next lemma establishes a Fourier expansion for a periodized function on $\mathbb{T}$. The lemma will be used to replace Poisson's summation formula for $L^2(\mathbb{R})$ when transferring proofs to the periodic setting.

Lemma 3.9. *Poisson summation for the circle.* *Let $f \in C(\mathbb{T})$ and $N \in \mathbb{N}$. Then*

$$\sum_{n=0}^{N-1} f\left(x + \frac{n}{N}\right) = N \sum_{k \in \mathbb{Z}} \hat{f}_{kN} e^{2\pi i k N x}, \quad x \in [0; \frac{1}{N}].$$

Proof. Let

$$\varphi(x) = \sum_{n=0}^{N-1} f\left(x + \frac{n}{N}\right), \quad x \in [0; \frac{1}{N}]$$

and

$$\psi(x) = \varphi\left(\frac{x}{N}\right), \quad x \in \mathbb{T}.$$

Since ψ is periodic with period 1, because φ is periodic with period $\frac{1}{N}$, then ψ can be expressed as a Fourier series with coefficients

$$\begin{aligned} \hat{\psi}(k) &= \int_0^1 \psi(x) e^{-2\pi i k x} dx \\ &= \int_0^1 \varphi\left(\frac{x}{N}\right) e^{-2\pi i k x} dx \\ &= N \int_0^{\frac{1}{N}} \varphi(x) e^{-2\pi i k N x} dx \\ &= N \int_0^{\frac{1}{N}} \left(\sum_{n=0}^{N-1} f\left(x + \frac{n}{N}\right) e^{-2\pi i k N x} \right) dx \\ &= N \sum_{n=0}^{N-1} \left(\int_0^{\frac{1}{N}} f\left(x + \frac{n}{N}\right) e^{-2\pi i k N (x + \frac{n}{N})} dx \right) \\ &= N \sum_{n=0}^{N-1} \left(\int_{\frac{n}{N}}^{\frac{n+1}{N}} f(x) e^{-2\pi i k N x} dx \right) \\ &= N \int_0^1 f(x) e^{-2\pi i k N x} dx \\ &= N \hat{f}_{kN}. \end{aligned}$$

In line five it is used that since the complex exponential function is periodic, then $e^{-2\pi i k N x} = e^{-2\pi i k N (x + \frac{n}{N})}$, and then the order of summation and integration is interchanged. This works well since both are finite. Now the summation and integration form a disjoint partition of the interval $[0; 1]$ (the overlap has measure zero) in line six, and so they are replaced by one combined integration.

With this, ψ can be expressed by its Fourier series

$$\begin{aligned} \psi(x) &= \sum_{k \in \mathbb{Z}} \hat{\psi}_k e^{2\pi i k x} \\ &= N \sum_{k \in \mathbb{Z}} \hat{f}_{kN} e^{2\pi i k x} \end{aligned}$$

for $x \in \mathbb{T}$ and this gives a similar expansion of φ :

$$\begin{aligned}\varphi(x) &= \psi(Nx) \\ &= N \sum_{k \in \mathbb{Z}} \hat{f}_{kN} e^{2\pi i k N},\end{aligned}$$

for $x \in [0; \frac{1}{N}]$. □

Remark 3.10. The lemma can be extended by density ([Groch01, A.1]) to cover all $f \in L^2(\mathbb{T})$, but then only with convergence almost everywhere $x \in [0; \frac{1}{N}]$.

Theorem 3.11. *Let $f \in C^{p-1}(\mathbb{T})$. Assume that $f^{(p)}$ is bounded and piecewise monotone. Then the Fourier series coefficients of f satisfy*

$$|c_k| \leq \frac{C}{|k|^{p+1}}, \forall k \in \mathbb{Z},$$

where C is a constant independent of k .

Proof. This is [Briggs95, theorem 6.2]. □

The Fourier series coefficients $\{c_k\}$ can be approximated by the DFT-coefficients $\{\frac{1}{\sqrt{L}}\hat{f}_k\}$. The error introduced by doing this is outlined in the following theorem.

Theorem 3.12. *Let $f \in C^{p-1}(\mathbb{T})$. Assume that $f^{(p)}$ is bounded and piecewise monotone. Then the error of the DFT coefficients $\hat{f}$ of order L as an approximation to the Fourier series coefficients c_k satisfy*

$$\left| \hat{f}_k - \frac{1}{\sqrt{L}} c_k \right| \leq \frac{C}{L^{p+1}}, \quad k = -\frac{L}{2} + 1, \dots, \frac{L}{2},$$

where C is a constant independent of k and L .

Proof. [Briggs95, theorem 6.3]. □

Remark 3.13. In [Briggs95] the theorem is extended to also cover the situation where f is not even continuous, but only bounded and piecewise monotone. In this case

$$(3.8) \quad \left| \hat{f}_k - \frac{1}{\sqrt{L}} c_k \right| \leq \frac{C}{L}$$

as could be expected from theorem 3.12.

4. GABOR FRAMES FOR $\mathbb{C}_p^L$

This section is a transfer of some of the results I showed in [Soen02] and the last part is based upon [Stroh98] but with more details.

4.1. Frames. In the first part of this project [Soen02] I showed some general theory of frames for Hilbert spaces. Since $\mathbb{C}_p^L$ is a Hilbert space all these results are applicable. I will briefly repeat some important results and definitions here.

Definition 4.1. A family of elements $\{e_j\}_{j \in J}$ is a *frame* for a separable Hilbert space $\mathcal{H}$ if constants $0 < A \leq B < \infty$ exists such that

$$(4.1) \quad A \|f\|_{\mathcal{H}}^2 \leq \sum_{j \in J} |\langle f, e_j \rangle_{\mathcal{H}}|^2 \leq B \|f\|_{\mathcal{H}}^2, \quad \forall f \in \mathcal{H}.$$

The constants A and B are denoted the lower and upper frame bound.

A sequence that only satisfies the upper frame bound is called a Bessel sequence:

Definition 4.2. A family of elements $\{e_j\}_{j \in J}$, $e_j \in \mathcal{H}$ is a *Bessel sequence* if

$$\sum_{j \in J} |\langle f, e_j \rangle_{\mathcal{H}}|^2 \leq B \|f\|_{\mathcal{H}}^2, \quad \forall f \in \mathcal{H}$$

for some constant $0 < B < \infty$. The constant B is called the *Bessel bound*.

Definition 4.3. Let $\{e_j\}_{j \in J}$ be a frame for a Hilbert space $\mathcal{H}$ and let $f \in \mathcal{H}$. Then the *analysis operator* is defined by

$$\mathcal{C}f = \{\langle f, e_j \rangle_{\mathcal{H}}\}_{j \in J}.$$

The analysis operator returns the inner products of the vector f with the frame elements $\{e_j\}_{j \in J}$.

Proposition 4.4. Let $\{e_j\}_{j \in J}$ be a frame for a Hilbert space $\mathcal{H}$. Then the analysis operator $\mathcal{C} : \mathcal{H} \rightarrow l^2(J)$ is a bounded operator $\mathcal{C} : \mathcal{H} \rightarrow l^2(J)$.

Proof. [Soen02] or [Groch01, prop. 5.1.1]. □

Definition 4.5. Let $\{e_j\}_{j \in J}$ be a frame for a Hilbert space $\mathcal{H}$. Then the *pre-frame operator* $\mathcal{D} : l^2(J) \rightarrow \mathcal{H}$ is defined by

$$\mathcal{D}\{c_j\} = \sum_{j \in J} c_j e_j.$$

The pre-frame operator assembles a vector f from its decomposition $\{c_j\}$ in the frame $\{e_j\}_{j \in J}$.

Proposition 4.6. Let $\{e_j\}_{j \in J}$ be a frame for a Hilbert space $\mathcal{H}$. Then the pre-frame operator $\mathcal{D}$ is the adjoint operator of the analysis operator $\mathcal{C}$.

Theorem 4.7. Let $\{e_j\}_{j \in J}$ be a frame for a Hilbert space $\mathcal{H}$. Then the frame operator defined as

$$\mathcal{S} : \mathcal{H} \rightarrow \mathcal{H}, \quad \mathcal{S}f = \mathcal{D}\mathcal{C}f = \sum_{j \in J} \langle f, e_j \rangle_{\mathcal{H}} e_j$$

is bounded, self-adjoint and invertible.

Proposition 4.8. Let $\{e_j\}_{j \in J}$ be a frame for a Hilbert space $\mathcal{H}$ with frame operator $\mathcal{S}$. The family $\{\mathcal{S}^{-1}e_j\}_{j \in J}$ is also a frame, the so-called dual frame.

Theorem 4.9. Let $\{e_j\}_{j \in J}$ be a frame for a Hilbert space $\mathcal{H}$. Then any $f \in \mathcal{H}$ can be written as

$$f = \sum_{j \in J} \langle f, \mathcal{S}^{-1}e_j \rangle_{\mathcal{H}} e_j.$$

The proofs of the preceding theorems and propositions can all be found in [Soen02, section 3.1] or in [Chr02, Groch01].

4.2. Basic properties. The definition of a Gabor frame for $\mathbb{C}_p^L$ requires the definition of the translation and modulation operators.

Definition 4.10. The translation $\mathcal{T}_n$ and the modulation $\mathcal{M}_m$ operators working on a sequence $\{f_k\} \in \mathbb{C}_p^L$ are defined by

$$\begin{aligned}\mathcal{T}_n \{f_k\} &= \{f_{k+n}\}_{k \in \mathbb{Z}} \\ \mathcal{M}_m \{f_k\} &= \left\{ e^{-2\pi i \frac{mk}{L}} f_k \right\}_{k \in \mathbb{Z}} = \left\{ \omega_L^{-mk} f_k \right\}_{k \in \mathbb{Z}}\end{aligned}$$

for all $m, n \in \mathbb{Z}$.

The translation of a periodic sequence is again a periodic sequence, since all the values have just been shifted by the same distance. Similarly for the modulation operator, the modulation of a periodic sequence is again a periodic sequence, since the modulation consist of pointwise multiplying the sequence with a periodic sequence, and all the periods are equal, namely L .

The operators are also periodic in their parameters:

$$\begin{aligned}\mathcal{T}_n \{f_k\} &= \mathcal{T}_{n \bmod L} \{f_k\} \\ \mathcal{M}_m \{f_k\} &= \mathcal{M}_{m \bmod L} \{f_k\}.\end{aligned}$$

The periodicity in m of $\mathcal{M}_m$ follows from the periodicity of the exponential function.

The inverse operators are simply the opposite action of the operators:

$$\begin{aligned}(\mathcal{T}_n)^{-1} \{f_k\} &= \mathcal{T}_{-n} \{f_k\} \\ (\mathcal{M}_m)^{-1} \{f_k\} &= \mathcal{M}_{-m} \{f_k\}.\end{aligned}$$

Composing the operators is straightforward:

$$\begin{aligned}\mathcal{T}_n \mathcal{T}_l \{f_k\} &= \mathcal{T}_{n+l} \{f_k\} \\ \mathcal{M}_m \mathcal{M}_l \{f_k\} &= \mathcal{M}_{m+l} \{f_k\}\end{aligned}$$

The two operators do not commute:

$$(4.2) \quad \mathcal{M}_m \mathcal{T}_n \{f_k\} = \left\{ e^{-2\pi i \frac{mk}{L}} f_{k+n} \right\}$$

$$(4.3) \quad = \left\{ e^{2\pi i \frac{mn}{L}} e^{-2\pi i \frac{m}{L} \cdot (k+n)} f_{k+n} \right\}$$

$$(4.4) \quad = e^{2\pi i \frac{mn}{L}} \mathcal{T}_n \mathcal{M}_m \{f_k\}.$$

It is now possible to define a Gabor system for $\mathbb{C}_p^L$.

Definition 4.11. A *Gabor system* in $\mathbb{C}_p^L$ is a family of sequences defined by:

$$\mathfrak{G} = G(\{g_k\}, a, \frac{b}{L}) = \{\mathcal{M}_{mb} \mathcal{T}_{na} g\}_{m=0, \dots, M-1, n=0, \dots, N-1},$$

with $\{g_k\} \in \mathbb{C}_p^L$ and where $a, b \in \mathbb{N}$ and $Mb = Na = L$. This means that L must be divisible by a and b .

Each element in the system is called a *Gabor atom*:

$$\{(\mathfrak{G}_{m,n})_k\} = \{\mathcal{M}_{mb} \mathcal{T}_{na} g\}_k = \left\{ e^{-2\pi i \frac{mb}{L} \cdot k} g_{k+na} \right\}_{k \in \mathbb{Z}}.$$

A Gabor system that is also a frame is called a *Gabor frame*.

Definition 4.12. The *density* of a Gabor system for $\mathbb{C}_p^L$ as defined in 4.11 is given by

$$(4.5) \quad \frac{L}{MN} = \frac{ab}{L} = \frac{a}{M} = \frac{b}{N}.$$

The density is the ratio of the length of the vectors, and the number of vectors in the Gabor system. The density divides Gabor systems for $\mathbb{C}_p^L$ in three categories:

$\frac{L}{MN} > 1$: The system cannot be a frame or a basis, as there is too few vectors to span the space $\mathbb{C}_p^L$

$\frac{L}{MN} = 1$: If the system is a frame, then it is also a basis, since L vectors spanning an L -dimensional space constitutes a basis.

$\frac{L}{MN} < 1$: The system cannot be a basis, but it can be a frame.

This categorization correspond to the similar case for Gabor frames for $L^2(\mathbb{R})$, [Soen02, corollary 3.14 and 3.15].

The next proposition will present the main result for Gabor frames for $\mathbb{C}_p^L$.

Proposition 4.13. Let $\mathfrak{G}(\{g_k\}, a, \frac{b}{L})$, $\{g_k\} \in \mathbb{C}_p^L$, $a, b, L \in \mathbb{N}$ be a Gabor frame for $\mathbb{C}_p^L$. Then there exists a sequence $\gamma^0 \in \mathbb{C}_p^L$ such that $\mathfrak{H}(\{\gamma_k^0\}, a, \frac{b}{L})$ is a dual frame of $\mathfrak{G}$.

Proof. The method is to show that both the translation and modulation operator commutes with the frame operator.

The frame operator can be written

$$(4.6) \quad Sf = \sum_{n=0}^{N-1} \sum_{m=0}^{M-1} \langle f, \mathcal{M}_{mb} \mathcal{T}_{na} g \rangle \mathcal{M}_{mb} \mathcal{T}_{na} g.$$

This is the definition of the frame operator from theorem 4.7 with the Gabor atoms inserted. Consider the following expression

$$(4.7) \quad (\mathcal{M}_{rb} \mathcal{T}_{sa})^{-1} S (\mathcal{M}_{rb} \mathcal{T}_{sa}) f$$

$$(4.8) \quad = \sum_{n=0}^{N-1} \sum_{m=0}^{M-1} \langle \mathcal{M}_{rb} \mathcal{T}_{sa} f, \mathcal{M}_{mb} \mathcal{T}_{na} g \rangle (\mathcal{M}_{rb} \mathcal{T}_{sa})^{-1} \mathcal{M}_{mb} \mathcal{T}_{na} g.$$

for $r, s \in \mathbb{Z}$. This was just a matter of inserting the definition of the frame operator (4.6) and moving the inverted translation and modulation inside the summation. This is perfectly legal since the summation is finite. The inner product can be rewritten as

$$\begin{aligned} \langle \mathcal{M}_{rb} \mathcal{T}_{sa} f, \mathcal{M}_{mb} \mathcal{T}_{na} g \rangle &= \langle f, \mathcal{T}_{-sa} \mathcal{M}_{(m-r)b} \mathcal{T}_{na} g \rangle \\ &= \left\langle f, e^{2\pi i ab(m-r)s} \mathcal{M}_{(m-r)b} \mathcal{T}_{-sa} \mathcal{T}_{na} g \right\rangle \\ &= e^{-2\pi i ab(m-r)s} \langle f, \mathcal{M}_{(m-r)b} \mathcal{T}_{(n-s)a} g \rangle, \end{aligned}$$

using the inverse operators of translation and modulation and using the commutation relation (4.4). The phase factor changes sign when pulled outside the inner product, because the inner product is conjugate linear in the second slot.

The last part of the expression can be rewritten in a similar manner

$$\begin{aligned} (\mathcal{M}_{rb} \mathcal{T}_{sa})^{-1} \mathcal{M}_{mb} \mathcal{T}_{na} g &= \mathcal{T}_{-sa} \mathcal{M}_{(m-r)b} \mathcal{T}_{na} g \\ &= e^{2\pi i ab(m-r)s} \mathcal{M}_{(m-r)b} \mathcal{T}_{-sa} \mathcal{T}_{na} g \\ &= e^{2\pi i ab(m-r)s} \mathcal{M}_{(m-r)b} \mathcal{T}_{(n-s)a} g. \end{aligned}$$

Inserting these two expressions in (4.8) shows that the phase factor cancel

$$\begin{aligned}
 (4.9) \quad & \sum_{n=0}^{N-1} \sum_{m=0}^{M-1} \langle \mathcal{M}_{rb} \mathcal{T}_{sa} f, \mathcal{M}_{mb} \mathcal{T}_{na} g \rangle (\mathcal{M}_{rb} \mathcal{T}_{sa})^{-1} \mathcal{M}_{mb} \mathcal{T}_{na} g \\
 &= \sum_{n=0}^{N-1} \sum_{m=0}^{M-1} \langle f, \mathcal{M}_{(m-r)b} \mathcal{T}_{(n-s)a} g \rangle \mathcal{M}_{(m-r)b} \mathcal{T}_{(n-s)a} g \\
 &= \sum_{n=0}^{N-1} \sum_{m=0}^{M-1} \langle f, \mathcal{M}_{mb} \mathcal{T}_{na} g \rangle \mathcal{M}_{mb} \mathcal{T}_{na} g \\
 &= \mathcal{S}f.
 \end{aligned}$$

The second line shows something that looks like the Gabor frame operator, except that all the summation indices are shifted. However, since both the translation and the modulation operator are L -periodic in their parameter, the second and the third line contain exactly the same terms, just summed in a different order.

To sum up, it has been shown that

$$(\mathcal{M}_{rb} \mathcal{T}_{sa})^{-1} \mathcal{S} (\mathcal{M}_{rb} \mathcal{T}_{sa}) = \mathcal{S}.$$

This can be transferred to the inverse operator:

$$\begin{aligned}
 \mathcal{S}^{-1} f &= \left((\mathcal{M}_{rb} \mathcal{T}_{sa})^{-1} \mathcal{S} (\mathcal{M}_{rb} \mathcal{T}_{sa}) \right)^{-1} f \\
 &= (\mathcal{M}_{rb} \mathcal{T}_{sa})^{-1} \mathcal{S}^{-1} (\mathcal{M}_{rb} \mathcal{T}_{sa}) f,
 \end{aligned}$$

so the same relation holds true for the inverse operator. A simple application of this commutation relation is

$$\begin{aligned}
 \mathcal{S}^{-1} \mathcal{M}_{rb} \mathcal{T}_{sa} g &= \mathcal{M}_{rb} \mathcal{T}_{sa} \mathcal{S}^{-1} g \\
 &= \mathcal{M}_{rb} \mathcal{T}_{sa} \gamma^0,
 \end{aligned}$$

where

$$\gamma^0 = \mathcal{S}^{-1} g.$$

Applying the inverse frame operator to a Gabor atom is the same as applying the translation and modulation directly to a new window function. This new window function $\gamma^0 = \mathcal{S}^{-1} g$ is called the *canonical dual sequence* of the Gabor frame. $\square$

The pre-frame (or synthesis) operator for a Gabor frame for $\mathbb{C}_p^L$ with window $\{g_n\}$ is given by

$$(4.10) \quad \mathcal{D}_g \{c_{m,n}\} = \sum_{n=0}^{N-1} \sum_{m=0}^{M-1} c_{m,n} \mathcal{M}_{mb} \mathcal{T}_{na} g,$$

from definition 4.5. The pre-frame operator for a Gabor frame for $\mathbb{C}_p^L$ will be called the *inverse discrete Gabor transform*, the IDGT.

The operator can be formulated using a matrix

$$(4.11) \quad \mathcal{D}_g \{c_{m,n}\} = \mathbf{D}_g c$$

if an ordering of the double indices m and n is chosen.

The ordering used in this thesis is that $c_{m,n}$ is placed in a vector c such that $c_{m+nM} = c_{m,n}$, that is $c_{0,0}$ is the first element followed by $c_{1,0}$ and the $M+1$ 'th element is $c_{0,1}$ followed by

$c_{1,1}$ and so on. This forces an ordering such that the Gabor atom $\mathcal{M}_0\mathcal{T}_0g$ is the first column in the matrix $\mathbf{D}_G$ followed by $\mathcal{M}_{1b}\mathcal{T}_0g$ as the second column and $\mathcal{M}_0\mathcal{T}_{1a}g$ as the $(M+1)$ 'th column.

This structure serves as the definition of a Gabor matrix:

Definition 4.14. A *Gabor matrix with window g* is a matrix $\mathbf{D}_g \in \mathbb{C}^{L,MN}$ with elements $(\mathbf{D}_g)_{m,n}$ given by:

$$\begin{aligned} (\mathbf{D}_g)_{m,n} &= \left(\mathcal{M}_{b\lfloor \frac{n}{M} \rfloor} \mathcal{T}_{a(n \bmod M)} g \right)_m \\ (4.12) \quad &= e^{-2\pi i b \lfloor \frac{n}{M} \rfloor \frac{m}{L}} g_{m+a(n \bmod M)}, \end{aligned}$$

where $\lfloor \frac{n}{M} \rfloor$ denotes the result of the division $\frac{n}{M}$ rounded towards zero.

The operator norm of the pre-frame operator can be easily calculated.

Proposition 4.15. Let $\{c_{m,n}\} \in \mathbb{C}^{M,N}$ and let $\{g_k\} \in \mathbb{C}_p^L$, $M, N, a, b, L \in \mathbb{N}$ with $Ma = Nb = L$. Then the pre-frame operator has operator norm

$$\|\mathcal{D}_{\{g\}}\{c_{m,n}\}\|_{\text{op}} \leq \sqrt{LMN} \max_j |g_j|$$

Proof. Define $f \in \mathbb{C}_p^L$ by

$$\{f_k\} = \sum_{n=0}^{N-1} \sum_{m=0}^{M-1} c_{m,n} \mathcal{M}_{mb} \mathcal{T}_{na} \{g_k\}.$$

Then

$$\begin{aligned} \|\{f_k\}\|_2^2 &= \sum_{k=0}^{L-1} \left| \sum_{n=0}^{N-1} \sum_{m=0}^{M-1} c_{m,n} e^{-2\pi i m b \frac{k}{L}} g_{k+na} \right|^2 \\ &\leq \sum_{k=0}^{L-1} \left(\sum_{n=0}^{N-1} \sum_{m=0}^{M-1} \left| c_{m,n} e^{-2\pi i m b \frac{k}{L}} g_{k+na} \right| \right)^2 \\ &\leq L \left(\max_j |g_j| \right)^2 \left(\sum_{n=0}^{N-1} \sum_{m=0}^{M-1} 1 \cdot |c_{m,n}| \right)^2 \\ &= L \left(\max_j |g_j| \right)^2 MN \|\{c_{m,n}\}\|_2^2. \end{aligned}$$

In line three the reference to $\{g_k\}$ is bounded by its maximum value. This can then be pulled outside the summations and the summation over k can be removed. Then Cauchy-Schwartz inequality can be used on the two sequences $\{1\}_{m,n}$ and $\{c_{m,n}\}$, and this gives the result. $\square$

Remark 4.16. The previous proposition does not require that $\{g_k\}$ generates a frame, so the statement can be used even if the pre-frame operator is not associated to a frame.

The analysis operator for a Gabor frame for $\mathbb{C}_p^L$ with window sequence $\{\gamma_n\}$ is given by

$$(4.13) \quad \mathcal{C}_\gamma f = \{\langle f, \mathcal{M}_{mb} \mathcal{T}_{na} \gamma \rangle\}_{m=0, \dots, M-1, n=0, \dots, N-1},$$

from definition 4.3. This operator can be formulated using a matrix

$$(4.14) \quad \mathcal{C}_\gamma f = \mathbf{D}_\gamma^* f,$$

where $\mathbf{D}_\gamma$ is a Gabor matrix with window γ . The analysis operator for a Gabor frame for $\mathbb{C}_p^L$ will be called the *discrete Gabor transform*, the DGT.

The next expression for $\mathcal{S}$ is the so-called *Walnut's representation* [Groch01, theorem 6.3.2] in the discrete, periodic case.

Proposition 4.17. *Walnut's representation.* *The Gabor frame operator $\mathcal{S}$ can be expressed as a matrix $\mathbf{S} \in \mathbb{C}^{L,L}$ with*

$$(4.15) \quad \mathbf{S}_{m,n} = \begin{cases} M \sum_{k=0}^{N-1} (\mathcal{T}_{ka}g)_m \overline{(\mathcal{T}_{ka}g)_n} & \text{if } (n-m) \bmod M = 0 \\ 0 & \text{otherwise} \end{cases}.$$

When $\{g_k\} \in \mathbb{R}_p^L$ then $\mathbf{S} \in \mathbb{R}^{L,L}$ and consequently the canonical dual is real: $\{(\gamma^0)_k\} \in \mathbb{R}_p^L$.

Proof. As the frame operator by definition (as defined in theorem 4.7) is the composition of the pre-frame operator and the analysis operator, then it follows from (4.11) and (4.14) that

$$(4.16) \quad \mathcal{S}f = \mathcal{D}_g \mathcal{C}_g f = \mathbf{D}_g \mathbf{D}_g^* f = \mathbf{S}f.$$

Fully written out, the components of $\mathbf{S}$ are:

$$\mathbf{S}_{m,n} = \sum_{k=0}^{N-1} \sum_{l=0}^{M-1} (\mathcal{M}_{lb} \mathcal{T}_{ka}g)_m \overline{(\mathcal{M}_{lb} \mathcal{T}_{ka}g)_n}.$$

This expression can be simplified:

$$\begin{aligned} \mathbf{S}_{m,n} &= \sum_{k=0}^{N-1} \sum_{l=0}^{M-1} e^{-2\pi i \frac{lb}{L} m} g_{m+ka} e^{2\pi i \frac{lb}{L} n} \overline{g_{n+ka}} \\ &= \sum_{k=0}^{N-1} g_{m+ka} \overline{g_{n+ka}} \sum_{l=0}^{M-1} e^{-2\pi i \frac{lb}{L} m} e^{2\pi i \frac{lb}{L} n} \\ &= \sum_{k=0}^{N-1} g_{m+ka} \overline{g_{n+ka}} \sum_{l=0}^{M-1} e^{2\pi i l \frac{n-m}{M}}. \end{aligned}$$

The final line contains the term $\sum_{l=0}^{M-1} e^{2\pi i l \frac{n-m}{M}}$. This is the inner product

$$\left\langle \sqrt{M} \theta_{n-m, M}, \sqrt{M} \theta_{0, M} \right\rangle = M \delta_{n-m \bmod M},$$

because $\sqrt{M} \theta_{0, M} = \{1\}_l$ and $\sqrt{M} \theta_{n-m, M} = \left\{ e^{2\pi i l \frac{n-m}{M}} \right\}_l$. The result follows from lemma 3.4.

Using this, the expression for $\mathbf{S}$ becomes:

$$\mathbf{S}_{m,n} = \begin{cases} M \sum_{k=0}^{N-1} (\mathcal{T}_{ka}g)_m \overline{(\mathcal{T}_{ka}g)_n} & \text{if } (n-m) \bmod M = 0 \\ 0 & \text{otherwise} \end{cases}$$

It shows, that since the matrix elements of $\mathbf{S}$ are composed only of translations of $\{g\}$ and no modulations are involved, then if $\{g\}$ is real then $\mathbf{S}$ is real as well. This means that also $\mathbf{S}^{-1}$ is a real matrix and so the statement about the canonical dual follows. $\square$

Theorem 4.18. *Janssens representation.* *Let $\mathfrak{G}(\{g_k\}, a, \frac{b}{L})$, $\{g_k\} \in \mathbb{C}_p^L$, be a Gabor frame for $\mathbb{C}_p^L$.*

$$\mathcal{S} \{f_k\} = \sum_{n=0}^{N-1} \sum_{m=0}^{N-1} \langle \{f_k\}, \mathcal{M}_{mb} \mathcal{T}_{na} \{g_k\} \rangle \mathcal{M}_{mb} \mathcal{T}_{na} g$$

can be written as

$$(4.17) \quad \mathcal{S}\{f_j\} = \frac{N}{b} \sum_{l=0}^{b-1} \sum_{r=0}^{a-1} \langle \{g_j\}, \mathcal{M}_{rN} \mathcal{T}_{lM} \{g_j\} \rangle \mathcal{M}_{rN} \mathcal{T}_{lM} \{f_j\}.$$

Proof. The matrix form of Walnut's expression for the frame operator can be written by summations:

$$(4.18) \quad \begin{aligned} (\mathcal{S}\{f_l\})_j &= \sum_{l=0}^{L-1} \mathbf{s}_{j,l} f_l \\ &= M \sum_{l=0}^{b-1} \left(\sum_{k=0}^{N-1} g_{j+ka} \bar{g}_{j+lM+ka} \right) f_{j+lM}. \end{aligned}$$

This last expression comes from only including the contributing terms from (4.15).

The sequences G_l , $l = 0, \dots, b-1$

$$\{G_l\}_j = \left\{ \sum_{k=0}^{N-1} g_{j+ka} \bar{g}_{j+lM+ka} \right\}_j$$

are periodic with period a and so they can be written using a discrete Fourier transform of length a :

$$\{G_l\}_j = \frac{1}{\sqrt{a}} \sum_{r=0}^{a-1} \left(\widehat{G}_l \right)_r e^{2\pi i r \frac{j}{a}}, \quad l = 0, \dots, a-1,$$

with

$$\begin{aligned} \left(\widehat{G}_l \right)_r &= \frac{1}{\sqrt{a}} \sum_{j=0}^{a-1} (G_l)_j e^{-2\pi i j \frac{r}{a}} \\ &= \frac{1}{\sqrt{a}} \sum_{j=0}^{a-1} \sum_{k=0}^{N-1} g_{j+ka} \bar{g}_{j+lM+ka} e^{-2\pi i j \frac{r}{a}} \\ &= \frac{1}{\sqrt{a}} \sum_{j=0}^{L-1} g_j \bar{g}_{j+lM} e^{-2\pi i j \frac{rN}{L}} \\ &= \frac{1}{\sqrt{a}} \langle \{g_j\}, \mathcal{M}_{rN} \mathcal{T}_{lM} \{g_j\} \rangle_{\mathbb{C}_p^L} \end{aligned}$$

In line three the summations over k and j are combined into one and $j + ka$ is replaced by j . Because the exponential function is periodic, then the term containing the exponential function does not change.

This gives an expression for G_l :

$$\{G_l\}_j = \frac{1}{a} \sum_{r=0}^{a-1} \langle \{g_j\}, \mathcal{M}_{rN} \mathcal{T}_{lM} \{g_j\} \rangle e^{2\pi i r \frac{jN}{L}}.$$

Inserting this into (4.18):

$$\begin{aligned} \mathcal{S}\{f_l\} &= \left\{ \frac{M}{a} \sum_{l=0}^{b-1} \left(\sum_{r=0}^{a-1} \langle \{g_j\}, \mathcal{M}_{rN} \mathcal{T}_{lM} \{g_j\} \rangle e^{2\pi i r \frac{jN}{L}} \right) f_{j+lM} \right\}_j \\ &= \frac{N}{b} \sum_{l=0}^{b-1} \sum_{r=0}^{a-1} \langle \{g_j\}, \mathcal{M}_{rN} \mathcal{T}_{lM} \{g_j\} \rangle \mathcal{M}_{rN} \mathcal{T}_{lM} \{f_j\}. \end{aligned}$$

The last line follows from $\frac{M}{a} = \frac{N}{b}$. $\square$

Proposition 4.19. *Let $\mathfrak{G}(g, a, \frac{b}{L})$, $g \in \mathbb{C}_p^L$, $a, b, L \in \mathbb{N}$ be a Gabor frame for $\mathbb{C}_p^L$. Let $\gamma \in \mathbb{C}_p^L$ be a dual window of g . Then*

$$(4.19) \quad \langle \mathcal{M}_{mb} \mathcal{T}_{na} g, \mathcal{M}_{rb} \mathcal{T}_{sa} \gamma \rangle = \begin{cases} 1 & \text{if } r = m \text{ modulo } M \text{ and } n = s \text{ modulo } N \\ 0 & \text{otherwise} \end{cases}$$

for all $m, n, r, s \in \mathbb{Z}$.

Proof. When γ is a dual window of g then $f = \mathcal{D}_g \mathcal{C}_\gamma f$ for all $f \in \mathbb{C}_k^L$. Using the Gabor matrix notation this is

$$f = \mathbf{D}_g \mathbf{D}_\gamma^* f, \forall f \in \mathbb{C}_k^L,$$

which by standard linear algebra gives

$$\mathbf{D}_g \mathbf{D}_\gamma^* = \mathbf{I}_L,$$

where $\mathbf{I}_L$ is the identity matrix of size L . Looking at the structure of the matrices shows the result for $m, r \in 0, \dots, M-1$ and $n, s \in 0, \dots, N-1$. The extension to $m, n, r, s \in \mathbb{Z}$ follows from the periodicity of the translation and modulation operators. $\square$

Remark 4.20. A more careful statement of the previous result would yield the Wexler-Raz bi-orthogonality relations for $\mathbb{C}_p^L$ (see [Groch01, theorem 7.1] for the $L^2(\mathbb{R}^d)$ relations). The relations characterize all possible dual window functions by requiring them to satisfy (4.19).

Considerations. The requirements for the numbers a, b, M, N, L for a Gabor system for $\mathbb{C}_p^L$ as defined in definition 4.11, that $Mb = Na = L$ and $a, b, M, N, L \in \mathbb{N}$ put some restrictions on the freedom to choose these numbers. A simple example: If one wishes to keep the density $\frac{L}{MN}$ and the ratio between the number of time shifts and frequency shifts $\frac{M}{N}$ fixed, then it is not possible to double L .

Another example is the very common practice of choosing L to be a power of 2. This will force M and N to also be powers of 2, and then the density can only be $\frac{1}{2^p}$ for some $p \in \mathbb{N}$, so one could only find a Gabor basis with redundancy 1, or a rather redundant Gabor frame with redundancy $\frac{1}{2}$ or $\frac{1}{4}$. It would not be possible to find frames with an intermediate density.

Good choices of a, b, M, N tends to be numbers divisible by different small prime factors 2, 3 and 5. An example would be

$$\begin{aligned} L &= 360 \\ M = N &= 20 \\ a = b &= 18, \end{aligned}$$

giving a square system with a redundancy of $\frac{360}{400} = 0.9$.

The proof of proposition 4.13 that shows how to find the canonical dual sequence resembles the corresponding result for Gabor frames for $L^2(\mathbb{R})$ [Soen02, prop. 3.11] very much. Similarly

in [Soen02] there is one definition of the translation- and modulation operators and of a Gabor frame, and in this thesis there is another as well.

It would be tempting to search for a common theory covering both the continuous and the discrete case of Gabor frames. Such a theory is described in [Groech98].

The combined theory for the discrete and the continuous theory abstracts the space $(\mathbb{C}_p^L \text{ or } L^2(\mathbb{R}))$ with a locally compact Abelian group. The group composition is then used to define the translation operator. The modulation operator is defined by the use of the characters, which for the two spaces considered are

$$\begin{aligned}\chi_y(x) &= e^{2\pi i y x}, \quad y \in \mathbb{R}, \quad x \in \mathbb{R} \\ \{(\chi_m)_k\} &= e^{2\pi i \frac{mk}{L}}, \quad m \in \mathbb{Z}, \quad k \in \mathbb{Z}.\end{aligned}$$

The modulation operator is then defined by point wise multiplication by a character.

In this manner translation and modulation and Gabor frames can be defined in a generic way.

I will not pursue the theory of locally compact Abelian groups in this thesis, instead I have just repeated the proof of proposition 4.13 adapted to $\mathbb{C}_p^L$. It does however seem quite fruitful to have one combined theory instead of two or more separate theories, but defining the theory in all details would be more work than simply repeating the few proofs needed for this thesis.

4.3. Algorithms. This section is written solely from [Stroh98], but the proofs for proposition 4.23, theorem 4.24 and theorem 4.25 are my own additions as there are only hinted at in the article. Most notably is the use of $\lfloor \frac{a}{b} \rfloor$ not used in the article.

This section will present some faster ways of calculating $c = \mathbf{D}_g^* f$, $f = \mathbf{D}_g c$ and $\gamma = \mathbf{S}^{-1} g$. To do so, some linear algebra tools and definitions are needed.

Definition 4.21. The *Kronecker product* $\mathbf{A} \otimes \mathbf{B}$ where $\mathbf{A}$ is a $j \times k$ matrix and $\mathbf{B}$ a $l \times m$ matrix is a $jl \times km$ matrix defined block wise by

$$\mathbf{A} \otimes \mathbf{B} = \begin{bmatrix} a_{0,0}\mathbf{B} & \cdots & a_{0,k-1}\mathbf{B} \\ \vdots & & \vdots \\ a_{j-1,0}\mathbf{B} & \cdots & a_{j-1,k-1}\mathbf{B} \end{bmatrix}.$$

For Kronecker products it holds that

$$(4.20) \quad (\mathbf{A} \otimes \mathbf{B})^* = \mathbf{A}^* \otimes \mathbf{B}^*,$$

see [vL92].

Another simple rule for Kronecker products that will be needed is

$$(4.21) \quad (\mathbf{I}_N \otimes \mathbf{A})(\mathbf{I}_N \otimes \mathbf{B}) = \mathbf{I}_N \otimes \mathbf{AB},$$

where $\mathbf{I}_N$ denotes the identity matrix of size N , $\mathbf{I}_N \in \mathbb{C}^{N,N}$ and $\mathbf{A} \in \mathbb{C}^{b,M}$, $\mathbf{B} \in \mathbb{C}^{M,b}$.

The matrix $\mathbf{I}_N \otimes \mathbf{A}$ is block-diagonal with the matrix $\mathbf{A}$ along the diagonal. By a little inspection of the matrix products it shows that multiplying matrices of this kind only multiplies corresponding blocks on the diagonal, and so the result follows.

Definition 4.22. The *mod p perfect shuffle permutation matrix* $\mathbf{P}_{p,L} \in \mathbb{C}^{L,L}$ where $p, q, L \in \mathbb{N}$ and $pq = L$ is defined by

$$(4.22) \quad \begin{aligned} (\mathbf{P}_{p,L})_{j,k} &= \begin{cases} 1 & \text{if } kp + \left\lfloor \frac{k}{q} \right\rfloor \bmod L = j \\ 0 & \text{otherwise} \end{cases} \\ &= \delta_{\left| kp + \left\lfloor \frac{k}{q} \right\rfloor - j \right| \bmod L}. \end{aligned}$$

or equivalently

$$(4.23) \quad \begin{aligned} (\mathbf{P}_{p,L})_{j,k} &= \begin{cases} 1 & \text{if } jq + \left\lfloor \frac{j}{p} \right\rfloor \bmod L = k \\ 0 & \text{otherwise} \end{cases} \\ &= \delta_{\left| kp + \left\lfloor \frac{j}{p} \right\rfloor - k \right| \bmod L}. \end{aligned}$$

The definition (4.22) can be reformulated using $k = rq + s$, where $r = 0, \dots, p-1$ and $s = 0, \dots, q-1$. Then

$$\begin{aligned} (\mathbf{P}_{p,L})_{j,rq+s} &= \delta_{\left| (rq+s)p + \left\lfloor \frac{rq+s}{q} \right\rfloor - j \right| \bmod L} \\ &= \delta_{\left| sp+r-j \right| \bmod L} \\ &= \delta_{sp+r-j}. \end{aligned}$$

The last line follows from the fact that $sp + r < L$ for all s, r . From this it follows that a matrix element is one if and only if $j = sp + r$, so the ones are placed in the positions

$$(\mathbf{P}_{p,L})_{sp+r, rq+s}.$$

If j is replaced by $sp + r$ in (4.23) the result is the same, and so the two definitions are equal.

Since $\mathbf{P}_{p,L}$ is a permutation matrix it follows from basic linear algebra that it is also a unitary matrix, i.e. $\mathbf{P}_{p,L}^* = \mathbf{P}_{p,L}^{-1}$.

The definition implies that $\mathbf{P}_{p,L}$ has exactly one element per column with the value 1 instead of 0. For the first $\frac{L}{p}$ columns this element is placed when $kp \bmod L = j$. For the next $\frac{L}{p}$ columns this pattern is repeated, except that all the columns are shifted by 1. The following picture may illuminate this

$$\mathbf{P}_{p,L} = \begin{matrix} p \text{ rows} \left\{ \right. \\ \\ \left. \right. \end{matrix} \underbrace{\begin{bmatrix} 1 & 0 & \cdots & 0 & 0 & \cdots & 0 \\ & & & & 1 & & \\ & 0 & & & & & \\ & 1 & & & & & \\ \vdots & & & & & & \vdots \\ & & & 1 & & & \\ 0 & & & & & & 1 \end{bmatrix}}_{q \text{ columns}}.$$

Some definitions for diagonal matrices are also needed. The notion $\text{diag}(\mathbf{W}_0, \mathbf{W}_1, \dots, \mathbf{W}_{M-1})$ will mean a block-diagonal matrix of size $Mm \times Mn$ if each $\mathbf{W}_k$ are of size $m \times n$:

$$\text{diag}(\mathbf{W}_0, \mathbf{W}_1, \dots, \mathbf{W}_{M-1}) = \begin{bmatrix} \mathbf{W}_0 & \mathbf{0} & \cdots & \mathbf{0} \\ \mathbf{0} & \mathbf{W}_1 & & \mathbf{0} \\ \vdots & & \ddots & \\ \mathbf{0} & \mathbf{0} & & \mathbf{W}_{M-1} \end{bmatrix},$$

where $\mathbf{0}$ is the zero-matrix of size $m \times n$.

Proposition 4.23. *Let $\mathbf{D}_g$ be a Gabor matrix with window g . Then*

$$(\mathbf{D}_g (\mathbf{I}_N \otimes \mathbf{F}_M))_{m,n} = \begin{cases} \sqrt{M} g_{m+a \lfloor \frac{n}{M} \rfloor} & \text{if } |m-n| \bmod M = 0 \\ 0 & \text{otherwise} \end{cases},$$

where $\mathbf{I}_N$ denotes the identity matrix of size N .

Proof. The matrix $\mathbf{I}_N \otimes \mathbf{F}_M$ is a block-diagonal matrix with $\mathbf{F}_M$ as each block. Multiplying $\mathbf{D}_g$ from the right with this matrix amounts to inner products between panels of the rows of $\mathbf{D}_g$ and the columns of $\mathbf{F}_M$, i.e. the modes of length M . The panels of the rows are those elements of the rows that corresponds to the same translation of the window function g , which are the elements $0, \dots, M-1, M, \dots, 2M-1$ and so on.

The j 'th panel of M elements of the m 'th row of the Gabor matrix $\mathbf{D}_g$ consists of

$$\begin{aligned} & [\mathcal{T}_{ja}g_m, (\mathcal{M}_{1b}\mathcal{T}_{ja}g)_m, (\mathcal{M}_{2b}\mathcal{T}_{ja}g)_m, \dots, (\mathcal{M}_{(M-1)b}\mathcal{T}_{ja}g)_m] \\ &= g_{m+ja} [\omega_L^0, \omega_L^{-mb}, \omega_L^{-2mb}, \dots, \omega_L^{-(M-1)mb}] \\ (4.24) \quad &= g_{m+ja} \sqrt{L} \theta_{-m,M}. \end{aligned}$$

Line three comes from $\omega_L^{-jmb} = e^{-2\pi i \frac{jmb}{L}} = e^{-2\pi i \frac{jn}{M}} = \omega_M^{-jn}$, since $\frac{L}{b} = M$.

With (4.24) and knowledge of the columns of the Fourier matrix, the term $(\mathbf{D}_g (\mathbf{I}_N \otimes \mathbf{F}_M))_{m,n}$ can be calculated:

$$\begin{aligned} (\mathbf{D}_g (\mathbf{I}_N \otimes \mathbf{F}_M))_{m,n} &= \sum_{k=0}^{M-1} \sqrt{L} g_{m+a \lfloor \frac{n}{M} \rfloor} (\theta_{-m,M})_k (\theta_{n \bmod M, M})_k \\ &= \sqrt{L} g_{m+a \lfloor \frac{n}{M} \rfloor} \langle \theta_{-m,M}, \theta_{-n,M} \rangle \\ &= \sqrt{L} g_{m+a \lfloor \frac{n}{M} \rfloor} \delta_{|m-n| \bmod M}. \end{aligned}$$

The panel index j in (4.24) has been replaced by $\lfloor \frac{n}{M} \rfloor$, as in (4.12). The last line follows from the orthogonality relations, lemma 3.4. $\square$

Rearranging the terms in $\mathbf{D}_g (\mathbf{I}_N \otimes \mathbf{F}_M)$ leads to the following proposition.

Theorem 4.24. *Let $\mathbf{D}_g \in \mathbb{C}^{L, MN}$ be a Gabor matrix with window g . Then*

$$\mathbf{D}_{diag} = \mathbf{P}_{M,L}^* \mathbf{D}_g (\mathbf{I}_N \otimes \mathbf{F}_M) \mathbf{P}_{N,MN}^*$$

where $\mathbf{D}_{diag} \in \mathbb{C}^{L, MN}$ is a block-diagonal matrix

$$\mathbf{D}_{diag} = \text{diag}(\mathbf{W}_0, \dots, \mathbf{W}_{M-1})$$

with $\mathbf{W}_k \in \mathbb{C}^{b,N}$:

$$(\mathbf{W}_k)_{m,n} = \sqrt{M} g_{k+mM+na}$$

for $m = 0, \dots, b-1$, $n = 0, \dots, N-1$ and $k = 0, \dots, M-1$.

Proof. Proposition 4.23 proves the sparse structure of $\mathbf{D}_g(\mathbf{I}_N \otimes \mathbf{F}_M)$.

Definition 4.22 gives the structure of the permutation matrices

$$\begin{aligned} (\mathbf{P}_{M,L}^*)_{m,n} &= \delta_{|mM + \lfloor \frac{m}{b} \rfloor - n| \bmod L} \\ (\mathbf{P}_{N,MN}^*)_{m,n} &= \delta_{|nN + \lfloor \frac{n}{N} \rfloor - m| \bmod MN}. \end{aligned}$$

The first line is derived from (4.22) and the second from (4.23).

Using this and proposition 4.23

$$\begin{aligned} & (\mathbf{D}_g(\mathbf{I}_N \otimes \mathbf{F}_M) \mathbf{P}_{N,MN}^*)_{m,n} \\ &= \sum_{k=0}^{MN-1} \sqrt{L} g_{m+a \lfloor \frac{k}{M} \rfloor} \delta_{|m-k| \bmod M} \delta_{|nM + \lfloor \frac{n}{N} \rfloor - k| \bmod MN} \\ &= \sqrt{L} g_{m+a \left\lfloor \frac{nM + \lfloor \frac{n}{N} \rfloor \bmod MN}{M} \right\rfloor} \delta_{|m - (nM + \lfloor \frac{n}{N} \rfloor \bmod MN)| \bmod M} \\ &= \sqrt{L} g_{m+an} \delta_{|m - \lfloor \frac{n}{N} \rfloor| \bmod M}, \end{aligned}$$

for $m = 0, \dots, L-1$ and $n = 0, \dots, MN-1$. The summation in line two is removed because $\mathbf{P}_{N,MN}^*$ is a permutation matrix, and so there is exactly one element per column that is non-

zero. Therefore k can be replaced by $nM + \lfloor \frac{n}{N} \rfloor \bmod MN$. The term $a \left\lfloor \frac{nM + \lfloor \frac{n}{N} \rfloor \bmod MN}{M} \right\rfloor$ reduces to an since $n < MN$ and therefore $\left\lfloor \frac{\lfloor \frac{n}{N} \rfloor \bmod MN}{M} \right\rfloor = 0$ always.

The full product can now explicitly be written out.

$$\begin{aligned} \mathbf{D}_{diag} &= (\mathbf{P}_{M,L}^* \mathbf{D}_g(\mathbf{I}_N \otimes \mathbf{F}_M) \mathbf{P}_{N,MN}^*)_{m,n} \\ &= \sum_{k=0}^{L-1} \delta_{|mM + \lfloor \frac{m}{b} \rfloor - k| \bmod L} \sqrt{L} g_{k+an} \delta_{|k - \lfloor \frac{n}{N} \rfloor| \bmod M} \\ &= \sqrt{L} g_{(mM + \lfloor \frac{m}{b} \rfloor \bmod L) + an} \delta_{|(mM + \lfloor \frac{m}{b} \rfloor \bmod L) - \lfloor \frac{n}{N} \rfloor| \bmod M} \\ &= \sqrt{L} g_{mM + \lfloor \frac{m}{b} \rfloor + an} \delta_{|\lfloor \frac{m}{b} \rfloor - \lfloor \frac{n}{N} \rfloor| \bmod M} \\ &= \sqrt{L} g_{mM + \lfloor \frac{m}{b} \rfloor + an} \delta_{|\lfloor \frac{m}{b} \rfloor - \lfloor \frac{n}{N} \rfloor|}, \end{aligned}$$

for $m = 0, \dots, L-1$ and $n = 0, \dots, MN-1$. This follows the same line of reasoning as the previous formulas, k is replaced by $mM + \lfloor \frac{m}{b} \rfloor \bmod L$. Since g is periodic with period L , the “ $\bmod L$ ” in the first part of line three is unnecessary. The second “ $\bmod L$ ” can be removed, since $\frac{L}{M} = b$, so M always divides L . The final “ $\bmod M$ ” can be removed since $m < L = bM$ and $N < MN$ so $|\lfloor \frac{m}{b} \rfloor - \lfloor \frac{n}{N} \rfloor|$ is always less than M .

The last line shows that the result is only non-zero if $\lfloor \frac{m}{b} \rfloor = \lfloor \frac{n}{N} \rfloor$. This is only the case for elements in the $b \times N$ blocks on the diagonal, so it proves the structure $\mathbf{D}_{diag} = \text{diag}(\mathbf{W}_0, \dots, \mathbf{W}_{N-1})$ where $\mathbf{W}_k \in \mathbb{C}^{b,N}$.

Setting $k = \lfloor \frac{m}{b} \rfloor = \lfloor \frac{n}{N} \rfloor$ shows that

$$\begin{aligned} (\mathbf{W}_k)_{m,n} &= \sqrt{L} g_{(m+kb)M+k+a(n+kN)} \\ &= \sqrt{L} g_{mM+k+an}, \end{aligned}$$

for $m = 0, \dots, b-1$ and $n = 0, \dots, N-1$. Going from indices m, n running over the whole matrix to indices running in the sub-matrices, it is necessary to add the terms kb and kN , but they can be removed because of the periodicity of g . $\square$

The factorization from theorem 4.24 can be used to factor $\mathbf{S}$ into a block-diagonal matrix.

Theorem 4.25. *Let $\mathbf{S}$ be the matrix of the frame operator S for a Gabor frame $\mathfrak{G}(g, a, b)$ for $\mathbb{C}_p^L$. Then*

$$\mathbf{S}_{diag} = \mathbf{P}_{M,L}^* \mathbf{S} \mathbf{P}_{M,L}$$

where $\mathbf{S}_{diag}$ is a block-diagonal matrix

$$\mathbf{S}_{diag} = \text{diag}(\mathbf{B}_0, \dots, \mathbf{B}_{M-1})$$

with

$$\mathbf{B}_k = \mathbf{W}_k \mathbf{W}_k^*$$

for $m = 0, \dots, b-1$, $n = 0, \dots, b-1$ and $k = 0, \dots, M-1$ and $\mathbf{W}_k$ defined as in theorem 4.24.

Proof. A rearrangement of the result from theorem 4.24 shows that

$$\mathbf{D}_g = \mathbf{P}_{M,L} \mathbf{D}_{diag} \mathbf{P}_{N,MN} (\mathbf{I}_N \otimes \mathbf{F}_M)^*.$$

Inserting this into the matrix expression for the frame operator (4.16)

$$\begin{aligned} S &= \mathbf{D}_g \mathbf{D}_g^* \\ &= \mathbf{P}_{M,L} \mathbf{D}_{diag} \mathbf{P}_{N,MN} (\mathbf{I}_N \otimes \mathbf{F}_M)^* (\mathbf{P}_{M,L} \mathbf{D}_{diag} \mathbf{P}_{N,MN} (\mathbf{I}_N \otimes \mathbf{F}_M)^*)^* \\ &= \mathbf{P}_{M,L} \mathbf{D}_{diag} \mathbf{P}_{N,MN} (\mathbf{I}_N \otimes \mathbf{F}_M^*) (\mathbf{I}_N \otimes \mathbf{F}_M) \mathbf{P}_{N,MN}^* \mathbf{D}_{diag}^* \mathbf{P}_{M,L}^* \\ &= \mathbf{P}_{M,L} \mathbf{D}_{diag} \mathbf{D}_{diag}^* \mathbf{P}_{M,L}^*. \end{aligned}$$

In line three (4.20) is used to put the transposition inside the Kronecker product $\mathbf{I}_N \otimes \mathbf{F}_M^*$. Then (4.21), as well as the fact that $\mathbf{F}_M$ is a unitary matrix, are used to remove the Fourier matrices. Similarly the permutation matrices are also unitary and the pair $\mathbf{P}_{N,MN} \mathbf{P}_{N,MN}^*$ can be removed.

From this and (4.21) the result follows:

$$\mathbf{P}_{M,L}^* \mathbf{S} \mathbf{P}_{M,L} = \mathbf{D}_{diag} \mathbf{D}_{diag}^* = \mathbf{S}_{diag}$$

$\square$

The factorization from theorem 4.24 provides a faster way of numerically computing the IDGT than explicitly calculating

$$f = \mathbf{D}_g c.$$

Using theorem 4.24 this can be expanded to

$$f = \mathbf{P}_{M,L} \mathbf{D}_{diag} \mathbf{P}_{N,MN} (\mathbf{I}_N \otimes \mathbf{F}_M^*) c.$$

This is more efficient to calculate as the matrices can be applied in turn to the vector c :

$$(4.25) \quad c_1 = (\mathbf{I}_N \otimes \mathbf{F}_M^*) c$$

$$(4.26) \quad c_2 = \mathbf{P}_{N,MN} c_1$$

$$(4.27) \quad c_3 = \mathbf{D}_{diag} c_2$$

$$(4.28) \quad f = \mathbf{P}_{M,L} c_3$$

Calculating c_1 amounts to N small FFTs of length M of panels of c , c_2 is just a reordering of c_1 , c_3 can be calculated by multiplying panels of c_2 by the matrices $\mathbf{W}_k$, and finally f is a reordering of c_3 . The total amount of floating point operations needed is

$$\mathcal{O}(NM \log(M)) + \mathcal{O}(MbN) = \mathcal{O}(NM \log(M)) + \mathcal{O}(LN).$$

The reordering from the permutation matrices have not been included, as the operations needed for this are memory moves and not floating point operations. They amount to a cost of $\mathcal{O}(MN)$.

A completely similar thing can be done for the DGT:

$$c = \mathbf{D}_g^* f.$$

Again using theorem 4.24, this expands to

$$c = (\mathbf{I}_N \otimes \mathbf{F}_M) \mathbf{P}_{N,MN}^* \mathbf{D}_{\gamma,diag}^* \mathbf{P}_{M,L}^* f.$$

This corresponds to everything done in the reverse order as it was done for the IDGT, and the computational cost is exactly the same.

The computation of the canonical dual window can similar be done more efficiently than

$$\gamma^0 = \mathbf{S}^{-1} g.$$

Using theorem 4.25 and inserting

$$\begin{aligned} \gamma^0 &= (\mathbf{P}_{M,L} \mathbf{S}_{diag} \mathbf{P}_{M,L}^*)^{-1} g \\ &= \mathbf{P}_{M,L} (\mathbf{S}_{diag})^{-1} \mathbf{P}_{M,L}^* g. \end{aligned}$$

Again it was used that permutation matrices are unitary matrices. Inverting a block-diagonal matrix is done by inverting the blocks, and so γ^0 can be calculated by

$$\begin{aligned} c_1 &= \mathbf{P}_{M,L}^* g \\ c_2 &= \text{diag}(\mathbf{B}_0^{-1}, \dots, \mathbf{B}_M^{-1}) c_1 \\ \gamma^0 &= \mathbf{P}_{M,L} c_2. \end{aligned}$$

Inverting and applying the blocks can be done in various ways. The total cost is

$$\mathcal{O}(Mb^3) = \mathcal{O}(Lb^2).$$

Perspectives. The algorithms presented in this section speed up the computation of the DGT and the IDGT from $\mathcal{O}(L^2)$ to $\mathcal{O}(NM \log(M)) + \mathcal{O}(LN)$. It depends on M , the number of modulations, which of the terms $\mathcal{O}(NM \log(M))$ and $\mathcal{O}(LN)$ will dominate. If the Gabor frame has many modulations then the first term $\mathcal{O}(NM \log(M))$ will dominate. If the density is kept close to one, then this corresponds to a situation with many modulations and few translations. For other choices of M and N the second term seems more likely to dominate.

Another aspect of these algorithms is that they require much less memory. The original algorithms $c = \mathbf{D}_\gamma^* \Gamma^* f$, $f = \mathbf{D}_g c$ and $\gamma = \mathbf{S}^{-1} g$ required that the matrices $\mathbf{D}_g$, $\mathbf{D}_\gamma$ and $\mathbf{S}$ were formed explicitly. This requires the storage of an $L \times L$ or $L \times MN$ complex matrix.

In an easy-to-use programming language like MATLAB, setting up the matrices $\mathbf{D}_g$ and $\mathbf{S}$ will actually take far longer time than computing the matrix product, so one could be tempted to claim, that one of the main reason for developing the improved algorithms was to avoid the big matrices.

The improved algorithms only require the storage of an $b \times N$ or an $N \times N$ real matrix. The matrices are real (if the window sequence is real) because the matrix $\mathbf{W}_k$ only consists of elements of g .

The article [Stroh98] presents some improvements of these algorithms, especially the algorithm for finding the canonical dual window. I choose to stop here, the presented algorithms are fast enough for my use. The purpose for which I will use the algorithms, solving partial differential equations, only requires a fixed Gabor frame. This means that the canonical dual window can be calculated once before the real computation starts, and then the DGT and the IDGT are computed numerous times. This is the reason why I have not focused more on developing algorithms for inverting $\mathbf{S}$.

There exists other kind of algorithms for these problems than matrix factorization algorithms. Most notable are iterative methods and methods using the Zak-transform. Since I found the algorithms in [Stroh98] adequate for my needs, I have not pursued any of the other kinds of algorithms.

The methods based on the Zak-transform avoids the use of the dual window altogether, and calculates the DGT directly. However, they only work for Gabor frames with a density of 1 (which is then a Gabor basis), and they can be extended to Gabor frames with integer oversampling, that is a Gabor frame with density $\frac{1}{p}$, $p \in \mathbb{N}$. An example of the Zak-transform method is in [Orr93]. An example of an conjugate-gradient and matrix factorization method is in [Qiu98].

5. GABOR FRAMES FOR $L^2(\mathbb{T})$.

5.1. Basic properties. To define a Gabor frame for $L^2(\mathbb{T})$, the translation and modulation operators will have to be defined once again.

Definition 5.1. The translation $\mathcal{T}_y$ and the modulation $\mathcal{M}_m$ operators for $L^2(\mathbb{T})$ are defined by

$$\begin{aligned} (\mathcal{T}_y f)(x) &= f(x + y) \\ (\mathcal{M}_m f)(x) &= e^{-2\pi i m x} f(x) \end{aligned}$$

for all $m \in \mathbb{Z}$ and $y \in \mathbb{T}$.

The symbols for the translation and modulation operators are identical to those of definition 4.10, but it should always be clear from the context which operator a symbol refers to.

The translation and modulation operators for $L^2(\mathbb{T})$ obeys similar relations as those for $\mathbb{C}_p^L$. The modulation operator is not periodic in its parameter, only the translation operator is:

$$(\mathcal{T}_y f)(x) = (\mathcal{T}_{y \bmod 1} f)(x).$$

The inverse and adjoint operators are the opposite action of the operators:

$$\begin{aligned} (\mathcal{T}_y)^{-1} f &= (\mathcal{T}_y)^* f = \mathcal{T}_{-y} f \\ (\mathcal{M}_m)^{-1} f &= (\mathcal{M}_m)^* f = \mathcal{M}_{-m} f. \end{aligned} \tag{5.1}$$

This also means that the translation and modulation operators are unitary operators on $L^2(\mathbb{T})$.

The commutation relation (similar to (4.4)) is slightly different:

$$\begin{aligned}
 (5.2) \quad (\mathcal{M}_m \mathcal{T}_y f)(x) &= e^{-2\pi i m x} f(x+y) \\
 (5.3) &= e^{2\pi i m y} e^{-2\pi i m (x+y)} f(x+y) \\
 (5.4) &= e^{2\pi i m y} (\mathcal{T}_y \mathcal{M}_m f)(x).
 \end{aligned}$$

The commutation relation can be used to commute pairs of the operators:

$$\begin{aligned}
 (\mathcal{M}_a \mathcal{T}_b \mathcal{M}_c \mathcal{T}_d f)(x) &= e^{-2\pi i b c} (\mathcal{M}_a \mathcal{M}_c \mathcal{T}_b \mathcal{T}_d f)(x) \\
 &= e^{-2\pi i b c} (\mathcal{M}_{a+c} \mathcal{T}_{b+d} f)(x) \\
 &= e^{-2\pi i b c} (\mathcal{M}_c \mathcal{M}_a \mathcal{T}_d \mathcal{T}_b f)(x) \\
 (5.5) \quad &= e^{2\pi i (ad-bc)} (\mathcal{M}_c \mathcal{T}_d \mathcal{M}_a \mathcal{T}_b f)(x)
 \end{aligned}$$

These operators are related to the corresponding operators for $\mathbb{C}_p^L$ as defined in definition 4.10 by the following relation:

$$(5.6) \quad \left(\mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g \right) \left(\frac{k}{L} \right) = (\mathcal{M}_{mb} \mathcal{T}_{Na} \{g_k\})_k,$$

with $\{g_k\}_{k=0, \dots, L-1} = g(\frac{k}{L})$ and $Na = L$.

Gabor frames for $L^2(\mathbb{T})$ can now be defined.

Definition 5.2. A *Gabor system* in $L^2(\mathbb{T})$ with window $g \in L^2(\mathbb{T})$ is a family of functions defined by:

$$\mathfrak{G} = G(g, \frac{1}{N}, b) = \left\{ \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g \right\}_{m \in \mathbb{Z}, n=0, \dots, N-1},$$

where $b, N \in \mathbb{N}$.

Each element in the system is called a *Gabor atom*:

$$\mathfrak{G}_{m,n}(x) = \left(\mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g \right)(x) = e^{-2\pi i m b x} g(x + \frac{n}{N}), \quad x \in \mathbb{T}.$$

For Gabor frames for $L^2(\mathbb{T})$ a canonical dual window function exists, just as the canonical, dual sequence for a Gabor frame for $\mathbb{C}_p^L$:

Proposition 5.3. Let $\mathfrak{G}(g, \frac{1}{N}, b)$, $g \in L^2(\mathbb{T})$, $M, b \in \mathbb{N}$ be a Gabor frame for $L^2(\mathbb{T})$. Then there exists a function $\gamma^0 \in L^2(\mathbb{T})$ such that $\mathfrak{H}(\gamma^0, \frac{1}{N}, b)$ is a dual frame of $\mathfrak{G}$.

Proof. The proof follows the exact same steps as the proof for proposition 4.13. The method is to show that both the translation and modulation operator commute with the frame operator. The notable differences are that $\mathcal{M}_{mb}$ is no longer periodic in m and that the sum over m in the definition of the frame operator, corresponding to (4.6), is an infinite sum.

The step corresponding to (4.9) is

$$\begin{aligned}
 & \left(\mathcal{M}_{rb} \mathcal{T}_{\frac{s}{N}} \right)^{-1} S \left(\mathcal{M}_{rb} \mathcal{T}_{\frac{s}{N}} \right) f \\
 &= \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} \left\langle f, \mathcal{M}_{(m-r)b} \mathcal{T}_{\frac{n-s}{N}} g \right\rangle \mathcal{M}_{(m-r)b} \mathcal{T}_{\frac{n-s}{N}} g \\
 &= \sum_{n=0}^{N-1} \sum_{m' \in \mathbb{Z}} \left\langle f, \mathcal{M}_{m'b} \mathcal{T}_{\frac{n}{N}} g \right\rangle \mathcal{M}_{m'b} \mathcal{T}_{\frac{n}{N}} g \\
 &= S f.
 \end{aligned}$$

A substitution of variables $m' = m - r$ and the fact that $\mathcal{T}_{\frac{n-s}{N}}$ is periodic with period N are used when going from line two to line three.

The rest of the proof follows exactly the proof for prop. 4.13. $\square$

The following definitions and theorems are adaptations to $L^2(\mathbb{T})$ of similar statements in [Groch01].

The first statements establish that the pre-frame, analysis and frame operators are bounded when $g \in L^\infty(\mathbb{T})$. They have been transferred from similar statements for $L^2(\mathbb{T})$ presented in [Groch01]. For the $L^2(\mathbb{T})$ theory the Wiener space W is used. A function is in the Wigner space if it is locally in L^∞ and globally in L^1 . Since decay properties are not meaningful for functions defined on $\mathbb{T}$, the circle-equivalent of W is the space $L^\infty(\mathbb{T})$.

Proposition 5.4. *Let $g \in L^\infty(\mathbb{T})$ and let $N, b \in \mathbb{N}$. Then*

$$(5.7) \quad \mathcal{D} \{c_{m,n}\} = \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} c_{m,n} \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g$$

is a bounded operator from $l^2(\mathbb{Z} \times \{0, \dots, N-1\})$ into $L^2(\mathbb{T})$ with operator norm

$$\|\mathcal{D}\|_{\text{op}} \leq \sqrt{N} \|g\|_\infty.$$

Proof. Consider $\{c_{m,n}\} \in l^2(\mathbb{Z} \times \{0, \dots, N-1\})$ with finite support, i.e. only finitely many $c_{m,n}$ are non-zero. Then the trigonometric polynomials

$$m_k(x) = \sum_{n \in \mathbb{Z}} c_{k,n} e^{-2\pi i b n x}, \quad x \in \mathbb{T}$$

$k = 0, \dots, N-1$ are well defined. Since the m_k 's are periodic with period $\frac{1}{b}$ and $\sqrt{b}e^{-2\pi i b n x}$ is an orthonormal basis for $L^2([0; \frac{1}{b}])$ then

$$\|m_k\|_{L^2([0; \frac{1}{b}])}^2 = \frac{1}{b} \sum_{n \in \mathbb{Z}} |c_{k,n}|^2.$$

Define the matrix $\Gamma \in \mathbb{C}^{N,N}$ by

$$\Gamma_{j,k}(x) = \sum_{r=0}^{b-1} \bar{g}(x + \frac{j}{N} + \frac{r}{b}) g(x + \frac{k}{N} + \frac{r}{b}).$$

Since $g \in L^\infty(\mathbb{T})$ then $|g(x)| \leq \|g\|_\infty$, a.e. $x \in \mathbb{T}$ and so

$$\sum_{j=0}^{N-1} |\Gamma_{j,k}(x)| \leq bN \|g\|_\infty^2, \quad \text{a.e. } x \in \mathbb{T}.$$

The same estimate holds with j replaced by k . Then it follows from Schur's test that the norm of Γ is bounded by

$$\|\Gamma\|_{\text{op}} \leq bN \|g\|_\infty^2.$$

Then, for almost all $x \in \mathbb{T}$:

$$(5.8) \quad 0 \leq \sum_{j=0}^{N-1} \sum_{k=0}^{N-1} m_k(x) \overline{m_j(x)} \Gamma_{j,k}(x) \leq bN \|g\|_\infty^2 \sum_{k=0}^{N-1} |m_k(x)|^2.$$

Now define f as

$$f = \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} c_{k,n} \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g = \sum_{k=0}^{N-1} m_k \mathcal{T}_{\frac{n}{N}} g.$$

Using a periodization then the norm of f can be calculated as

$$\begin{aligned} \|f\|_2^2 &= \int_0^1 \left(\sum_{k=0}^{N-1} m_k(x) g(x + \frac{n}{N}) \right) \overline{\left(\sum_{j=0}^{N-1} m_j(x) g(x + \frac{n}{N}) \right)} dx \\ &= \int_0^{\frac{1}{b}} \sum_{k=0}^{N-1} \sum_{j=0}^{N-1} m_k(x) \overline{m_j(x)} \sum_{r=0}^{b-1} g(x + \frac{k}{N} + \frac{r}{b}) \overline{g(x + \frac{j}{N} + \frac{r}{b})} dx \\ &= \int_0^{\frac{1}{b}} \sum_{j=0}^{N-1} \sum_{k=0}^{N-1} m_k(x) \overline{m_j(x)} \Gamma_{j,k}(x) dx \end{aligned}$$

Using (5.8)

$$\begin{aligned} \|f\|^2 &= \int_0^{\frac{1}{b}} \sum_{j=0}^{N-1} \sum_{k=0}^{N-1} m_k(x) \overline{m_j(x)} \Gamma_{j,k}(x) dx \\ &\leq bN \|g\|_\infty^2 \sum_{k=0}^{N-1} \int_0^{\frac{1}{b}} |m_k(x)|^2 dx \\ &= N \|g\|_\infty^2 \sum_{k=0}^{N-1} \sum_{n \in \mathbb{Z}} |c_{k,n}|^2 dx = N \|g\|_\infty^2 \|c_{k,n}\|^2 dx \end{aligned}$$

By the density principle [Groch01, A.1] the result extends to all $\{c_{m,n}\} \in l^2(\mathbb{Z} \times \{0, \dots, N-1\})$ and the results is proven. $\square$

Corollary 5.5. *Let $g \in L^\infty(\mathbb{T})$ and let $N, b \in \mathbb{N}$. Then*

$$\mathcal{C}f = \left\{ \left\langle f, \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g \right\rangle \right\}_{m,n}$$

is a bounded operator from $L^2(\mathbb{T})$ into $l^2(\mathbb{Z} \times \{0, \dots, N-1\})$ with operator norm

$$\|\mathcal{C}\|_{\text{op}} \leq \sqrt{N} \|g\|_\infty,$$

and

$$(5.9) \quad \mathcal{S}f = \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} \left\langle f, \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g \right\rangle \mathcal{M}_{mb} \mathcal{N}_{\frac{n}{N}} g$$

is a bounded operator from $L^2(\mathbb{T})$ into $L^2(\mathbb{T})$ with operator norm

$$\|\mathcal{S}\|_{\text{op}} \leq N \|g\|_\infty^2.$$

Proof. The operators $\mathcal{D}$, $\mathcal{C}$ and $\mathcal{S}$ are the pre-frame-, analysis- and frame- operators for a Gabor system $\mathfrak{G}(g, \frac{1}{N}, b)$. In [Groch01, prop. 5.1.1] it is shown that $\mathcal{C}$ is the adjoint operator of $\mathcal{D}$ and so $\|\mathcal{D}\|_{\text{op}} = \|\mathcal{C}\|_{\text{op}} = \sqrt{\|\mathcal{S}\|_{\text{op}}}$ and from this the result follows. $\square$

Definition 5.6. Let $g \in L^2(\mathbb{T})$ and let $b, N \in \mathbb{N}$. Then the correlation functions of g are defined as

$$G_n(x) = \sum_{k=0}^{N-1} \bar{g}\left(x + \frac{n}{b} + \frac{k}{N}\right) g\left(x + \frac{k}{N}\right), \quad x \in \mathbb{T}.$$

for $n = 0, \dots, b-1$.

Lemma 5.7. Let $g \in L^2(\mathbb{T})$, $b, N \in \mathbb{N}$ and let G_n , $n = 0, \dots, b-1$ be the correlation functions as defined in definition 5.6. Then G_n has the Fourier series

$$G_n(x) = N \sum_{k \in \mathbb{Z}} \left\langle g, \mathcal{M}_{-kN} \mathcal{T}_{\frac{n}{b}} g \right\rangle e^{2\pi i k N x}, \quad \text{a.e. } x \in [0; \frac{1}{N}].$$

Proof. Let

$$h_n(x) = \bar{g}\left(x + \frac{n}{b}\right) g(x) = \left(\mathcal{T}_{\frac{n}{b}} \bar{g} g \right)(x), \quad x \in \mathbb{T}.$$

The correlation function G_n are a periodization of h_n with period $\frac{1}{N}$ and $h_n \in L^2([0; \frac{1}{N}])$ since $g \in L^2(\mathbb{T})$. Using lemma 3.9:

$$\begin{aligned} G_n(x) &= \sum_{k=0}^{N-1} \bar{g}\left(x + \frac{n}{b} + \frac{k}{N}\right) g\left(x + \frac{k}{N}\right) \\ (5.10) \quad &= N \sum_{k \in \mathbb{Z}} \left(\hat{h}_n \right)_{kN} e^{2\pi i k N x}, \quad \text{a.e. } x \in [0; \frac{1}{N}] \end{aligned}$$

The kN 'th Fourier coefficients of h_n are:

$$\begin{aligned} \left(\hat{h}_n \right)_{kN} &= \int_0^1 \bar{g}\left(x + \frac{n}{b}\right) g(x) e^{-2\pi i k N x} dx \\ &= \left\langle g, \mathcal{M}_{-kN} \mathcal{T}_{\frac{n}{b}} g \right\rangle. \end{aligned}$$

Inserting this in (5.10):

$$G_n(x) = N \sum_{k \in \mathbb{Z}} \left\langle g, \mathcal{M}_{-kN} \mathcal{T}_{\frac{n}{b}} g \right\rangle e^{2\pi i k N x}, \quad \text{a.e. } x \in [0; \frac{1}{N}].$$

□

Theorem 5.8. Walnut's representation. Let $g \in L^\infty(\mathbb{T})$ and let $b, N \in \mathbb{N}$. Then the operator $\mathcal{S} : L^2(\mathbb{T}) \rightarrow L^2(\mathbb{T})$

$$\mathcal{S}f = \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} \left\langle f, \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g \right\rangle \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g$$

can be written using definition 5.6:

$$(5.11) \quad \mathcal{S}f = \frac{1}{b} \sum_{k=0}^{b-1} \left(G_k \mathcal{T}_{\frac{k}{b}} f \right).$$

Proof. The fact that $\mathcal{S}$ is bounded on $L^2(\mathbb{T})$ follows from corollary 5.5.

First an important identity for frequency-time shifts is needed:

$$\begin{aligned} \left\langle f, \mathcal{M}_{-mb} \mathcal{T}_{\frac{n}{N}} g \right\rangle &= \int_0^1 f(x) \overline{g}\left(x + \frac{n}{N}\right) e^{-2\pi i m b x} dx \\ &= \int_0^1 \left(f(x) \overline{g}\left(x + \frac{n}{N}\right) \right) e^{-2\pi i m b x} dx \\ &= \left(\mathcal{F} \left(f \mathcal{T}_{\frac{n}{N}} \overline{g} \right) \right)_{mb}. \end{aligned}$$

Line two is recognized as a Fourier coefficient, and from this the result follows.

Consider $\mathcal{C}_g f$, where

$$\mathcal{C}_g f = \left\{ \left\langle f, \mathcal{M}_{-mb} \mathcal{T}_{\frac{n}{N}} g \right\rangle \right\}_{m \in \mathbb{Z}, n=0, \dots, N-1}.$$

Then $\left\{ \left\langle f, \mathcal{M}_{-mb} \mathcal{T}_{\frac{n}{N}} g \right\rangle \right\}_{m \in \mathbb{Z}, n=0, \dots, N-1} \in l^2(\mathbb{Z} \times \{0, \dots, N-1\})$ from corollary 5.5. Therefore, the Fourier series $m_n(x)$

$$(5.12) \quad m_n(x) = \sum_{m \in \mathbb{Z}} \left\langle f, \mathcal{M}_{-mb} \mathcal{T}_{\frac{n}{N}} g \right\rangle e^{2\pi i m b x} = \sum_{m \in \mathbb{Z}} \left(\mathcal{F} \left(f \mathcal{T}_{\frac{n}{N}} \overline{g} \right) \right)_{mb} e^{2\pi i m b x}, \text{ a.e. } x \in \mathbb{T}.$$

are in $L^2(\mathbb{T})$ by Plancherel's theorem, theorem 3.8. Also $m_n(x)$ is periodic with period $\frac{1}{b}$ since $e^{2\pi i m b x}$ is periodic with period $\frac{1}{b}$.

The Fourier series $m_n(x)$ can now be rewritten using lemma 3.9:

$$\begin{aligned} m_n(x) &= \frac{1}{b} \left(b \sum_{m \in \mathbb{Z}} \left(\mathcal{F} \left(f \mathcal{T}_{\frac{n}{N}} \overline{g} \right) \right)_{mb} e^{2\pi i m b x} \right) \\ &= \frac{1}{b} \sum_{k=0}^{b-1} \left(f \mathcal{T}_{\frac{n}{N}} \overline{g} \right) \left(x + \frac{k}{b} \right) \\ &= \sum_{m \in \mathbb{Z}} \left\langle f, \mathcal{M}_{-mb} \mathcal{T}_{\frac{n}{N}} g \right\rangle e^{2\pi i m b x}, \text{ a.e. } x \in [0; \frac{1}{b}]. \end{aligned}$$

The last line is (5.12) again. Substituting this into the definition $\mathcal{S}$

$$\begin{aligned} \mathcal{S}f(x) &= \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} \left\langle f, \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g \right\rangle \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g \\ &= \sum_{n=0}^{N-1} \left(\sum_{m \in \mathbb{Z}} \left\langle f, \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g \right\rangle e^{-2\pi i m b x} \right) g\left(x + \frac{n}{N}\right) \\ &= \sum_{n=0}^{N-1} \left(\frac{1}{b} \sum_{k=0}^{b-1} f\left(x + \frac{k}{b}\right) \overline{g}\left(x + \frac{k}{b} + \frac{n}{N}\right) \right) g\left(x + \frac{n}{N}\right), \text{ a.e. } x \in \mathbb{T}. \end{aligned}$$

Since both sums are finite they can be interchanged

$$\mathcal{S}f(x) = \sum_{k=0}^{b-1} \left(\frac{1}{b} \sum_{n=0}^{N-1} \overline{g}\left(x + \frac{k}{b} + \frac{n}{N}\right) g\left(x + \frac{n}{N}\right) \right) f\left(x + \frac{k}{b}\right), \text{ a.e. } x \in \mathbb{T},$$

or in short notation using definition 5.6 with $\gamma = g$:

$$\mathcal{S}f(x) = \frac{1}{b} \sum_{k=0}^{b-1} G_k(x) f(x + \frac{k}{b}), \text{ a.e. } x \in \mathbb{T}.$$

□

Corollary 5.9. *Let $\mathfrak{G}(g, \frac{1}{N}, b)$, $b, N \in \mathbb{N}$ be a Gabor frame for $L^2(\mathbb{T})$ with a real window function $g \in L^\infty(\mathbb{T})$. Then the canonical dual window $\gamma^0 \in L^2(\mathbb{T})$ is also a real function.*

Proof. This follows from Walnut's representation. When g is real, the correlation functions are also real, and so it follows from Walnut's representation that $\mathcal{S}$ is a real operator. The inverse operator of a real operator is also real, and because

$$\gamma^0 = \mathcal{S}^{-1}g$$

then γ^0 is real. □

Remark 5.10. Theorem 5.8 and corollary 5.9 are the $L^2(\mathbb{T})$ equivalent of proposition 4.17.

5.2. Sampling. The following subsection is a conversion of results in [Jans97] to a periodic setting. It will present two main results: That a Gabor frame from $\mathbb{C}_p^L$ can be obtained from a Gabor frame for $L^2(\mathbb{T})$ by sampling, and that the canonical dual windows for these two frames are closely related.

The following definitions will ensure that the process of sampling is well defined.

Definition 5.11. Let f be a measurable function on $\mathbb{T}$. Then x_0 is a *Lebesgue point* of f if

$$\lim_{\varepsilon \rightarrow 0} \frac{1}{\varepsilon} \int_{-\frac{1}{2}\varepsilon}^{\frac{1}{2}\varepsilon} |f(x + x_0) - f(x_0)| dx = 0.$$

With a little more care Lebesgue points can be defined for any measurable space.

Lemma 5.12. *Let $f \in L^1(\mathbb{T})$. If f is continuous at a point x_0 , then x_0 is a Lebesgue point of f .*

Proof. Since f is continuous at x_0 then for all $\varepsilon > 0$ there exist a δ such that

$$|f(x + x_0) - f(x_0)| < \varepsilon, \forall x \in [-\frac{1}{2}\delta, \frac{1}{2}\delta].$$

Now

$$\frac{1}{\delta} \int_{-\frac{1}{2}\delta}^{\frac{1}{2}\delta} |f(x + x_0) - f(x_0)| dx \leq \sup_{x \in [-\frac{1}{2}\delta, \frac{1}{2}\delta]} |f(x + x_0) - f(x_0)| < \varepsilon.$$

Here the integral is bounded by the maximum value. This shows that when $\varepsilon \rightarrow 0$ then $\frac{1}{\delta} \int_{-\frac{1}{2}\delta}^{\frac{1}{2}\delta} |f(x + x_0) - f(x_0)| dx \rightarrow 0$ and this is exactly the definition of a Lebesgue point. □

Lemma 5.13. *Let $f \in L^1(\mathbb{T})$. If x_0 is a Lebesgue point of f then*

$$f(x_0) = \lim_{\varepsilon \rightarrow 0} \frac{1}{\varepsilon} \int_{-\frac{1}{2}\varepsilon}^{\frac{1}{2}\varepsilon} f(x + x_0) dx.$$

Proof. Consider for $\varepsilon > 0$

$$\begin{aligned} \left| f(x_0) - \frac{1}{\varepsilon} \int_{-\frac{1}{2}\varepsilon}^{\frac{1}{2}\varepsilon} f(x + x_0) dx \right| &= \left| \frac{1}{\varepsilon} \int_{-\frac{1}{2}\varepsilon}^{\frac{1}{2}\varepsilon} f(x_0) dx - \frac{1}{\varepsilon} \int_{-\frac{1}{2}\varepsilon}^{\frac{1}{2}\varepsilon} f(x + x_0) dx \right| \\ &\leq \frac{1}{\varepsilon} \int_{-\frac{1}{2}\varepsilon}^{\frac{1}{2}\varepsilon} |f(x_0) - f(x + x_0)| dx. \end{aligned}$$

Since x_0 is a Lebesgue point then $\frac{1}{\varepsilon} \int_{-\frac{1}{2}\varepsilon}^{\frac{1}{2}\varepsilon} |f(x_0) - f(x + x_0)| dx \rightarrow 0$ as $\varepsilon \rightarrow 0$, and therefore

$$\lim_{\varepsilon \rightarrow 0} \left| f(x_0) - \frac{1}{\varepsilon} \int_{-\frac{1}{2}\varepsilon}^{\frac{1}{2}\varepsilon} f(x + x_0) dx \right| = 0,$$

and this proves the lemma. $\square$

Proposition 5.14. *Let $f, g \in L^2(\mathbb{T})$ such that $f = g$ in L^2 -sense. If x_0 is a Lebesgue point of f and g then $f(x_0) = g(x_0)$.*

Proof. Since $f = g$ in L^2 -sense, then $f(x) = g(x)$ except on a set of measure zero, and so $\int_E f(x) dx = \int_E g(x) dx$ for any set E . Using lemma 5.13 then

$$f(x_0) = \lim_{\varepsilon \rightarrow 0} \frac{1}{\varepsilon} \int_{-\frac{1}{2}\varepsilon}^{\frac{1}{2}\varepsilon} f(x + x_0) dx = \lim_{\varepsilon \rightarrow 0} \frac{1}{\varepsilon} \int_{-\frac{1}{2}\varepsilon}^{\frac{1}{2}\varepsilon} g(x + x_0) dx = g(x_0).$$

$\square$

Definition 5.15. Let f be a measurable function on $\mathbb{T}$ and let $L \in \mathbb{N}$. Then f satisfies *condition R* if:

$$(5.13) \quad \lim_{\varepsilon \rightarrow 0} \sum_{k=0}^{L-1} \frac{1}{\varepsilon} \int_{-\frac{1}{2}\varepsilon}^{\frac{1}{2}\varepsilon} \left| f\left(x + \frac{k}{L}\right) - f\left(\frac{k}{L}\right) \right| dx = 0.$$

Remark 5.16. It shows that f satisfies condition R if and only if all the points $x = \frac{k}{L}$, $k = 0, \dots, L-1$ are Lebesgue points of f . From lemma 5.12 it then follows, that a continuous function satisfies condition R, since all points of a continuous function are Lebesgue points.

With a condition to ensure that sampling is well defined, the first main theorem of this section can be stated.

Theorem 5.17. *Let $M, N, a, b, L \in \mathbb{N}$ such that $Mb = Na = L$. Let*

$$\mathfrak{G} \left(g, \frac{1}{N}, b \right), \quad g \in L^\infty(\mathbb{T})$$

be a Gabor frame for $L^2(\mathbb{T})$ as defined in definition 5.2 with frame bounds $0 < A \leq B < \infty$, and let g satisfy condition R (definition 5.15). Let

$$\mathfrak{H} \left(\{g_k\}, a, \frac{b}{L} \right), \quad \{g_k\} \in \mathbb{C}_p^L$$

be a Gabor system in $\mathbb{C}_p^L$ as defined in definition 4.11 such that $\{g_k\} = g(\frac{k}{L})$.

Then $\mathfrak{H}$ is a frame for $\mathbb{C}_p^L$ with frame bounds AL and BL .

Proof. Let $0 < \varepsilon < \frac{1}{L}$ and let

$$\begin{aligned}\delta^\varepsilon(x) &= \frac{1}{\varepsilon} \chi_{]-\frac{1}{2}\varepsilon; \frac{1}{2}\varepsilon[}(x) \\ \delta_k^\varepsilon(x) &= \frac{1}{\varepsilon} \chi_{]-\frac{1}{2}\varepsilon; \frac{1}{2}\varepsilon[}(x + \frac{k}{L})\end{aligned}$$

for $k \in 0, \dots, L-1$. Let $\{c_k\}, \{d_k\} \in \mathbb{C}_p^L$ and set

$$\begin{aligned}f^\varepsilon &= \sum_{r=0}^{L-1} c_r \delta_r^\varepsilon \\ h^\varepsilon &= \sum_{s=0}^{L-1} d_s \delta_s^\varepsilon.\end{aligned}$$

Then $f^\varepsilon \in L^2(\mathbb{T})$ since the sum over r is finite and

$$\begin{aligned}\|f^\varepsilon\|_2^2 &= \int_0^1 \left| \sum_{r=0}^{L-1} c_r \delta_r^\varepsilon \right|^2 dx \\ &= \sum_{k=0}^{L-1} \int_0^1 |c_r \delta_r^\varepsilon|^2 dx \\ (5.14) \quad &= \frac{1}{\varepsilon} \sum_{k=0}^{L-1} |c_r|^2 dx = \frac{1}{\varepsilon} \|c\|_2^2.\end{aligned}$$

Line two follows from the fact that the δ_r^ε have disjoint support. Define

$$\Upsilon_{m,n}(\varepsilon) = \varepsilon \sum_{l \in \mathbb{Z}} \left\langle f^\varepsilon, \mathcal{M}_{-lL+mb} \mathcal{T}_{\frac{n}{N}} g \right\rangle \left\langle \mathcal{M}_{-lL+mb} \mathcal{T}_{\frac{n}{N}} g, h^\varepsilon \right\rangle,$$

with $m = 0, \dots, M-1$, $n = 0, \dots, N-1$ and $\varepsilon \in [-\frac{1}{2L}; \frac{1}{2L}]$. Since $\mathfrak{G}$ is a frame then $\left\{ \left\langle f^\varepsilon, \mathcal{M}_{-lL+mb} \mathcal{T}_{\frac{n}{N}} g \right\rangle \right\}_l \in l^2(\mathbb{Z})$ by the frame inequality. Then the series defining Υ is in $l^1(\mathbb{Z})$ from Cauchy-Schwartz inequality. Therefore $\Upsilon_{m,n}(\varepsilon)$ is well-defined.

It follows by rearranging the summations that

$$(5.15) \quad \sum_{n=0}^{N-1} \sum_{m=0}^{M-1} \Upsilon_{m,n}(\varepsilon) = \varepsilon \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} \left\langle f^\varepsilon, \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g \right\rangle \left\langle \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g, h^\varepsilon \right\rangle.$$

The function $\Upsilon_{m,n}$ can be written as

$$\begin{aligned}\Upsilon_{m,n}(\varepsilon) &= \varepsilon \sum_{l \in \mathbb{Z}} \left\langle \sum_{r=0}^{L-1} c_r \delta_r^\varepsilon, \mathcal{M}_{-lL+mb} \mathcal{T}_{\frac{n}{N}} g \right\rangle \left\langle \mathcal{M}_{-lL+mb} \mathcal{T}_{\frac{n}{N}} g, \sum_{s=0}^{L-1} d_s \delta_s^\varepsilon \right\rangle \\ (5.16) \quad &= \sum_{r=0}^{L-1} \sum_{s=0}^{L-1} c_r \bar{d}_s \varepsilon \sum_{l \in \mathbb{Z}} \left\langle \delta_r^\varepsilon, \mathcal{M}_{-lL+mb} \mathcal{T}_{\frac{n}{N}} g \right\rangle \left\langle \mathcal{M}_{-lL+mb} \mathcal{T}_{\frac{n}{N}} g, \delta_s^\varepsilon \right\rangle.\end{aligned}$$

The summations over r and s can be pulled outside the inner product as they are both finite.

Consider the $\frac{1}{L}$ -periodic functions

$$\alpha_t(u) = \sum_{j=0}^{L-1} \delta_t^\varepsilon(u - \frac{j}{L}) \overline{\mathcal{M}_{mb} \mathcal{T}_N^n g}(u - \frac{j}{L}), \quad t = 0, \dots, L-1, \quad n = 0, \dots, N-1.$$

Using lemma 3.9 then α_t has the expansion

$$\alpha_t(u) = L \sum_{k \in \mathbb{Z}} \left(\mathcal{F} \left(\delta_t^\varepsilon \overline{\mathcal{M}_{mb} \mathcal{T}_N^n g} \right) \right)_{kL} e^{2\pi i k L x}, \quad a.e. x \in [0; \frac{1}{L}],$$

with

$$\begin{aligned} \left(\mathcal{F} \left(\delta_t^\varepsilon \overline{\mathcal{M}_{mb} \mathcal{T}_N^n g} \right) \right)_{kL} &= \int_0^1 \delta_t^\varepsilon(x) \overline{\mathcal{M}_{mb} \mathcal{T}_N^n g}(x) e^{-2\pi i k L x} dx \\ &= \left\langle \delta_t^\varepsilon, \mathcal{M}_{-kL+mb} \mathcal{T}_N^n g \right\rangle_{L^2(\mathbb{T})}. \end{aligned}$$

Since the complex exponentials $\sqrt{L} e^{2\pi i k L x}$, $k \in \mathbb{Z}$ form an orthonormal basis for $L^2([-\frac{1}{2L}; \frac{1}{2L}])$ then Parseval's equation can be used on the expansion of α_t :

$$\begin{aligned} &\varepsilon \sum_{l \in \mathbb{Z}} \left\langle \delta_r^\varepsilon, \mathcal{M}_{-lL+mb} \mathcal{T}_N^n g \right\rangle_{L^2(\mathbb{T})} \left\langle \mathcal{M}_{-lL+mb} \mathcal{T}_N^n g, \delta_s^\varepsilon \right\rangle_{L^2(\mathbb{T})} \\ &= \varepsilon \frac{1}{L} \int_{-\frac{1}{2L}}^{\frac{1}{2L}} \alpha_r(x) \bar{\alpha}_s(x) dx \\ &= \varepsilon \frac{1}{L} \int_{-\frac{1}{2L}}^{\frac{1}{2L}} \left(\sum_{k \in \mathbb{Z}} \delta_r^\varepsilon(x - \frac{k}{L}) \overline{\mathcal{M}_{mb} \mathcal{T}_N^n g}(x - \frac{k}{L}) \right) \overline{\left(\sum_{l \in \mathbb{Z}} \delta_s^\varepsilon(x - \frac{l}{L}) \overline{\mathcal{M}_{mb} \mathcal{T}_N^n g}(x - \frac{l}{L}) \right)} dx \\ &= \frac{1}{\varepsilon} \frac{1}{L} \int_{-\frac{1}{2\varepsilon}}^{\frac{1}{2\varepsilon}} \overline{\mathcal{M}_{mb} \mathcal{T}_N^n g}(x + \frac{r}{L}) \mathcal{M}_{mb} \mathcal{T}_N^n g(x + \frac{s}{L}) dx \end{aligned}$$

The summations over k and l have been removed because

$$\delta_r^\varepsilon(u + \frac{k}{L}) = \begin{cases} \frac{1}{\varepsilon} & \text{if } -\frac{k}{L} + \frac{r}{L} \in [-\frac{1}{\varepsilon}; \frac{1}{\varepsilon}] \\ 0 & \text{otherwise.} \end{cases}$$

and since $0 < \varepsilon < \frac{1}{L}$ the terms in the sums are only non-zero when $k = r$ and $l = s$. Similarly, the range over which to integrate is restricted to the support of the indicator functions, which then can be removed.

Since $\frac{r}{L}$ and $\frac{s}{L}$ are Lebesgue points of $\mathcal{M}_{mb} \mathcal{T}_N^n g$ and because $g \in L^\infty(\mathbb{T})$ is essentially bounded then

$$\begin{aligned} &\lim_{\varepsilon \rightarrow 0} \frac{1}{\varepsilon} \frac{1}{L} \int_{-\frac{1}{2\varepsilon}}^{\frac{1}{2\varepsilon}} \overline{\mathcal{M}_{mb} \mathcal{T}_N^n g}(x + \frac{r}{L}) \mathcal{M}_{mb} \mathcal{T}_N^n g(x + \frac{s}{L}) dx \\ &= \frac{1}{L} \overline{\mathcal{M}_{mb} \mathcal{T}_N^n g}(\frac{r}{L}) \mathcal{M}_{mb} \mathcal{T}_N^n g(\frac{s}{L}). \end{aligned}$$

This can be inserted into (5.16):

$$\begin{aligned}
\lim_{\varepsilon \rightarrow 0} \Upsilon_{m,n}(\varepsilon) &= \lim_{\varepsilon \rightarrow 0} \sum_{r=0}^{L-1} \sum_{s=0}^{L-1} c_r \bar{d}_s \varepsilon \sum_{l \in \mathbb{Z}} \left\langle \delta_r^\varepsilon, \mathcal{M}_{-lL+mb} \mathcal{T}_{\frac{n}{N}} g \right\rangle \left\langle \mathcal{M}_{-lL+mb} \mathcal{T}_{\frac{n}{N}} g, \delta_s^\varepsilon \right\rangle \\
&= \frac{1}{L} \sum_{r=0}^{L-1} \sum_{s=0}^{L-1} c_r \bar{d}_s \overline{\mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g} \left(\frac{r}{L} \right) \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g \left(\frac{s}{L} \right)
\end{aligned}$$

From this

$$\begin{aligned}
&\lim_{\varepsilon \rightarrow 0} \varepsilon \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} \left\langle f^\varepsilon, \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g \right\rangle_{L^2(\mathbb{T})} \left\langle \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g, h^\varepsilon \right\rangle_{L^2(\mathbb{T})} \\
&= \lim_{\varepsilon \rightarrow 0} \sum_{n=0}^{N-1} \sum_{m=0}^{M-1} \Upsilon_{m,n}(\varepsilon) \\
(5.17) \quad &= \sum_{n=0}^{N-1} \sum_{m=0}^{M-1} \lim_{\varepsilon \rightarrow 0} \Upsilon_{m,n}(\varepsilon) \\
&= \frac{1}{L} \sum_{n=0}^{N-1} \sum_{m=0}^{M-1} \sum_{r=0}^{L-1} \sum_{s=0}^{L-1} c_r \bar{d}_s \overline{\mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g} \left(\frac{r}{L} \right) \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g \left(\frac{s}{L} \right) \\
(5.18) \quad &= \frac{1}{L} \sum_{n=0}^{N-1} \sum_{m=0}^{M-1} \langle \{c_k\}, \mathcal{M}_{mb} \mathcal{T}_{na} \{g_k\} \rangle_{\mathbb{C}_p^L} \langle \mathcal{M}_{mb} \mathcal{T}_{na} \{g_k\}, \{d_k\} \rangle_{\mathbb{C}_p^L}.
\end{aligned}$$

Second line comes from (5.15). The translation and modulation operators in the last line are the discrete versions. Since $\|f^\varepsilon\|^2 = \frac{1}{\varepsilon} \|c\|^2$ from (5.14) and

$$\varepsilon A \|f^\varepsilon\|^2 = A \|c\|^2 \leq \varepsilon \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} \left| \left\langle f^\varepsilon, \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g \right\rangle_{L^2(\mathbb{T})} \right|^2 \leq \varepsilon B \|f^\varepsilon\|^2 = B \|c\|^2$$

from the frame condition (4.1), then

$$AL \|c\|^2 \leq \sum_{n=0}^{N-1} \sum_{m=0}^{M-1} \left| \langle \{c_k\}, \mathcal{M}_{mb} \mathcal{T}_{na} \{g_k\} \rangle_{\mathbb{C}_p^L} \right|^2 \leq BL \|c\|^2,$$

from (5.18). Since $\{c_k\}$ was arbitrarily chosen, the proof is complete. $\square$

Remark 5.18. In (5.17) the limit is pulled inside the two summations. This is not possible in the similar proof for non-periodic frames, so the proof of [Jans97, prop. 2] deviates at this point.

To prove the relationship between the canonical duals of the frames $\mathfrak{G}$ and $\mathfrak{H}$ from theorem 5.17 another representation of the Gabor frame operator is needed. This is Janssens representation. To make sure it is convergent an additional condition is needed.

Definition 5.19. Let $g \in L^2(\mathbb{T})$ and $N, b \in \mathbb{N}$. The g satisfies *condition A* if:

$$(5.19) \quad \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} \left| \left\langle g, \mathcal{M}_{mN} \mathcal{T}_{\frac{n}{b}} g \right\rangle \right| < \infty.$$

The corresponding condition A for $L^2(\mathbb{R})$ is delicate to satisfy [Groch01, p. 132], but for $L^2(\mathbb{T})$ the situation is simpler.

Proposition 5.20. *Let $g \in C^1(\mathbb{T})$, $N, b \in \mathbb{N}$. Then g satisfies condition A.*

Proof. Consider

$$\begin{aligned} \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} \left| \left\langle g, \mathcal{M}_{mN} \mathcal{T}_{\frac{n}{b}} g \right\rangle \right| &= \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} \left| \int_0^1 g\left(x + \frac{n}{b}\right) \bar{g}\left(x + \frac{n}{b}\right) e^{2\pi i m N x} dx \right| \\ &= \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} \left| \left(\mathcal{F} \left(\mathcal{T}_{\frac{n}{b}} (g\bar{g}) \right) \right)_{-mN} \right| \end{aligned}$$

The term in line two is bounded if

$$\mathcal{F} \left(\mathcal{T}_{\frac{n}{b}} (g\bar{g}) \right) \in l^1(\mathbb{Z}), \forall n \in \{0, \dots, N-1\}.$$

From theorem 3.11 then this is satisfied if $\mathcal{T}_{\frac{n}{b}} (g\bar{g})$ is continuous with a bounded and piecewise monotone derivative. This is easily satisfied if requiring $g \in C^1(\mathbb{T})$. $\square$

The previous proposition is by no means meant to be exhaustive, it is only meant to provide a simple, easy to check condition to be used for the construction of Gabor frames.

With this condition then Janssens representation can be defined.

Theorem 5.21. Janssen's representation. *Let $\mathfrak{G}(g, \frac{1}{N}, b)$, $g \in L^2(\mathbb{T})$, $N, b \in \mathbb{N}$ be a Gabor frame for $L^2(\mathbb{T})$ and let g satisfy condition A. Then the Gabor frame operator*

$$\mathcal{S}f = \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} \left\langle f, \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g \right\rangle \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g$$

can be written as

$$(5.20) \quad \mathcal{S}f = \frac{N}{b} \sum_{s=0}^{b-1} \sum_{r \in \mathbb{Z}} \left\langle g, \mathcal{M}_{rN} \mathcal{T}_{\frac{s}{b}} g \right\rangle \mathcal{M}_{rN} \mathcal{T}_{\frac{s}{b}} f$$

with unconditional convergence.

Proof. Define the operator $\tilde{\mathcal{S}}$ by

$$\tilde{\mathcal{S}}f = \frac{N}{b} \sum_{s=0}^{b-1} \sum_{r \in \mathbb{Z}} \left\langle g, \mathcal{M}_{-rN} \mathcal{T}_{\frac{s}{b}} g \right\rangle \mathcal{M}_{-rN} \mathcal{T}_{\frac{s}{b}} f.$$

The operator $\tilde{\mathcal{S}}$ is equal to (5.20), except for a trivial change in the sign of r .

It must be shown that $\tilde{\mathcal{S}}$ equals $\mathcal{S}$. Since g satisfy condition A, the series defining $\tilde{\mathcal{S}}$ is absolutely convergent and therefore independent of the order of summation.

$$\begin{aligned} (\tilde{\mathcal{S}}f)(x) &= \frac{1}{b} \sum_{s=0}^{b-1} \left(N \sum_{r \in \mathbb{Z}} \left\langle g, \mathcal{M}_{-rN} \mathcal{T}_{\frac{s}{b}} g \right\rangle e^{2\pi i r N x} \right) \left(\mathcal{T}_{\frac{s}{b}} f \right)(x) \\ &= \frac{1}{b} \sum_{s=0}^{b-1} \left(G_s \mathcal{T}_{\frac{s}{b}} f \right)(x), \text{ a.e. } x \in \mathbb{T}. \end{aligned}$$

The last line comes from the Fourier series expansion of the correlation functions, lemma 5.7. The term in the last line is Walnut's representation of the frame operator, theorem 5.8, and from this the result follows. $\square$

Remark 5.22. The proof of the preceding theorem is similar to, but conceptually more difficult than the proof of theorem 4.17, because Janssens representation of the frame operator for a Gabor frame for $\mathbb{C}_p^L$ does not require a condition A to converge.

Before the second main theorem can be proven, an additional lemma and proposition is needed.

Lemma 5.23. *Let $f \in L^2(\mathbb{T})$ satisfy condition R and assume that the Gabor system in $L^2(\mathbb{T})$: $\mathfrak{G}(g, \frac{1}{N}, b)$, $g \in L^\infty(\mathbb{T})$, $N, b \in \mathbb{N}$ is a Bessel sequence with Bessel bound B . Let $\{c_{k,l}\} \in l^1(\mathbb{Z} \times 0, \dots, N-1)$.*

Then

$$\varphi = \sum_{l=0}^{N-1} \sum_{k \in \mathbb{Z}} c_{k,l} \mathcal{M}_{kN} \mathcal{T}_{\frac{l}{b}} g$$

satisfies condition R and the Gabor system in $L^2(\mathbb{T})$: $\mathfrak{H}(\varphi, \frac{1}{N}, b)$ is a Bessel sequence with Bessel bound

$$\tilde{B} = B \sum_{l=0}^{N-1} \sum_{k \in \mathbb{Z}} |c_{k,l}|^2.$$

Proof. The function φ is well-defined since $\{c_{k,l}\} \in l^1(\mathbb{Z} \times 0, \dots, N-1)$.

For $f \in L^2(\mathbb{T})$ there holds

$$\begin{aligned} & \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} \left| \left\langle f, \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} \varphi \right\rangle \right|^2 \\ &= \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} \left| \left\langle f, \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} \left(\sum_{l=0}^{N-1} \sum_{k \in \mathbb{Z}} c_{k,l} \mathcal{M}_{kN} \mathcal{T}_{\frac{l}{b}} g \right) \right\rangle \right|^2 \\ &= \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} \left| \sum_{l=0}^{N-1} \sum_{k \in \mathbb{Z}} \left\langle f, c_{k,l} \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} \mathcal{M}_{kN} \mathcal{T}_{\frac{l}{b}} g \right\rangle \right|^2 \\ &= \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} \left| \sum_{l=0}^{N-1} \sum_{k \in \mathbb{Z}} \tilde{c}_{k,l} e^{-2\pi i(ml-nk)} e^{-2\pi ikl \frac{N}{b}} \left\langle \mathcal{M}_{-kN} \mathcal{T}_{-\frac{l}{b}} f, \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g \right\rangle \right|^2 \\ &\leq \left(\sum_{l=0}^{N-1} \sum_{k \in \mathbb{Z}} |c_{k,l}| \left(\sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} \left| \left\langle \mathcal{M}_{-kN} \mathcal{T}_{-\frac{l}{b}} f, \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g \right\rangle \right|^2 \right)^{\frac{1}{2}} \right)^2 \\ &\leq \left(\sum_{l=0}^{N-1} \sum_{k \in \mathbb{Z}} |c_{k,l}| \left(B \left\| \mathcal{M}_{-kN} \mathcal{T}_{-\frac{l}{b}} f \right\|^2 \right)^{\frac{1}{2}} \right)^2 = B \|f\|^2 \left(\sum_{l=0}^{N-1} \sum_{k \in \mathbb{Z}} |c_{k,l}| \right)^2 \end{aligned}$$

In the third line the summation can be pulled outside the inner product since φ is well-defined and the inner product is continuous. In the fourth line the commutation relations (5.5) and (5.4) and the adjoint operators of the translation and modulation operators (5.1) are used. In the fifth line the triangle inequality

$$\left\| \sum_{k,l} c_{kl} \psi_{k,l} \right\| \leq \sum |c_{k,l}| \|\psi_{k,l}\|$$

for $\psi_{k,l} \in l^2(\mathbb{Z} \times 0, \dots, N-1)$ is used. In the last line the frame condition (4.1) is used and the fact that the translation and modulation operators are unitary operators on $L^2(\mathbb{T})$.

This calculation shows that $\mathfrak{H}(\varphi, \frac{1}{N}, b)$ has a finite upper frame bound $\tilde{B}$. To show that φ satisfies condition R then consider:

$$\begin{aligned}
 & \sum_{j=0}^{N-1} \frac{1}{\varepsilon} \int_{\frac{1}{2}\varepsilon}^{\frac{1}{2}\varepsilon} |\varphi(j+x) - \varphi(j)|^2 dx \\
 &= \sum_{j=0}^{N-1} \frac{1}{\varepsilon} \int_{\frac{1}{2}\varepsilon}^{\frac{1}{2}\varepsilon} \left| \sum_{l=0}^{N-1} \sum_{k \in \mathbb{Z}} c_{k,l} \left(\mathcal{M}_{kN} \mathcal{T}_{\frac{l}{b}} g(j+x) - \mathcal{M}_{kN} \mathcal{T}_{\frac{l}{b}} g(j) \right) \right|^2 dx \\
 (5.21) \quad &\leq \left(\sum_{l=0}^{N-1} \sum_{k \in \mathbb{Z}} |c_{k,l}| \left(\sum_{j=0}^{N-1} \frac{1}{\varepsilon} \int_{\frac{1}{2}\varepsilon}^{\frac{1}{2}\varepsilon} \left| \mathcal{M}_{kN} \mathcal{T}_{\frac{l}{b}} g(j+x) - \mathcal{M}_{kN} \mathcal{T}_{\frac{l}{b}} g(j) \right|^2 dx \right)^{\frac{1}{2}} \right)^2.
 \end{aligned}$$

The last line follows from the triangle inequality for $L^2(0, \dots, N-1 \times [-\frac{1}{2}\varepsilon; \frac{1}{2}\varepsilon])$. The part inside the integral can be estimated

$$\begin{aligned}
 & \left| \mathcal{M}_{kN} \mathcal{T}_{\frac{l}{b}} g(j+x) - \mathcal{M}_{kN} \mathcal{T}_{\frac{l}{b}} g(j) \right| \\
 &\leq \left| e^{-2\pi i x N} g(j+x+\frac{l}{b}) - g(j+\frac{l}{b}) \right| \\
 &\leq \left| (e^{-2\pi i x N} - 1) g(j+x+\frac{l}{b}) \right| + \left| g(j+x+\frac{l}{b}) - g(j+\frac{l}{b}) \right|.
 \end{aligned}$$

In the second line the term $e^{-2\pi i j N}$ is pulled outside the absolute signs and bounded by one.

From this

$$\begin{aligned}
 & \sum_{j=0}^{N-1} \frac{1}{\varepsilon} \int_{\frac{1}{2}\varepsilon}^{\frac{1}{2}\varepsilon} \left| \mathcal{M}_{kN} \mathcal{T}_{\frac{l}{b}} g(j+x) - \mathcal{M}_{kN} \mathcal{T}_{\frac{l}{b}} g(j) \right|^2 dx \\
 &\leq \int_{\frac{1}{2}\varepsilon}^{\frac{1}{2}\varepsilon} \frac{|e^{-2\pi i x N} - 1|^2}{\varepsilon} \sum_{j=0}^{N-1} \left| g(j+x+\frac{l}{b}) \right|^2 dx \\
 (5.22) \quad &+ \sum_{j=0}^{N-1} \frac{1}{\varepsilon} \int_{\frac{1}{2}\varepsilon}^{\frac{1}{2}\varepsilon} \left| g(j+x+\frac{l}{b}) - g(j+\frac{l}{b}) \right|^2 dx.
 \end{aligned}$$

Since $\frac{|e^{-2\pi i x N} - 1|^2}{\varepsilon} \rightarrow 0$ as $\varepsilon \rightarrow 0$ and since $\sum_{j=0}^{N-1} \left| g(j+x+\frac{l}{b}) \right|^2 \leq N \|g\|_\infty^2$ it shows that the term in line two goes to zero as $\varepsilon \rightarrow 0$. The term in line three goes to zero because g satisfies condition R. Both of the terms are bounded in k and l since g is essentially bounded. Since $\{c_{k,l}\} \in l^1(\mathbb{Z} \times 0, \dots, N-1)$ then it follows from (5.21) that φ satisfies condition R. $\square$

Proposition 5.24. *Let $\mathfrak{G}(g, \frac{1}{N}, b)$ and $\mathfrak{H}(h, \frac{1}{N}, b)$, $g \in L^\infty(\mathbb{T})$, $h \in L^2(\mathbb{T})$, $N, b \in \mathbb{N}$ be Gabor frames for $L^2(\mathbb{T})$. Assume that g satisfy condition R and condition A, definition 5.19. Let $L = Ma = Nb$ with $L, M, a \in \mathbb{N}$. Let $\mathcal{S}_g$ be the frame operator corresponding to $\mathfrak{G}$ and let $\mathcal{S}_{\{g\}}$ be the frame operator corresponding to the Gabor frame for $\mathbb{C}_p^L: \mathfrak{J}(\{g(\frac{k}{L})\}_k, a, \frac{b}{L})$.*

Then $\mathcal{S}h$ satisfies condition R, and the Gabor system in $L^2(\mathbb{T})$: $\mathfrak{K}(\mathcal{S}_g h, a, \frac{b}{L})$ is a Bessel sequence, and

$$(\mathcal{S}_g h) \left(\frac{k}{L} \right) = \frac{1}{L} S_{\{g\}} \left\{ h \left(\frac{k}{L} \right) \right\}_k, \quad \forall k \in \{0, \dots, L-1\}.$$

Proof. The claim that $\mathfrak{J}$ is a Gabor frame for $\mathbb{C}_p^L$ follows from theorem 5.17.

Since g satisfies condition A then

$$\sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} \left| \left\langle g, \mathcal{M}_{mN} \mathcal{T}_{\frac{n}{b}} g \right\rangle \right| < \infty$$

which means that $\{c_{m,n}\} = \left\{ \left\langle g, \mathcal{M}_{mN} \mathcal{T}_{\frac{n}{b}} g \right\rangle \right\}_{m,n} \in l^1(\mathbb{Z} \times \{0, \dots, N-1\})$. Using lemma 5.23 on h with the sequence $\{c_{m,n}\}$ shows that $\mathcal{S}h$ satisfies condition R and that $\mathfrak{K}$ is a Bessel sequence.

Since g satisfies condition A, then $\mathcal{S}$ can be expressed using Janssens representation, and since $\mathcal{S}h$ satisfies condition R, then it can be sampled:

$$\begin{aligned} \mathcal{S}h \left(\frac{k}{L} \right) &= \frac{N}{b} \sum_{n=0}^{b-1} \sum_{m \in \mathbb{Z}} \left\langle g, \mathcal{M}_{mN} \mathcal{T}_{\frac{n}{b}} g \right\rangle \mathcal{M}_{mN} \mathcal{T}_{\frac{n}{b}} h \left(\frac{k}{L} \right) \\ &= \frac{N}{b} \sum_{n=0}^{b-1} \sum_{m=0}^{a-1} \sum_{l \in \mathbb{Z}} \left\langle g, \mathcal{M}_{mN+lL} \mathcal{T}_{\frac{n}{b}} g \right\rangle \mathcal{M}_{mN+lL} \mathcal{T}_{\frac{n}{b}} h \left(\frac{k}{L} \right) \\ (5.23) \quad &= \frac{N}{b} \sum_{n=0}^{b-1} \sum_{m=0}^{a-1} \sum_{l \in \mathbb{Z}} \left\langle g, \mathcal{M}_{mN+lL} \mathcal{T}_{\frac{n}{b}} g \right\rangle \mathcal{M}_{mN} \mathcal{T}_{\frac{n}{b}} h \left(\frac{k}{L} \right). \end{aligned}$$

In line two the summation over m is replaced by summations over m and l . The simplification in the modulation operator in line three is possible because of the sampling.

Next consider the $\frac{1}{L}$ periodic function ψ defined by

$$\psi(x) = \sum_{k=0}^{L-1} g \left(\frac{k}{L} + x \right) \bar{g} \left(\frac{k}{L} + x + \frac{n}{b} \right) e^{-2\pi i m N \left(\frac{k}{L} + x \right)}.$$

for $m = 0, \dots, M-1$ and $n = 0, \dots, N-1$ Then

$$\begin{aligned} &\frac{1}{\varepsilon} \int_{-\frac{1}{2}\varepsilon}^{\frac{1}{2}\varepsilon} |\psi(x) - \psi(0)| dx \\ &\leq \frac{1}{\varepsilon} \sum_{k=0}^{L-1} \int_{-\frac{1}{2}\varepsilon}^{\frac{1}{2}\varepsilon} \left| g \left(\frac{k}{L} + x \right) \bar{g} \left(\frac{k}{L} + x + \frac{n}{b} \right) e^{-2\pi i m N x} - g \left(\frac{k}{L} \right) \bar{g} \left(\frac{k}{L} + \frac{n}{b} \right) \right| dx \\ &\leq \frac{1}{\varepsilon} \sum_{k=0}^{L-1} \int_{-\frac{1}{2}\varepsilon}^{\frac{1}{2}\varepsilon} \left| g \left(\frac{k}{L} + x \right) \bar{g} \left(\frac{k}{L} + x + \frac{n}{b} \right) \right| |e^{-2\pi i m N x} - 1| dx \\ &\quad + \frac{1}{\varepsilon} \sum_{k=0}^{L-1} \int_{-\frac{1}{2}\varepsilon}^{\frac{1}{2}\varepsilon} \left| g \left(\frac{k}{L} + x \right) \bar{g} \left(\frac{k}{L} + x + \frac{n}{b} \right) - g \left(\frac{k}{L} \right) \bar{g} \left(\frac{k}{L} + \frac{n}{b} \right) \right| dx. \end{aligned}$$

The last two lines go to zero as $\varepsilon \rightarrow 0$ because of the exact same reasoning as for (5.22). This shows that

$$\lim_{\varepsilon \rightarrow 0} \frac{1}{\varepsilon} \int_{-\frac{1}{2}\varepsilon}^{\frac{1}{2}\varepsilon} |\psi(x) - \psi(0)| dx = 0$$

which means that 0 is a Lebesgue point of ψ . A Fourier expansion for ψ can be found from lemma 3.9:

$$\psi(x) = L \sum_{l \in \mathbb{Z}} \left(\mathcal{F} \left(g(x) \bar{g}(x + \frac{n}{b}) e^{-2\pi i m N x} \right) \right)_{lL} e^{2\pi i l L x}, \text{ a.e. } x \in [0; \frac{1}{L}],$$

where

$$\begin{aligned} \left(\mathcal{F} \left(g(x) \bar{g}(x + \frac{n}{b}) e^{-2\pi i m N x} \right) \right)_{lL} &= \int_0^1 g(x) \bar{g}(x + \frac{n}{b}) e^{-2\pi i m N x} e^{-2\pi i l L x} dx \\ &= \left\langle g, \mathcal{M}_{-mN-lL} \mathcal{T}_{\frac{n}{b}} g \right\rangle, \end{aligned}$$

so

$$(5.24) \quad \psi(x) = L \sum_{l \in \mathbb{Z}} \left\langle g, \mathcal{M}_{-mN-lL} \mathcal{T}_{\frac{n}{b}} g \right\rangle e^{2\pi i l L x}, \text{ a.e. } x \in [0; \frac{1}{L}].$$

The right hand side of (5.24) is a continuous function, since $\left\{ \left\langle g, \mathcal{M}_{-mN-lL} \mathcal{T}_{\frac{n}{b}} g \right\rangle_l \right\} \in l^1(\mathbb{Z})$, [BuNe71, prop. 4.1.5]. Since the right hand side is continuous, it coincides with ψ at all Lebesgue points. This follows from lemma 5.12 and proposition 5.14. Therefore it also coincides for $x = 0$. From this

$$\begin{aligned} \psi(0) &= L \sum_{l \in \mathbb{Z}} \left\langle g, \mathcal{M}_{-mN-lL} \mathcal{T}_{\frac{n}{b}} g \right\rangle = \sum_{k=0}^{L-1} g\left(\frac{k}{L}\right) \bar{g}\left(\frac{k}{L} + \frac{n}{b}\right) e^{-2\pi i m N \frac{k}{L}} \\ (5.25) \quad &= \left\langle \{g_k\}, \mathcal{M}_{-mN} \mathcal{T}_{\frac{n}{b}} \{g_k\} \right\rangle. \end{aligned}$$

This can be inserted in (5.23):

$$\begin{aligned} \mathcal{S}h\left(\frac{k}{L}\right) &= \frac{N}{b} \sum_{n=0}^{b-1} \sum_{m=0}^{a-1} \sum_{l \in \mathbb{Z}} \left\langle g, \mathcal{M}_{mN+lL} \mathcal{T}_{\frac{n}{b}} g \right\rangle \mathcal{M}_{mN} \mathcal{T}_{\frac{n}{b}} h\left(\frac{k}{L}\right) \\ &= \frac{1}{L} \left(\frac{N}{b} \sum_{n=0}^{b-1} \sum_{m=0}^{a-1} \langle \{g_k\}, \mathcal{M}_{mN} \mathcal{T}_{\frac{n}{b}} \{g_k\} \rangle \mathcal{M}_{mN} \mathcal{T}_{\frac{n}{b}} \{h_k\} \right)_k \\ &= \frac{1}{L} (\mathcal{S}_{\{g\}} \{h_k\})_k. \end{aligned}$$

Line two comes from inserting (5.25), and using (5.6) to express the term using the translation and modulation operators for $\mathbb{C}_p^L$. The second line is Janssens representation of the frame operator for $\mathfrak{J}$ (4.17), and so the result follows. $\square$

Remark 5.25. The factor $\frac{1}{L}$ between $\mathcal{S}h$ and $(\mathcal{S}_{\{g\}} \{h_k\})_k$ does not appear in the similar statement for non-periodic spaces $L^2(\mathbb{R})$ and $l^2(\mathbb{Z})$ presented in [Jans97, proposition 3]. It is a result of the need to normalize the sampled functions. If a function $f \in L^2(\mathbb{T})$ has $\|f\| = 1$ then it does not hold that $\|\{f_k\}\| = 1$. The same normalization constant appears in the definitions of the DFT and the IDFT, definition 3.1 and corollary 3.6 and in proposition 4.15.

To prove the main theorem and additional result is needed, which I will state as a conjecture.

Conjecture 5.26. *Let $\mathfrak{G}(g, \frac{1}{N}, b)$, $g \in L^2(\mathbb{T})$, $N, b \in \mathbb{N}$ be a Gabor frame for $L^2(\mathbb{T})$ and assume that g satisfies condition A. Then the canonical dual window γ^0 also satisfies condition A.*

For Gabor frames for $\mathfrak{G}(g, \alpha, \beta)$ for $L^2(\mathbb{R}^d)$ with rational oversampling, that is $\alpha\beta \in \mathbb{Q}$, the corresponding result is proven in [Jans95a]. The proof utilizes the Zak-transform and the so-called Zibulski-Zeevi representation of the Gabor frame operator. I have chosen not to include this theory in the thesis, but based on my experience of translating the other results in this thesis, I see no reason why it should not be possible. The Zak transform can be defined for general locally compact Abelian groups, and therefore also for $\mathbb{T}$. In the following theorem it is made an assumption that the canonical dual window should satisfy condition A, but it is most likely a superfluous assumption.

Theorem 5.27. *Let $\mathfrak{G}$, $\mathfrak{H}$, g be as defined in theorem 5.17. Furthermore assume that g satisfies condition A (definition 5.19). (Assume that canonical dual window for $\mathfrak{G}$: γ^0 also satisfies condition A).*

Then γ^0 satisfies condition R, and $\{L\gamma_k^0\} = \{L\gamma^0(\frac{k}{L})\}_{k=0, \dots, L-1}$ is the canonical dual sequence for $\mathfrak{H}$.

Proof. Let $\mathcal{S}_g$ be the frame operator of $\mathfrak{G}$ and let $\mathcal{S}_{\{g\}}$ be the frame operator of $\mathfrak{H}$.

Since γ^0 is the canonical dual of g then $\gamma^0 = \mathcal{S}_g^{-1}g$. Since $\mathcal{S}_g^{-1} = \mathcal{S}_{\gamma^0}$ is the frame operator for the dual frame $\mathfrak{J}(\gamma^0, \frac{1}{N}, b)$ and since γ^0 satisfies condition A, then $\mathcal{S}_g^{-1}$ can be written as Janssen's representation (theorem 5.23). This gives the expression

$$\gamma^0 = \mathcal{S}_g^{-1}g = \mathcal{S}_\gamma f = \frac{N}{b} \sum_{s=0}^{b-1} \sum_{r \in \mathbb{Z}} \left\langle \gamma^0, \mathcal{M}_{rN} \mathcal{T}_{\frac{s}{b}} \gamma^0 \right\rangle \mathcal{M}_{rN} \mathcal{T}_{\frac{s}{b}} g.$$

By lemma 5.23 then γ^0 satisfies condition R. From $g = \mathcal{S}_g \gamma^0$ and from proposition 5.24 then

$$\left\{ \frac{1}{L} g_k \right\} = \mathcal{S}_{\{g\}} \{ \gamma_k^0 \},$$

which is the same as

$$\{L\gamma_k^0\} = (\mathcal{S}_{\{g\}})^{-1} \{g_k\}.$$

This shows that $\{L\gamma_k^0\}$ is the canonical dual sequence of $\mathfrak{H}$. □

5.3. Bounds. This section is greatly inspired by the corresponding proofs for Fourier series presented in [Briggs95].

Theorem 5.28. *Let $f \in C^{p-1}(\mathbb{T})$, $p \geq 1$. Assume that $f^{(p)}$ is bounded and piecewise monotone and let $\mathfrak{G}(g, \frac{1}{N}, b)$, $N, b \in \mathbb{N}$ be a Gabor frame for $L^2(\mathbb{T})$ with window function $g \in L^2(\mathbb{T})$ and canonical dual window function $\gamma \in C^\infty(\mathbb{T})$. Then the Gabor coefficients of f ,*

$$\{c_{m,n}\} = \mathcal{C}_\gamma f$$

satisfy

$$|c_{m,n}| \leq \frac{C}{|mb|^{p+1}}, \quad \forall m \in \mathbb{Z}, \forall n \in 0, \dots, N-1$$

where C is a constant independent of m and n .

Proof. The coefficient $c_{m,n}$ can be integrated by parts

$$\begin{aligned}
c_{m,n} &= \left\langle f, e^{-2\pi imbx} \gamma \left(x + \frac{n}{N} \right) \right\rangle \\
&= \int_0^1 f(x) e^{2\pi imbx} \bar{\gamma} \left(x + \frac{n}{N} \right) dx \\
&= \left[\frac{e^{2\pi imbx}}{2\pi imb} f(x) \bar{\gamma} \left(x + \frac{n}{N} \right) \right]_0^1 - \frac{1}{2\pi imb} \int_0^1 e^{2\pi imbx} \left(f(x) \bar{\gamma} \left(x + \frac{n}{N} \right) \right)' dx \\
&= -\frac{1}{2\pi imb} \int_0^1 e^{2\pi imbx} \left(f(x) \bar{\gamma} \left(x + \frac{n}{N} \right) \right)' dx.
\end{aligned}$$

The first term in the third line goes away because of the periodicity of f , γ and the complex exponential. Repeating this process a total of p times gives

$$c_{m,n} = \left(-\frac{1}{2\pi imb} \right)^p \int_0^1 e^{2\pi imbx} \left(f(x) \bar{\gamma} \left(x + \frac{n}{N} \right) \right)^{(p)} dx.$$

According to [Briggs95, theorem 6.2], the integral can be bounded by

$$\int_0^1 e^{2\pi imbx} \left(f(x) \bar{\gamma} \left(x + \frac{n}{N} \right) \right)^{(p)} dx \leq \frac{M_n}{mp}$$

where M is a constant independent of m . This leads to the bound

$$|c_{m,n}| \leq \frac{M_n}{|2\pi imb|^{p+1}} = \frac{C_n}{|mb|^{p+1}},$$

where C_n is a constant independent of m but dependent on n . Choosing the constant C as $C = \max_{n=0,\dots,N-1} C_n$ gives the final result. $\square$

Proposition 5.29. *Let $f \in C(\mathbb{T})$ and assume that $f^{(1)}$ is piecewise monotone. Let $\mathfrak{G}$, $\mathfrak{H}$, g and γ^0 be as defined in theorem 5.27, and assume that $\gamma \in C^\infty(\mathbb{T})$. Let*

$$\{c_{m,n}\} = \left\{ \left\langle f, \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} \gamma \right\rangle \right\}_{m \in \mathbb{Z}, n \in 0, \dots, N-1}$$

and

$$\{\tilde{c}_{m,n}\} = \{ \langle \{f_k\}, \mathcal{M}_{mb} \mathcal{N}_{na} \{L\gamma_k\} \rangle \}_{m=0, \dots, M-1, n=0, \dots, N-1},$$

where $f_k = f(\frac{k}{L})$ and $\gamma_k = \gamma(\frac{k}{L})$ for all $k = 0, \dots, L-1$. Then

$$\tilde{c}_{r,s} = \frac{1}{L} \sum_{j \in \mathbb{Z}} c_{r+jM,s},$$

for all $r = 0, \dots, M-1$ and $s = 0, \dots, N-1$.

Note 5.30. This note has been appended after the project was handed in and evaluated. The following proof is incorrect, but the result can be proved in a different way, using the Poisson summation formula.

Proof. The expansion of f in $\mathfrak{G}$ is

$$(5.26) \quad f(x) = \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} c_{m,n} e^{-2\pi imbx} g\left(x + \frac{n}{N}\right), \quad x \in \mathbb{T},$$

with $\{c_{m,n}\} \in l^1(\mathbb{Z} \times \{0, \dots, N-1\})$, by theorem 5.28. This means that the series (5.26) is dominated by an absolutely convergent series

$$\|g\|_\infty \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} |c_{m,n}|,$$

and is therefore uniformly and pointwise convergent to $f(x)$. The expansion of f therefore holds pointwise in the samples $x_k = \frac{k}{L}$, $k = 0, \dots, L-1$:

$$\begin{aligned} f_k = f\left(\frac{k}{L}\right) &= \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} c_{m,n} e^{-2\pi i m b \frac{k}{L}} g\left(\frac{k}{L} + \frac{n}{N}\right) \\ (5.27) \qquad &= \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} c_{m,n} e^{-2\pi i m b \frac{k}{L}} g_{k+na} \end{aligned}$$

for $k = 0, \dots, L-1$. Notice that $g(\frac{k}{L} + \frac{n}{N}) = g_{k+na}$ since $L = Na$.

Since the requirements of theorem 5.27 are satisfied, then $\{L\gamma_k\}$ is the canonical dual window of $\mathfrak{H}$. Analyzing $\{f_k\}$ in $\mathfrak{H}$ gives:

$$\begin{aligned} \tilde{c}_{r,s} &= \left\langle \{f_k\}, \left\{ e^{-2\pi i r b \frac{k}{L}} L\gamma_{k+sa} \right\}_k \right\rangle \\ &= L \sum_{k=0}^{L-1} f_k e^{2\pi i r b \frac{k}{L}} \bar{\gamma}_{k+sa} \\ &= L \sum_{k=0}^{L-1} \left(\sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} c_{m,n} e^{-2\pi i m b \frac{k}{L}} g_{k+na} \right) e^{2\pi i r b \frac{k}{L}} \bar{\gamma}_{k+sa}. \end{aligned}$$

In the final line (5.27) was inserted. Since the sum over k is finite, then it can be moved inside the sum over m :

$$\begin{aligned} \tilde{c}_{r,s} &= L \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} c_{m,n} \sum_{k=0}^{L-1} e^{-2\pi i m b \frac{k}{L}} g_{k+na} e^{2\pi i r b \frac{k}{L}} \bar{\gamma}_{k+sa} \\ &= \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} c_{m,n} \langle \mathcal{M}_{mb} \mathcal{T}_{na} \{g_k\}, \mathcal{M}_{rb} \mathcal{T}_{sa} \{L\gamma_k\} \rangle. \end{aligned}$$

Together with the bi-orthogonality relations for Gabor frames for $\mathbb{C}_p^L$, proposition 4.19, this shows that the inner product in the final line is equal to one if $r = m \bmod M$ and $s = n \bmod N$ and zero otherwise. Using this, the result follows:

$$\tilde{c}_{r,s} = \sum_{j \in \mathbb{Z}} c_{r+jM,s}$$

□

Remark 5.31. It should be possible to loosen the assumptions in the previous theorem. The requirements of $\gamma^0 \in L^\infty(\mathbb{T})$ and $f \in C(\mathbb{T})$ with $f^{(1)}$ being piecewise monotone, are only needed to ensure that the Gabor expansion of f is pointwise convergent. Hopefully this can be ensured in a more clever way.

The Gabor coefficients $\{c_{m,n}\}$ obtained from expanding f Gabor frame for $L^2(\mathbb{T})$ can be approximated by expanding $\{f_k\}$ in a Gabor frame from $\mathbb{C}_p^L$. The error introduced by doing this is outlined in the following theorem.

Theorem 5.32. *Let $f \in C^{p-1}(\mathbb{T})$, $p \geq 1$. Assume that $f^{(p)}$ is bounded and piecewise monotone. Let $\mathfrak{G}$, $\mathfrak{H}$, g , γ^0 be as defined in theorem 5.27, and assume that $\gamma^0 \in L^\infty(\mathbb{T})$. Let*

$$\{c_{m,n}\}_{m \in \mathbb{Z}, n=0, \dots, N-1} = \mathcal{C}_{\gamma^0} f$$

and

$$\{\tilde{c}_{m,n}\}_{m=-\frac{M}{2}+1, \dots, \frac{M}{2}, n=0, \dots, N-1} = \mathcal{C}_{\{L\gamma_k^0\}} \{f_k\}.$$

Then the error of using $\{\tilde{c}_{m,n}\}$ as an approximation of $\{c_{m,n}\}$ is

$$|c_{m,n} - \tilde{c}_{m,n}| \leq \frac{C}{L^{p+1}}, \forall m = -\frac{L}{2} + 1, \dots, \frac{L}{2}, \forall n = 0, \dots, N-1$$

where C is a constant independent of k and L .

Proof. Using proposition 5.29:

$$|c_{m,n} - \tilde{c}_{m,n}| = \left| \sum_{j \in \mathbb{Z} \setminus \{0\}} c_{m+jM,n} \right|.$$

The decay of the Gabor coefficients of f expanded in $\mathfrak{G}$ are known from theorem 5.28, and so the sum can be bounded as

$$\begin{aligned} \left| \sum_{j \in \mathbb{Z} \setminus \{0\}} c_{m+jM,n} \right| &= \left| \sum_{j \in \mathbb{N}} c_{m-jM,n} + c_{m+jM,n} \right| \\ &\leq \sum_{j \in \mathbb{N}} \frac{C_1}{|(m-jM)b|^{p+1}} + \frac{C_1}{|(m+jM)b|^{p+1}} \\ &= \frac{C_1}{L^p} \sum_{j \in \mathbb{N}} \frac{1}{\left|\frac{m}{M} - j\right|^{p+1}} + \frac{1}{\left|\frac{m}{M} + j\right|^{p+1}} \\ &\leq \frac{C_2}{L^{p+1}}, \end{aligned}$$

for all $m = -\frac{L}{2} + 1, \dots, \frac{L}{2}$ and $n = 0, \dots, N-1$ and $p \geq 1$. The constant C_1 is the one occurring in theorem 5.28 and C_2 is independent of L . $\square$

Lemma 5.33. *Let $g \in C(\mathbb{T})$ and $N, b \in \mathbb{N}$. Assume that $\{c_{m,n}\} \in l^1(\mathbb{Z})$. Then the summation*

$$\sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} c_{m,n} \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g$$

is uniformly convergent.

Proof. This is proven by dominated convergence of an absolute convergent series. For all $x \in \mathbb{T}$ it holds

$$\left| c_{m,n} e^{-2\pi i m b x} g\left(x + \frac{n}{N}\right) \right| \leq |c_{m,n}| \left| \sup_{x \in \mathbb{T}} g(x) \right|.$$

Since $\{c_{m,n}\} \in l^1(\mathbb{Z})$ then it holds that

$$\sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} |c_{m,n}| \left| \sup_{x \in \mathbb{T}} g(x) \right| = \left| \sup_{x \in \mathbb{T}} g(x) \right| \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} |c_{m,n}| < \infty.$$

Now the result follows from [Jens92, p. 176]. $\square$

5.4. Existence. All previous statements have characterized Gabor frames, but so far nothing has been mentioned about the existence of Gabor frames. Theorem 5.17 specifies under which conditions a Gabor frame for $\mathbb{C}_p^L$ can be obtained by sampling a Gabor frame for $L^2(\mathbb{T})$. However, this still leaves the question about when a Gabor system in $L^2(\mathbb{T})$ is a frame.

This subsection will show some results of when a Gabor system is a frame.

For the following results let

$$G_{j,l}(x) = \sum_{k=0}^{N-1} \bar{g}(x + \frac{l}{b} + \frac{k}{N}) g(x + \frac{j}{b} + \frac{k}{N}),$$

for $j, l = 0, \dots, b-1$. These $G_{j,l}$ are identical to the correlation functions (definition 5.6, denoted by only one index on G : G_n when $j = 0$: $G_n(x) = G_{0,n}(x)$). Observe that

$$(5.28) \quad G_n(x + \frac{j}{b}) = G_{j,j+n}(x).$$

The following corollary of Walnut's representation, theorem 5.8, is needed to prove the next theorem.

Corollary 5.34. *Let $g \in L^\infty(\mathbb{T})$ and let $b, N \in \mathbb{N}$. Then the operator $\mathcal{S} : L^2(\mathbb{T}) \rightarrow L^2(\mathbb{T})$*

$$\mathcal{S}f = \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} \left\langle f, \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g \right\rangle \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} g$$

can be written as

$$(5.29) \quad \langle \mathcal{S}f, h \rangle = \sum_{l=0}^{b-1} \sum_{k=0}^{b-1} \frac{1}{b} \int_0^{\frac{1}{b}} \left(G_{l,k}(x) \mathcal{T}_{\frac{k}{b}} f(x) \right) \overline{\mathcal{T}_{\frac{l}{b}} h(x)} dx$$

for all $f, h \in L^2(\mathbb{T})$.

Proof. The fact that $\mathcal{S}$ is bounded on $L^2(\mathbb{T})$ follows from corollary 5.5. Substituting Walnut's representation (5.11) for $\mathcal{S}f$:

$$\begin{aligned} \langle \mathcal{S}f, h \rangle &= \frac{1}{b} \int_0^1 \left(\sum_{k=0}^{b-1} G_k(x) f(x + \frac{k}{b}) \right) \overline{h(x)} dx \\ &= \frac{1}{b} \int_0^{\frac{1}{b}} \left(\sum_{l=0}^{b-1} \sum_{k=0}^{b-1} G_k(x + \frac{l}{b}) f(x + \frac{l+k}{b}) \right) \overline{h(x + \frac{l}{b})} dx \\ &= \sum_{l=0}^{b-1} \sum_{k=0}^{b-1} \frac{1}{b} \int_0^{\frac{1}{b}} \left(G_{l,l+k}(x) \mathcal{T}_{\frac{l+k}{b}} f(x) \right) \overline{\mathcal{T}_{\frac{l}{b}} h(x)} dx \\ &= \sum_{l=0}^{b-1} \sum_{k=0}^{b-1} \frac{1}{b} \int_0^{\frac{1}{b}} \left(G_{l,k}(x) \mathcal{T}_{\frac{k}{b}} f(x) \right) \overline{\mathcal{T}_{\frac{l}{b}} h(x)} dx. \end{aligned}$$

Line two comes from splitting the interval $[0; 1]$ in b pieces. In line three the summations are rearranged, but this is fine since they are finite. The correlation function G_k is replaced as in (5.28). In line four $k + l$ is replaced by k . This is valid since $G_{l, l+k}$ and $\mathcal{T}_{\frac{l+k}{b}}$ are both periodic in $l + k$ with period b , and since the summation k runs over $0, \dots, b - 1$. After the replacement, the terms are just added in different order. $\square$

Theorem 5.35. *Let $g \in L^2(\mathbb{T})$ and let $N, b \in \mathbb{N}$. Define for each $x \in \mathbb{T}$ the operator $\mathcal{G}(x) : \mathbb{C}^b \rightarrow \mathbb{C}^b$ by*

$$(5.30) \quad \left\{ \mathcal{G}(x) \{c_l\}_{l=0, \dots, b-1} \right\}_{j=0, \dots, b-1} = \sum_{l=0}^{b-1} G_{j,l}(x) c_l = \mathbf{G}(x) c,$$

where $\mathbf{G}(x) \in \mathbb{C}^{b,b}$ is the matrix defined by

$$(5.31) \quad \mathbf{G}(x)_{j,l} = G_{j,l}(x).$$

Then $\mathcal{S}$ as defined in theorem 5.8 is invertible on $L^2(\mathbb{T})$ if and only if

$$(5.32) \quad \mathcal{G}(x) \geq a \mathcal{I}_{\mathbb{C}^b}, \text{ a.e. } x \in [0; \frac{1}{N}],$$

for some constant $a > 0$.

Note 5.36. Note by the author after the project was evaluated. The logic in the following proof is incorrect, the proof does not show a bi-implication, only an implication. The logic is also flawed in the similar proof presented in [Groch01].

Proof. Since all $G_{j,l}$ are periodic with period $\frac{1}{N}$, then $\mathcal{G}(x) = \mathcal{G}(x + \frac{k}{N})$, $k = 0, \dots, N + 1$. This means that (5.32) is equivalent to $\mathcal{G}(x) \geq a \mathcal{I}_{\mathbb{C}^b}$, a.e. $x \in \mathbb{T}$, because of the periodicity of the $\mathcal{G}(x)$.

Assume that

$$(5.33) \quad \mathcal{G}(x) \geq a \mathcal{I}_{\mathbb{C}^b}, \text{ a.e. } x \in \mathbb{T}.$$

Then $\langle \mathcal{G}(x) \{c_k\}, \{c_k\} \rangle \geq a \|\{c_k\}\|^2$ holds for any sequence $\{c_k\} \in \mathbb{C}^b$ and so it also holds for the sequence $\{c_j\} = \left\{ \mathcal{T}_{\frac{j}{b}} f(x) \right\}_j$ for every $f \in L^\infty(\mathbb{T})$. From this

$$\sum_{j=0}^{b-1} \sum_{l=0}^{b-1} G_{j,l}(x) \mathcal{T}_{\frac{l}{b}} f(x) \overline{\mathcal{T}_{\frac{j}{b}} f(x)} \geq a \sum_{j=0}^{b-1} \left| \mathcal{T}_{\frac{j}{b}} f(x) \right|^2.$$

Integrating this over the interval $[0; \frac{1}{b}]$, the left hand side is recognized as part of (5.29). From this it follows

$$\langle \mathcal{S}f, f \rangle \geq \frac{a}{b} \int_0^{\frac{1}{b}} \sum_{j=0}^{b-1} \left| \mathcal{T}_{\frac{j}{b}} f(x) \right|^2 dx = \frac{a}{b} \|f\|^2.$$

This shows that $\mathcal{S}$ is a positive operator. From the spectral theorem it follows that $\mathcal{S}$ is invertible.

To prove the “only if” part, assume (5.33) as before and furthermore assume that $\mathcal{S}$ is not invertible. Then there exists a function $f \neq 0$ such that $\mathcal{S}f = 0$ and this function can be approximated by a sequence of bounded functions. From this sequence a subsequence $f_n \in L^\infty(\mathbb{T})$ can be chosen such that

$$\langle \mathcal{S}f_n, f_n \rangle < \frac{1}{n} \|f_n\|_2^2.$$

Using (5.29) then

$$\begin{aligned}
0 &< \frac{1}{n} \|f_n\|_2^2 - \langle \mathcal{S} f_n, f_n \rangle \\
&= \frac{1}{n} \int_0^1 |f_n(x)|^2 dx - \sum_{l=0}^{b-1} \sum_{k=0}^{b-1} \frac{1}{b} \int_0^{\frac{1}{b}} \left(G_{l,k}(x) \mathcal{T}_{\frac{k}{b}} f_n(x) \right) \overline{\mathcal{T}_{\frac{l}{b}} f_n(x)} dx \\
&= \int_0^{\frac{1}{b}} \left(\frac{1}{n} \sum_{j=0}^{b-1} \left| f_n(x + \frac{j}{b}) \right|^2 - \frac{1}{b} \sum_{l=0}^{b-1} \sum_{k=0}^{b-1} \left(G_{l,k}(x) \mathcal{T}_{\frac{k}{b}} f_n(x) \right) \overline{\mathcal{T}_{\frac{l}{b}} f_n(x)} \right) dx.
\end{aligned}$$

The third line comes from writing the first integration in line two as an integration over $[0; \frac{1}{b}]$ of a periodization of f_n , and then collecting everything in one integration.

Therefore there exists a set $E_n \subseteq [0; \frac{1}{b}]$ with positive measure such that

$$\frac{1}{n} \sum_{j=0}^{b-1} \left| f_n(x + \frac{j}{b}) \right|^2 - \frac{1}{b} \sum_{l=0}^{b-1} \sum_{k=0}^{b-1} \left(G_{l,k}(x) \mathcal{T}_{\frac{k}{b}} f_n(x) \right) \overline{\mathcal{T}_{\frac{l}{b}} f_n(x)} > 0, \forall x \in E_n.$$

Writing $\mathcal{T}_{\frac{l}{b}} f(x)$ as c_j then it follows that

$$\langle \mathcal{G}c, c \rangle < \frac{b}{n} \sum_{j=0}^{b-1} |c_j|_2^2, \forall x \in E_n,$$

or differently phrased

$$\mathcal{G}(x) < \frac{b}{n} \mathcal{I}, \forall x \in E_n.$$

This contradicts (5.33) because E_n has positive measure, and therefore $\mathcal{S}$ must be invertible. $\square$

Remark 5.37. If $\mathcal{S}$, g , N, b are as defined in theorem 5.8 and $\mathcal{S}$ is a bounded and invertible operator then $\mathfrak{G}(g, \frac{1}{N}, b)$ is a frame for $L^2(\mathbb{T})$. This follows from standard frame theory, see [Chr02].

The operator $\mathcal{G}(x)$ is a positive, semi-definite operator:

Proposition 5.38. *The operator $\mathcal{G}(x)$ as defined in (5.30) is a positive semi-definite operator in the sense that*

$$\langle \mathcal{G}(x)c, c \rangle \geq 0, \forall c \in \mathbb{C}^b.$$

Proof. By direct calculation

$$\begin{aligned}
\langle \mathcal{G}(x)c, c \rangle &= \left\langle \left\{ \sum_{l=0}^{b-1} G_{j,l}(x)c_l \right\}_j, \{c_j\} \right\rangle \\
&= \sum_{j=0}^{b-1} \sum_{l=0}^{b-1} G_{j,l}(x)c_l \bar{c}_j \\
&= \sum_{k=0}^{N-1} \left(\sum_{j=0}^{b-1} \bar{c}_j g\left(x + \frac{j}{b} + \frac{k}{N}\right) \right) \left(\sum_{l=0}^{b-1} c_l \bar{g}\left(x + \frac{l}{b} + \frac{k}{N}\right) \right) \\
&= \sum_{k=0}^{N-1} \left| \sum_{j=0}^{b-1} \bar{c}_j g\left(x + \frac{j}{b} + \frac{k}{N}\right) \right|^2 \geq 0.
\end{aligned}$$

In line three the expression (5.30) is inserted and the summations are rearranged. Since the summations over j and l are independent, the expressions in the parenthesis in line three is a number and its complex conjugate, and from this the result follows. $\square$

Theorem 5.35 provides a way to test if a set g, b, N with $g \in L^\infty(\mathbb{T})$, $b, N \in \mathbb{N}$ will generate a Gabor frame. How to use the test in practice will be discussed in section 7.

6. INTERPOLATION AND NUMERICAL DIFFERENTIATION

This is my own collection of results from standard books on analysis, and for the second part, also from [Briggs95, Tref].

On a computer it is impossible to represent a general, continuous function f , as it would require infinite storage. A common practice is instead to store a finite number of samples of f : $\{f_n\}$ at grid points $\{x_n\}$, and use the samples to represent the function f . As there exist infinitely many continuous functions interpolating a given set of discrete points, then one must choose a method to interpolate the points, and then hope that this interpolating function $\tilde{f}$ through the samples $\{f_n\}$ closely match the original sampled function f .

The process of numerical differentiation consists of differentiating the interpolating function $\tilde{f}$, and using this derivative $\tilde{f}'$ as an approximation to the derivative f' of the original function f .

All considered functions will belong to $L^2(\mathbb{T})$, and all grids will be equidistantly spaced on $\mathbb{T}$ denoted by the interval $[0; 1]$. For a grid $\{x_n\}$, $n = 0, \dots, L - 1$:

$$x_k = kh = \frac{k}{L}, \forall k \in 0, \dots, L - 1,$$

so h is

$$h = \frac{1}{L}.$$

The next section defines two ways of using polynomials to do numerical differentiation, and derives an error bound for one of them.

6.1. Polynomial differentiation. Polynomial differentiation could be done by interpolating the samples $\{f_n\}$ of the function f at the points $\{x_n\}$ by a polynomial, differentiating this polynomial, and sampling the derivative. However, the method described here is even simpler, because no global interpolating function is constructed. Instead, for each $f_k = f(x_k)$, a local interpolating polynomial is found, and then this polynomial is differentiated and sampled at the point x_k .

Definition 6.1. Let f be a continuous function $f \in C(\mathbb{T})$ with $hL = 1$ and let $\{f_k\} = f(x_k)$ with $x_k = \frac{k}{L}$. The *second-order centered³ finite difference* is given by

$$\left\{ \tilde{f}'_k \right\} = \left\{ \frac{f_{k+1} - f_{k-1}}{2h} \right\}_k.$$

Similarly the fourth-order centered finite difference can be defined:

Definition 6.2. Let f be a continuous, periodic function $f \in C^0(\mathbb{T})$ with period $hL = 1$ and let $\{f_k\} = f(x_k)$ with $x_k = hk$. The *fourth-order centered finite difference* is given by:

$$\left\{ \tilde{f}'_k \right\} = \left\{ \frac{-f_{k+2} + 8f_{k+1} - 8f_{k-1} + f_{k-2}}{12h} \right\}_k.$$

An error bound for the use of the second-order difference is:

Proposition 6.3. Let $f \in C^3(\mathbb{T})$ with $hL = 1$. Then the relative, global error of using the second-order centered finite difference as defined in definition 6.1 to approximate the sampled derivative $\{f'_k\} = f'(x_k)$ of f is:

$$\frac{\left\| \{f'_n\} - \{\tilde{f}'_n\} \right\|}{\left\| \{f'_n\} \right\|} \leq Ch^2,$$

where C is a constant independent of h .

Proof. The Taylor expansion for f_{k+1} developed around the point x_k is

$$f_{k+1} = f_k + h \frac{df}{dx}(x_k) + \frac{h^2}{2} \frac{d^2f}{dx^2}(x_k) + \frac{h^3}{6} \frac{d^3f}{dx^3}(\xi_1),$$

where $\xi_1 \in [x_k; x_{k+1}]$. Similarly for f_{k-1} :

$$f_{k-1} = f_k - h \frac{df}{dx}(x_k) + \frac{h^2}{2} \frac{d^2f}{dx^2}(x_k) - \frac{h^3}{6} \frac{d^3f}{dx^3}(\xi_2),$$

where $\xi_2 \in [x_{k-1}; x_k]$. Constructing the second-order centered finite difference from these gives

$$(6.1) \quad \frac{f_{k+1} - f_{k-1}}{2h} = \frac{df}{dx}(x_k) + \frac{h^2}{12} \left(\frac{d^3f}{dx^3}(\xi_1) + \frac{d^3f}{dx^3}(\xi_2) \right).$$

The term $\frac{d^3f}{dx^3}(\xi_1) + \frac{d^3f}{dx^3}(\xi_2)$ can always be bounded from above by a constant

$$(6.2) \quad C = 2 \left(\sup_{y \in [0;1]} \left| \frac{d^3f}{dx^3}(y) \right| \right),$$

since f is assumed to be both continuous and periodic.

³“Centered” refers to the approximation of the derivative being calculated in the same place as one of the grid points.

This shows, that the error $\left|f'(x_i) - \tilde{f}'_i\right|$ decreases as h^2 if the function f is sampled on a finer grid.

The global, relative error $\frac{\|\{f'_n\} - \{\tilde{f}'_n\}\|}{\|\{f'_n\}\|}$ also decreases as h^2 since

$$\frac{\|\{f'_n\} - \{\tilde{f}'_n\}\|}{\|\{f'_n\}\|} \leq \frac{\|\{h^2 C\}_n\|}{\|\{C_1 C\}_n\|} = \frac{1}{C_1} h^2 = C_2 h^2$$

where C is as in (6.2) and $C_1 = \frac{1}{C} \inf_{y \in [0;1]} \left| \frac{df}{dx}(y) \right|$. Both C and C_1 are independent of h . $\square$

A similar analysis requiring f to be five times differentiable would show that the error of fourth-order centered finite difference decreases as h^4 .

The assumption of proposition 6.3 is that f is differentiable. This assumption excludes many interesting situations, where f might only be piece-wise continuous. I will not analyze this situation just make the following remark:

Remark 6.4. The computation of the second order finite difference is very localized. To compute an approximation to the derivative at x_k it only uses the values f_{k+1} and f_{k-1} , so if just f is differentiable in a region around x_{k-1} , x_k and x_{k+1} , then it does not influence the computation that f might have discontinuities elsewhere. Same argument goes for the fourth order finite difference, except that the region around x_k where f must be differentiable is larger.

6.2. Spectral differentiation. Spectral differentiation is done by interpolating the samples $\{f_k\}$ of a function f with the interpolant

$$p(x) = \frac{1}{\sqrt{L}} \sum_{k=-\frac{L}{2}+1}^{\frac{L}{2}} \hat{f}_k e^{2\pi i k x}, \quad x \in [0;1].$$

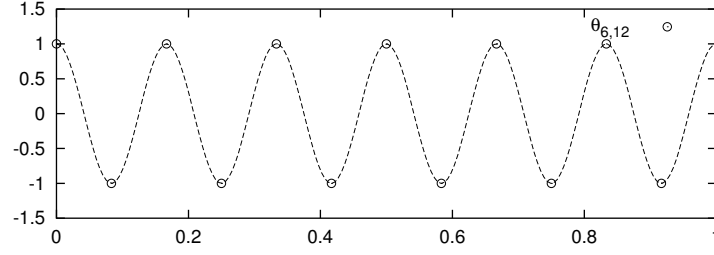
This is the IDFT from (3.7) extended to a continuous function on $\mathbb{T}$. It can easily be differentiated and sampled

$$\{p'_n\} = \left\{ \frac{2\pi}{\sqrt{L}} \sum_{k=-\frac{L}{2}+1}^{\frac{L}{2}} i k \hat{f}_k e^{2\pi i \frac{k n}{L}} \right\}_n.$$

Here the different terms in the summation has been straightforward differentiated. This is not quite the way do it - the problem is the modes with highest and lowest wavenumber: $\theta_{\frac{L}{2},L}$ and $\theta_{-\frac{L}{2},L}$. These modes are equal, and could both equally well be included in the summation. Sampling $\theta_{\frac{L}{2},L} = \theta_{-\frac{L}{2},L}$ on the grid gives

$$e^{\pm 2\pi i \frac{n}{2}} = \begin{cases} 1 & \text{if } n \text{ is even} \\ -1 & \text{if } n \text{ is odd,} \end{cases}$$

a real sequence, as shown on figure 6.1. Therefore the derivative of the function represented by this sequence should be real and zero at the grid points, and not a complex exponential. This lead to the definition:

FIGURE 6.1. The mode $\sqrt{L}\theta_{6,12}$ and its interpolant $\cos(2\pi i \frac{x}{12})$.

Definition 6.5. Let $f \in C(\mathbb{T})$. The *spectral derivative* of f is defined by

$$(6.3) \quad \left\{ \tilde{f}'_n \right\} = \left\{ \frac{2\pi}{\sqrt{L}} \sum_{k=-\frac{L}{2}+1}^{\frac{L}{2}-1} ik \hat{f}_k e^{2\pi i \frac{kn}{L}} \right\}_n.$$

The $\frac{L}{2}$ -term has been removed, because it should not contribute.

To analyze the validity and accuracy of definition 6.5, f and f' will be expanded in Fourier series.

Now the accuracy of the spectral derivative can be stated.

Theorem 6.6. Let $f \in C^{p-1}(\mathbb{T})$. Assume that $f^{(p)}$ is bounded and piecewise monotone. Then the global error of using the spectral derivative $\{\tilde{f}'_k\}$ as an approximation to the sampled derivative $\{f'_k\}$ is

$$\left\| \left\{ \tilde{f}'_k \right\} - \left\{ f'_k \right\} \right\| \leq C \frac{1}{L^{p-1}},$$

where L is the number of samples and C is a constant independent of L .

Proof. The Fourier expansions of f are:

$$f(x) = \sum_{k \in \mathbb{Z}} c_k e^{2\pi i k x},$$

where

$$c_k = \int_0^1 f(x) e^{-2\pi i k x} dx, \quad \forall k \in \mathbb{Z}.$$

Similarly then f' can be expanded as

$$(6.4) \quad f'(x) = \sum_{k \in \mathbb{Z}} d_k e^{2\pi i k x},$$

where

$$d_k = \int_0^1 f'(x) e^{-2\pi i k x} dx, \quad \forall k \in \mathbb{Z}.$$

Using partial integration on d_k :

$$(6.5) \quad \begin{aligned} d_k &= \left[f(x) e^{-2\pi i k x} \right]_0^1 + 2\pi i k \int_0^1 f(x) e^{-2\pi i k x} dx \\ &= 2\pi i k c_k, \end{aligned}$$

for all $k \in \mathbb{Z}$. This shows that the Fourier series for f' : $\{d_k\}$ can be found from the Fourier series for f , $\{c_k\}$. Theorem 3.12 and (6.5) shows that an approximation $\{\tilde{d}_k\}$ to the Fourier series $\{d_k\}$ for the derivative of f can be calculated from the DFT of f :

$$\tilde{d}_k = 2\pi i k \frac{1}{\sqrt{L}} \hat{f}_k.$$

Inserting this in (6.4) and truncating the summation and sampling at the grid point gives the expression

$$\left\{ \tilde{f}'_n \right\} = \left\{ \frac{1}{\sqrt{L}} \sum_{k=-\frac{L}{2}+1}^{\frac{L}{2}-1} 2\pi i k \hat{f}_k e^{2\pi i k \frac{n}{L}} \right\}_{n \in \mathbb{Z}}.$$

This is the spectral derivative as defined in definition 6.5.

The error in using the spectral derivative as an approximation to the sampled derivative is the residual sequence $\{r_n\} = \{f'_n - \tilde{f}'_n\}$. This can be analyzed using (6.4) and sampled on the grid:

$$\begin{aligned} \{r_n\} = \{f'_n - \tilde{f}'_n\} &= \left\{ \sum_{k \in \mathbb{Z}} d_k e^{2\pi i k \frac{n}{L}} \right\}_n - \left\{ \frac{1}{\sqrt{L}} \sum_{k=-\frac{L}{2}+1}^{\frac{L}{2}-1} 2\pi i k \hat{f}_k e^{2\pi i k \frac{n}{L}} \right\}_n \\ &= \left\{ \sum_{k=-\frac{L}{2}+1}^{\frac{L}{2}-1} \left(d_k - \frac{1}{\sqrt{L}} 2\pi i k \hat{f}_k \right) e^{2\pi i k \frac{n}{L}} \right\}_n + \left\{ \sum_{k \in \mathbb{Z} \setminus \{-\frac{L}{2}+1, \dots, \frac{L}{2}-1\}} d_k e^{2\pi i k \frac{n}{L}} \right\}_n \\ &= \left\{ \sum_{k \in \mathbb{Z}} r_k e^{2\pi i k \frac{n}{L}} \right\}_n \\ &= \left\{ \sum_{j \in \mathbb{Z}} \sum_{k=0}^{L-1} r_{k+jL} e^{2\pi i (k+jL) \frac{n}{L}} \right\}_n = \left\{ \sum_{j \in \mathbb{Z}} \sum_{k=0}^{L-1} r_{k+jL} e^{2\pi i k \frac{n}{L}} \right\}_n, \end{aligned}$$

where

$$r_k = \begin{cases} d_k - 2\pi i k \frac{1}{\sqrt{L}} \hat{f} = 2\pi i k (c_k - \frac{1}{\sqrt{L}} \hat{f}_k) & \text{if } k \in \{-\frac{L}{2}+1, \dots, \frac{L}{2}-1\} \\ d_k & \text{otherwise} \end{cases}, \quad \forall k \in \mathbb{Z}.$$

Since the exponential function is periodic, the last term can be simplified.

The terms $|r_k|$ can be estimated using theorem 3.11 and theorem 3.12:

$$|r_k| = \begin{cases} |k| \frac{C}{L^{p+1}} & \text{if } k \in \{-\frac{L}{2}+1, \dots, \frac{L}{2}-1\} \\ \frac{C}{|k|^{p+1}} & \text{otherwise} \end{cases}, \quad \forall k \in \mathbb{Z},$$

where the constant C is the largest of the constants in the two theorems and the factor 2π is included. Notice that theorem 5.28 is not directly applicable to $d_k - 2\pi i k \frac{1}{\sqrt{L}} \hat{f}_k$ but instead to $2\pi i k (c_k - \frac{1}{\sqrt{L}} \hat{f}_k)$, and this makes the term $|k|$ enter the bound.

Since the modes form an orthonormal basis for $\mathbb{C}_p^L$ then Parseval's equation and the estimate for $|r_k|$ can be used to calculate $\|\{r_k\}\|_2$:

$$\begin{aligned}
\|\{r_k\}\|_2^2 &= L \left\| \sum_{j \in \mathbb{Z}} \sum_{k=0}^{L-1} r_{k+jL} \frac{1}{\sqrt{L}} e^{2\pi i k \frac{n}{L}} \right\|_2^2 \\
&\leq L \sum_{j \in \mathbb{Z}} \sum_{k=0}^{L-1} |r_{k+jL}|^2 \\
&= L \sum_{k \in \mathbb{Z}} |r_k|^2 \\
&= CL \left(\left| \frac{1}{L^{p+1}} \right|^2 \sum_{k \in \{-\frac{L}{2}+1, \dots, \frac{L}{2}-1\}} |k|^2 + \sum_{k \in \mathbb{Z} \setminus \{-\frac{L}{2}+1, \dots, \frac{L}{2}-1\}} \left| \frac{1}{|k|^{p+1}} \right|^2 \right) \\
&\leq C_1 L \left(\frac{1}{L^{2p+2}} L^3 + \frac{1}{L^{2p+2}} \right) \\
&\leq C_2 \frac{1}{L^{2p-2}}.
\end{aligned}$$

The constants C_1 and C_2 are independent of L . This gives the result:

$$\|\{r_k\}\| \leq C_2 \frac{1}{L^{p-1}}.$$

□

From this it can be seen that the error of the spectral differentiation decays faster than any polynomial when f is smooth. When f is not smooth the rate of convergence is limited by how many continuous derivatives f have. Further work is done in [Tref, theorem 4], where the result is improved for analytical and band-limited functions.

To compare with remark 6.4 it is clear that computing the spectral derivative at a point x_k involves all samples of a function f . Therefore if f has a discontinuity, then this discontinuity affects the computation of all $\tilde{f}'_k$, and not just those $\tilde{f}'_k$ close to the discontinuity.

6.3. The Gabor derivative. Gabor differentiation is done in the same manner as the spectral differentiation: By interpolating the samples $\{f_k\}$ of a function f with an interpolant

$$\tilde{f}(x) = \sum_{n=0}^{N-1} \sum_{m=-\frac{M}{2}+1}^{\frac{M}{2}-1} \tilde{c}_{m,n} e^{-2\pi i m b x} g(x + \frac{n}{N}), \quad x \in \mathbb{T},$$

which is the IDGT from (4.10) extended to a continuous function on $\mathbb{T}$ and $\tilde{c}_{m,n}$ are the corresponding Gabor coefficients. Differentiating this yields:

$$(6.6) \quad \tilde{f}'(x) = \sum_{n=0}^{N-1} \sum_{m=-\frac{M}{2}+1}^{\frac{M}{2}-1} -2\pi i m b \tilde{c}_{m,n} e^{-2\pi i m b x} g(x + \frac{n}{N}) + \tilde{c}_{m,n} e^{-2\pi i m b x} g'(x + \frac{n}{N})$$

Sampling this leads to the definition:

Definition 6.7. Let $f \in C(\mathbb{T})$ and $\mathfrak{H}(\{g_k\}, a, b)$ be a Gabor system for $\mathbb{C}_p^L$ with canonical dual sequence $\{\gamma_k^0\}$. The *Gabor derivative* of f is defined by

$$\{\tilde{f}'_k\} = \left\{ \sum_{n=0}^{N-1} \sum_{m=-\frac{M}{2}+1}^{\frac{M}{2}-1} \left(-2\pi i m b \tilde{c}_{m,n} e^{-2\pi i \frac{m}{M} k} g_{k+na} + \tilde{c}_{m,n} e^{-2\pi i \frac{m}{M} k} g'_{k+na} \right) \right\}_k,$$

where

$$\tilde{c}_{m,n} = \langle \{f'_k\}_k, \mathcal{M}_{mb} \mathcal{T}_{na} \{\gamma_k^0\} \rangle$$

are the coefficients of the expansion of $\{f'_k\}$ in $\mathfrak{H}$ and $f_k = f(\frac{k}{L})$, $k = 0, \dots, L-1$.

Lemma 6.8. Let $f \in C^1(\mathbb{T})$ and assume that $f^{(2)}$ is piecewise monotone. let $\mathfrak{G}(g, \frac{1}{N}, b)$, $g \in C^1(\mathbb{T})$. Assume that the canonical dual window of $\mathfrak{G}$: $\gamma^0 \in L^\infty(\mathbb{T})$. Then

$$(6.7) \quad f'(x) = \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} \left(-2\pi i m b c_{m,n} e^{-2\pi i m b x} g(x + \frac{n}{N}) + c_{m,n} e^{-2\pi i m b x} g'(x + \frac{n}{N}) \right),$$

where

$$c_{m,n} = \langle f, \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} \gamma^0 \rangle.$$

Proof. Since $\mathfrak{G}$ is a Gabor frame then f has the expansion

$$(6.8) \quad f(x) = \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} c_{m,n} e^{-2\pi i m b x} g(x + \frac{n}{N}).$$

A standard result from analysis [Jens92, prop. 23.3] is that if (6.8) is pointwise convergent and if the right hand side of (6.7) is uniformly convergent, the right hand side of (6.7) converges to $f'(x)$ for all $x \in \mathbb{T}$. Since $f \in C^1(\mathbb{T})$ and $f^{(2)}$ is piecewise monotone then from theorem 5.28 the decay of $c_{m,n}$ is $|c_{m,n}| < \frac{C}{|mb|^3}$, and so $\{c_{m,n}\} \in l^1(\mathbb{Z} \times \{0, \dots, N-1\})$. Lemma 5.33 then shows that (6.8) is uniformly convergent, and therefore also pointwise convergent.

The right hand side of (6.7) can be written as

$$(6.9) \quad \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} -2\pi i m b c_{m,n} e^{-2\pi i m b x} g(x + \frac{n}{N}) + \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} c_{m,n} e^{-2\pi i m b x} g'(x + \frac{n}{N}).$$

The sequence $\{-2\pi i m b c_{m,n}\} \in l^1(\mathbb{Z} \times \{0, \dots, N-1\})$ since $|-2\pi i m b c_{m,n}| < \frac{C}{|mb|^2}$ and since $g' \in C(\mathbb{T})$, then lemma 5.33 is satisfied for both summations in (6.9) and so the right hand side of (6.7) is uniformly convergent and converges to $f'(x)$. $\square$

Theorem 6.9. Let $f \in C^{p-1}(\mathbb{T})$, $p \geq 2$. Assume that $f^{(p)}$ is piecewise monotone. Let $\mathfrak{G}$, $\mathfrak{H}$, g , γ^0 be as defined in theorem 5.27. Assume that $g \in C^1(\mathbb{T})$, and that $\gamma^0 \in C^\infty(\mathbb{T})$. Then the global error in using the Gabor derivative $\{\tilde{f}'_k\}$ as defined in definition 6.7 as an approximation to $\{f'_k\}$ is

$$(6.10) \quad \left\| \{f'_k\} - \{\tilde{f}'_k\} \right\| \leq C \frac{1}{\sqrt{L}} \frac{MN}{L^{p-1}},$$

where C is a constant independent of L, M and N .

Proof. Consider the expansion of f in $\mathfrak{G}$:

$$f(x) = \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} c_{m,n} e^{-2\pi i m b x} g(x + \frac{n}{N}),$$

where

$$c_{m,n} = \left\langle f, \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} \gamma^0 \right\rangle,$$

and

$$\tilde{c}_{m,n} = \left\langle \{f_k\}, \mathcal{M}_{mb} \mathcal{T}_{\frac{n}{N}} \{L\gamma_k^0\} \right\rangle$$

Using lemma 6.8 then $f'(x)$ can be written as:

$$f'(x) = \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} \left(-2\pi i m b c_{m,n} e^{-2\pi i m b x} g(x + \frac{n}{N}) + c_{m,n} e^{-2\pi i m b x} g'(x + \frac{n}{N}) \right).$$

The residual sequence $\{r_k\} = \{f'(\frac{k}{L})\} - \{\tilde{f}'_k\}$ where $\{\tilde{f}'_k\}$ is the Gabor derivative from definition 6.7) can now be calculated

$$\begin{aligned} r_k &= f'(\frac{k}{L}) - \tilde{f}'_k \\ &= \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} \left(-2\pi i m b c_{m,n} e^{-2\pi i m b \frac{k}{L}} g(\frac{k}{L} + \frac{n}{N}) + c_{m,n} e^{-2\pi i m b \frac{k}{L}} g'(\frac{k}{L} + \frac{n}{N}) \right) \\ &\quad - \sum_{n=0}^{N-1} \sum_{m=-\frac{M}{2}+1}^{\frac{M}{2}-1} \left(-2\pi i m b \tilde{c}_{m,n} e^{-2\pi i m b \frac{k}{L}} g_{k+na} + \tilde{c}_{m,n} e^{-2\pi i m b \frac{k}{L}} g'_{k+na} \right) \\ &= \sum_{n=0}^{N-1} \sum_{j \in \mathbb{Z}} \sum_{m=0}^{M-1} -2\pi i m b s_{m+jM,n} e^{-2\pi i (m+jM)b \frac{k}{L}} g_{k+na} + \sum_{n=0}^{N-1} \sum_{j \in \mathbb{Z}} \sum_{m=0}^{M-1} s_{m+jM,n} e^{-2\pi i (m+jM)b \frac{k}{L}} g'_{k+na} \\ &= \sum_{j \in \mathbb{Z}} \mathcal{D}_{\{g\}} \{ -2\pi i (m+jM) b s_{m+jM,n} \}_{m=0,\dots,M-1,n} + \sum_{j \in \mathbb{Z}} \mathcal{D}_{\{g'\}} \{ s_{m+jM,n} \}_{m=0,\dots,M-1,n}, \end{aligned}$$

where

$$s_{m,n} = \begin{cases} c_{m,n} - \tilde{c}_{m,n} & \text{if } m \in \{-\frac{M}{2} + 1, \dots, \frac{M}{2} - 1\} \\ c_{m,n} & \text{otherwise} \end{cases}, \quad m \in \mathbb{Z}, n = 0, \dots, N-1,$$

and $\{s_{m,n}\} \in l^2(\mathbb{Z} \times \{0, \dots, N-1\})$ because $\{c_{m,n}\} \in l^2(\mathbb{Z} \times \{0, \dots, N-1\})$. In line four the summations are rearranged, and because the exponential functions are periodic, then line five follows. The operator $\mathcal{D}_{\{g\}}$ is the pre-frame operator of the frame $\mathfrak{H}$ and $\mathcal{D}_{\{g'\}}$ is an operator as defined in proposition 4.15.

From proposition 4.15, the operator norm of the two pre-frame operators are bounded by $C_1 = \sqrt{LMN} \max_k |g_k|$ and $C_2 = \sqrt{LMN} \max_k |g'_k|$ respectively. Using this

$$\begin{aligned}
\|\{r_k\}\|_2^2 &\leq \sum_{j \in \mathbb{Z}} \|\mathcal{D}_{\{g\}} \{-2\pi i(m+jM)bs_{m+jM,n}\}\|_2^2 + \sum_{j \in \mathbb{Z}} \|\mathcal{D}_{\{g'\}} \{s_{m+jM,n}\}\| \\
&\leq 4\pi^2 C_1^2 LMN \sum_{j \in \mathbb{Z}} \sum_{n=0}^{N-1} \sum_{m=0}^{M-1} |(m+jM)bs_{m+jM,n}|^2 + C_1^2 LMN \sum_{j \in \mathbb{Z}} \sum_{n=0}^{N-1} \sum_{m=0}^{M-1} |s_{m+jM,n}|^2 \\
&= LMN \sum_{n=0}^{N-1} \sum_{m \in \mathbb{Z}} (4\pi^2 Bm^2 b^2 + C^2) |s_{m,n}|^2 \\
&\leq LMN^2 \left(\sum_{m=-\frac{M}{2}+1}^{\frac{M}{2}-1} (4\pi^2 Bm^2 b^2 + C^2) \frac{C_1}{L^{2p+2}} + \sum_{m \in \mathbb{Z} \setminus \{-\frac{M}{2}+1, \dots, \frac{M}{2}-1\}} \frac{C_2(4\pi^2 Bm^2 b^2 + C^2)}{|mb|^{2p+2}} \right) \\
&\leq LMN^2 \left(4\pi^2 \frac{C_3}{L^{2p+2}} ML^2 + \frac{C_4}{L^{2p}} \right) \\
&\leq C \frac{M^2 N^2}{L} \frac{1}{L^{2p-2}}
\end{aligned}$$

In line three the summations over m and j are collapsed to one summation. In line five $s_{m,n}$ is expanded and replaced by the bounds from theorem 5.32 and theorem 5.28, and the summation over n is removed because the bounds are independent of n . The first term in line five comes from $b^2 \sum_m m^2 \leq CM^3 b^2 = cML^2$. All constants are independent of M, N and L . $\square$

Perspectives. In the discussion so far it has been assumed that a known function f is sampled, and that the samples then are used to construct an approximation to the derivative of f . There is little point in doing this if f is already known analytically. Then it would probably be much easier and more accurate to just differentiate f analytically.

On the other hand, if f is unknown, and only samples of f are known, there exists infinitely many functions interpolating the samples, so asking for the accuracy of the method make little sense: It might just as well be a cubic spline interpolating the samples as it might be a superposition of sine and cosine functions. In the first case, a fourth order finite difference would give the exact result, and in the second case a spectral method would give the exact result.

When solving partial differential equations the situation is a mixture: The function f might be an analytically known initial condition to a partial differential equation

$$\mathcal{L}u = g$$

where

$$u(x, 0) = f(x),$$

but as time evolves then $u(x, t)$ for t fixed is no longer known analytically. In this case the error bounds make perfect sense. They allow for an estimate of the error

$$\|\{u(x_k, t)\}_k - \{\tilde{u}_k(t)\}_k\|$$

where $\tilde{u}_k(t)$ is the numerical approximation to $u(x_k, t)$.

7. NUMERICAL EXPERIMENTS.

7.1. Dual windows. This section will show graphs and results for Gabor frames for $L^2(\mathbb{T})$. All these results have been numerically computed by sampling a Gabor frame for $L^2(\mathbb{T})$ to obtain a Gabor frame for $\mathbb{C}_p^L$. This process works because of theorem 5.17 and theorem 5.27 are satisfied, so the assumptions of these theorems must also be satisfied.

To construct Gabor frames for $L^2(\mathbb{T})$ the following class of window functions will be used

$$\mathcal{P}\varphi_W(x) = \sum_{n \in \mathbb{Z}} \varphi_W(x+n) = 2^{\frac{1}{4}} \sum_{n \in \mathbb{Z}} e^{-W\pi(x+n)^2}, \quad x \in \mathbb{T}$$

These are the periodized Gaussian functions. Since the Gaussians are in the Schwartz class of smooth, faster than polynomial decaying functions, the Gaussians are also in $L^1(\mathbb{R})$, and therefore their periodizations are well defined.

Lemma 7.1. *The periodized Gaussians are smooth, $\mathcal{P}\varphi_W \in C^\infty(\mathbb{T})$.*

Proof. The summation defining the periodized Gaussian can be bounded by

$$\begin{aligned} \sum_{n \in \mathbb{Z}} \varphi_W(x+n) &= \sum_{n \in \mathbb{N} \cup \{0\}} \varphi_W(x+n) + \sum_{n \in \mathbb{N}} \varphi_W(x-n) \\ &\leq \sum_{n \in \mathbb{N}} \varphi_W(n) + \sum_{n \in \mathbb{N}} \varphi_W(-n+1). \end{aligned}$$

In line two the summation is split in two, and each interval is bounded by its largest element.

This shows that the summation can be bounded by an absolutely convergent series, and the series is therefore uniformly convergent. Differentiating each term in the summation will multiply each term by a factor $-2W\pi(x+n)$, but this is just a polynomial, so the new series can be bounded by an absolutely convergent series as well:

$$\sum_{n \in \mathbb{Z}} -2\pi W(x+n)\varphi_W(x+n) \leq \sum_{n \in \mathbb{N}} 2\pi W n \varphi_W(n) + \sum_{n \in \mathbb{N}} 2\pi W n \varphi_W(-n+1).$$

This shows that the series is also absolutely convergent, and therefore converges to the derivative of $\mathcal{P}\varphi_W$.

This argument was enough to prove that $\mathcal{P}\varphi_W$ is differentiable. Repeating it on the derivatives of $\mathcal{P}\varphi_W$ shows that $\mathcal{P}\varphi_W$ is smooth. $\square$

Since the periodized Gaussians are continuous, they satisfy condition R (remark 5.16) and condition A (proposition 5.20). This means that the periodized Gaussians satisfies all the requirements of theorem 5.27, except from the requirement that the canonical dual window have to satisfy condition A. I choose to believe in conjecture 5.26, since I cannot verify the condition.⁴

The parameter W controls the “width” of the function. If W is close to zero the function is almost constant, and if W is large the function has one large spike. Examples of periodized Gaussians of different width can be seen on figure 7.1.

To compute the function numerically, it is only necessary to do enough iterations to make sure that $e^{-W\pi(x+n)^2}$ has decayed to the machine precision.

Some examples of periodized Gaussians and their dual window can be seen on figure 7.1.

⁴Looking at the calculated canonical dual windows many of them look differentiable, so by proposition 5.20 they would satisfy condition A. The problem is that they cannot be calculated without using theorem 5.27, so this is not the way to solve the problem.

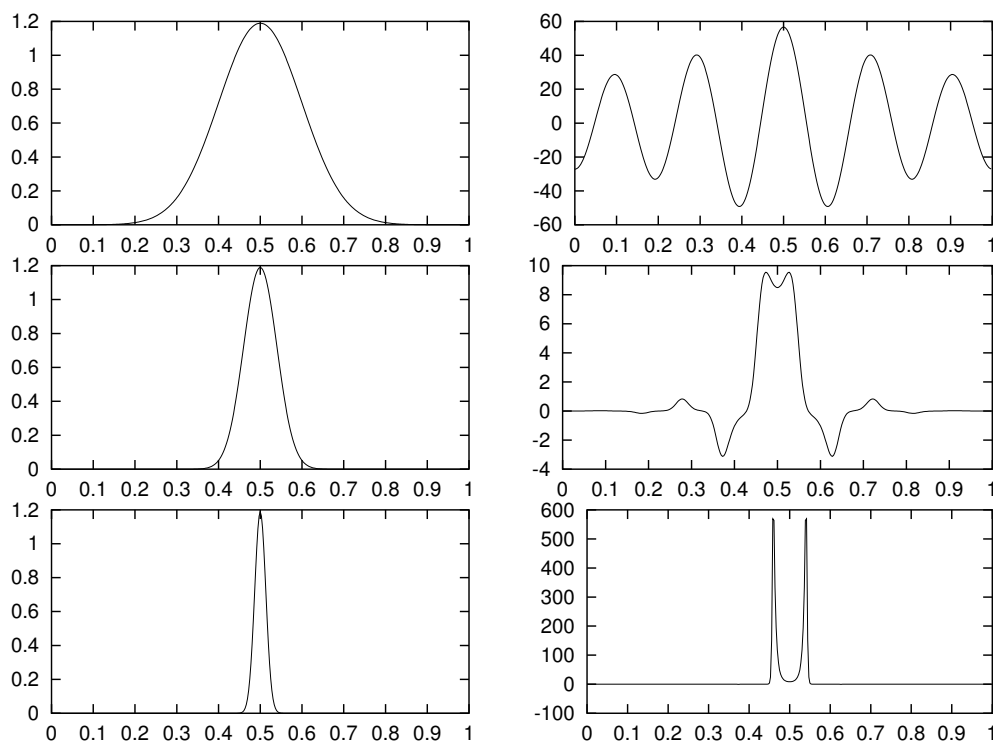


FIGURE 7.1. Comparison of canonical dual windows for three different window functions. The windows are on the right hand side and their corresponding canonical dual windows are on the left hand side. The Gabor frames has the parameters $N = 12$ and $b = 10$. The three windows are $\mathcal{P}\varphi_W$ with values of W of 4, 10 and 30.

The figure shows three sets of windows and dual windows. The parameters N and b have fixed values, and the windows and dual windows are shown for three different values of W .

The first row shows a very wide window with its dual, which has big, regular oscillations spread out across the whole interval. The dual window in the middle is smooth, quite narrow and oscillates a little bit, but the oscillations die out quickly. The last row shows a narrow window. The dual window is narrow as well, and have no oscillations, but instead it has two sharp spikes. From the graph it can not be seen if the window is at all differentiable.

It has not yet be verified that the windows and parameters shown on figure 7.1 actually generate Gabor frames. Since the periodized Gaussians are bounded, the operator $\mathcal{S}$ from (5.9) is bounded by corollary 5.5. To show that $\mathcal{S}$ is also invertible and prove that $\mathbb{G}(\mathcal{P}\varphi_W, \frac{1}{N}, b)$ is a Gabor frame for $L^2(\mathbb{T})$, the test from theorem 5.35 can be done with the parameters for each plot.

The test from theorem 5.35 requires that the minimum eigenvalue of $\mathbf{G}(x)$ is found over all $x \in [0; \frac{1}{N}]$. As a start this can be done by doing a fine grained sampling of the interval $[0; \frac{1}{N}]$, and plotting the result. This is done on figure 7.2. Three of the four plot corresponds to the windows and dual windows on figure 7.1. The figure shows that the minimum eigenvalue is non-zero for all four plots.

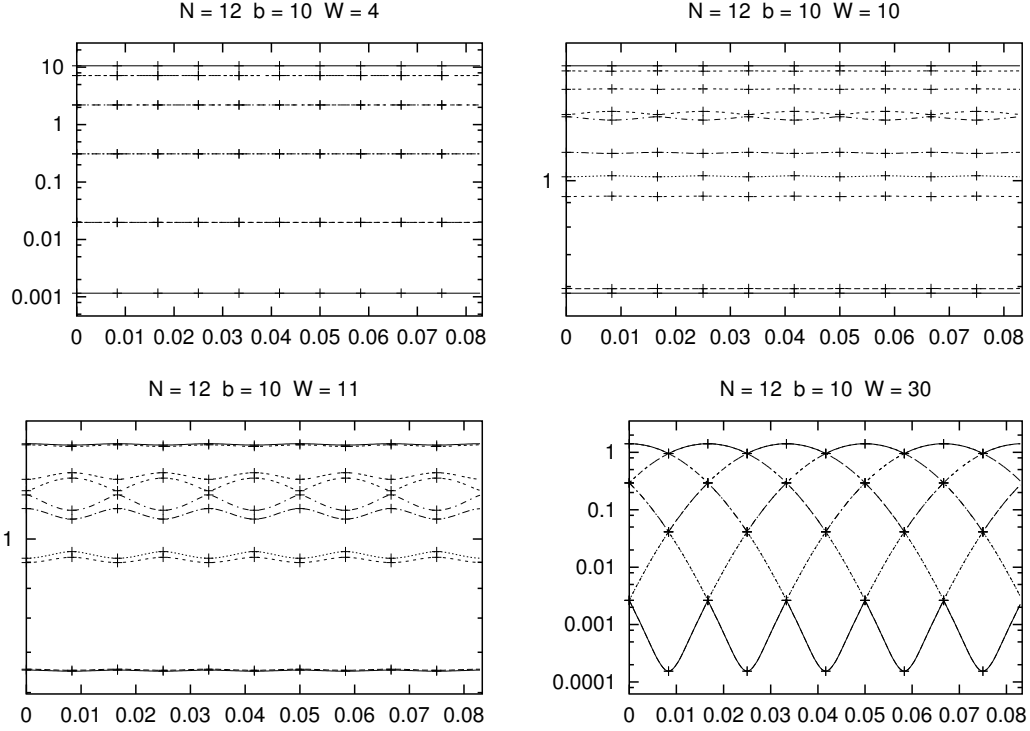


FIGURE 7.2. The figure shows all the eigenvalues of the matrices $\mathbf{G}(x)$ over the range $x \in [0; \frac{1}{N}]$. The window functions used to generate $\mathbf{G}(x)$ are periodized Gaussians of width W .

Since the eigenvalues have only been calculated for certain samples of $x \in [0; \frac{1}{N}]$, then it could be possible that $\mathbf{G}(x)$ for some x not among the samples had a zero eigenvalue. Judging from the shapes of the curves of the eigenvalues it seems highly unlikely. The curves are flat (to within the margin of the machine accuracy) or have regular oscillations. When W is small, corresponding to a wide window, the eigenvalues are constant over all x , and as the width of the pulse decreases the eigenvalues start to oscillate.

Looking at figure 7.2 it seems like the eigenvalues have a periodic behavior over $x \in [0; \frac{1}{N}]$. I have not tried to prove this mathematically, but all the plots I have looked at show the same structure. By counting the periods over a range of parameters N, b and window functions g , I have found that there are always $\frac{b}{\gcd(b, N)}$ periods on one interval $[0; \frac{1}{N}]$, where $\gcd(b, N)$ denotes the greatest common divisor of b and N . Another way of saying this is that there is $\text{lcm}(b, N)$ periods on the interval $[0; 1]$, where $\text{lcm}(b, N)$ is the least common multiple of N and b .

The “sampling” approach for finding the minimum eigenvalue is not adequate, because there is no guarantee that the minimum eigenvalue has actually been found. Two possible solutions are the following:

- Use a numerical method to find the minimum eigenvalue. Since a matrix $\mathbf{G}(x)$ can be constructed for any x , it is a simple matter of using a numerical method to find a local minimum. If the behavior of the eigenvalues are indeed periodic, then considering an

interval $x \in [0; \frac{\gcd(N,b)}{Nb}]$ would suffice, and a local minimum on this interval would most likely be a global minimum.

- Assume that g must satisfy certain nice properties: g should be continuous, symmetric around $x = 0$ and monotone on the interval $[0; \frac{1}{2}]$. Combined with the hypothesis that the eigenvalues are periodic, it would probably be enough to show, that the minimum eigenvalue for all $\mathbf{G}(x)$, $x \in [0; \frac{1}{N}]$ is found for $\mathbf{G}(x)$ with $x = 0$ or $x = \frac{1}{2} \frac{\gcd(N,b)}{Nb}$. This is illustrated on figure 7.2 by the crosses on the lines. They are all placed in points corresponding to $x = 0$ or $x = \frac{1}{2} \frac{\gcd(N,b)}{Nb}$.

In this thesis I have just calculated the minimum eigenvalue of $\mathbf{G}(x)$ over a fine grained sampling of the interval $[0; \frac{1}{N}]$, and made sure that the points include the points just mentioned. This should produce the correct result.

Figure 7.2 show that the minimum eigenvalue is highest for the plots with $W = 10$ and $W = 11$ and considerably lower for the other two. To study this effect figure 7.3 shows the minimum eigenvalue of a range of W .

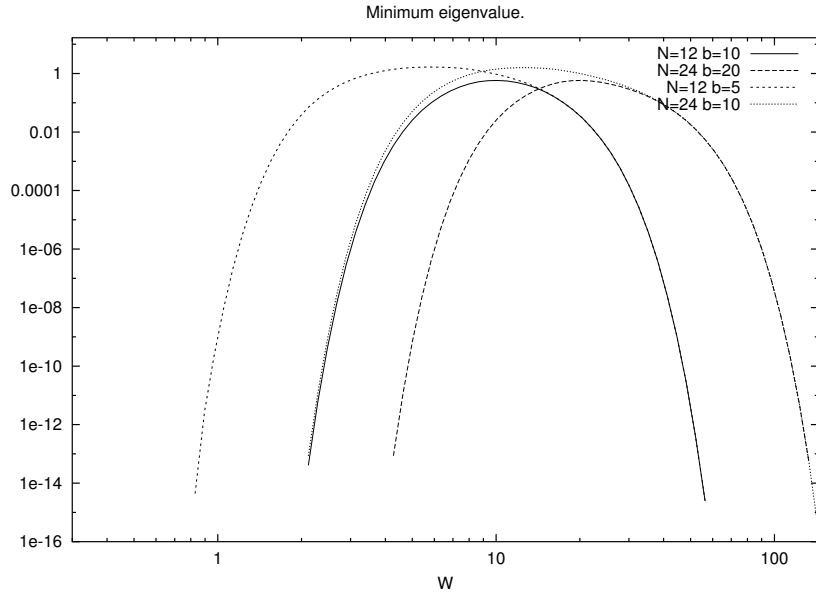


FIGURE 7.3. Plot of the minimum eigenvalue of the matrix G for different sets of parameters N and b over a range of window functions. The window functions are all periodized Gaussians $\mathcal{P}\varphi_W$ and so the eigenvalues are plotted against the value W . Note that the graphs have been calculated over all displayed values of W , but only the non-zero, non-negative results have been displayed.

Figure 7.3 shows results from five tests. The parameters N and b are kept fixed, and the width of the window W is varied. For each set N, b, W the minimum eigenvalue of $\mathbf{G}(x)$ is computed and plotted. The two first curves represent frames with little redundancy $\frac{b}{N} = 0.83$ and the next two represent frames with the double redundancy $\frac{b}{N} = 0.42$.

All five curves have been calculated for values of W in the interval $[0.3; 120]$, but some of the results turn out negative and close to the machine precision, so the graphs do not extend over the whole axis. However, from proposition 5.38 $\mathcal{G}$ cannot have negative eigenvalues, so

the negative results indicates numerical errors, in this case an eigenvalue that is zero within the machine precision.

I cannot tell whether the matrix $\mathbf{G}(x)$ really has a zero eigenvalue, or whether the curves would continue toward zero without hitting it if the computations could be done in infinite precision. Numerically however, it does not make a difference. If $\mathbf{G}(x)$ has an eigenvalue that is numerically zero, then $\mathbf{S}$ cannot be numerically inverted and the dual window cannot be calculated.

The regions where the graphs are not displayed correspond to intervals of W where the periodized Gaussian does not generate a frame for that particular set N, b . The regions are very unstable. The slightest deviation from the shape of the periodized Gaussian, like failing to compute the periodization to a sufficiently high precision, will result in the minimum eigenvalue growing several orders of magnitude. Since I am not acquainted with perturbation theory for Gabor frames, I will not pursue the interesting question about what happens in these regions.

All five graphs on figure 7.3 rise and fall in the same manner. The three windows on figure 7.1 correspond to points on the rise, straight in the middle and on the fall of the graphs. The three dual windows are quite representative for the different varieties of dual windows for the periodized Gaussians. Compare with the two top dual windows in the right column of figure 7.4. The top one represents a point on the fall of a graph, and the middle one represent a

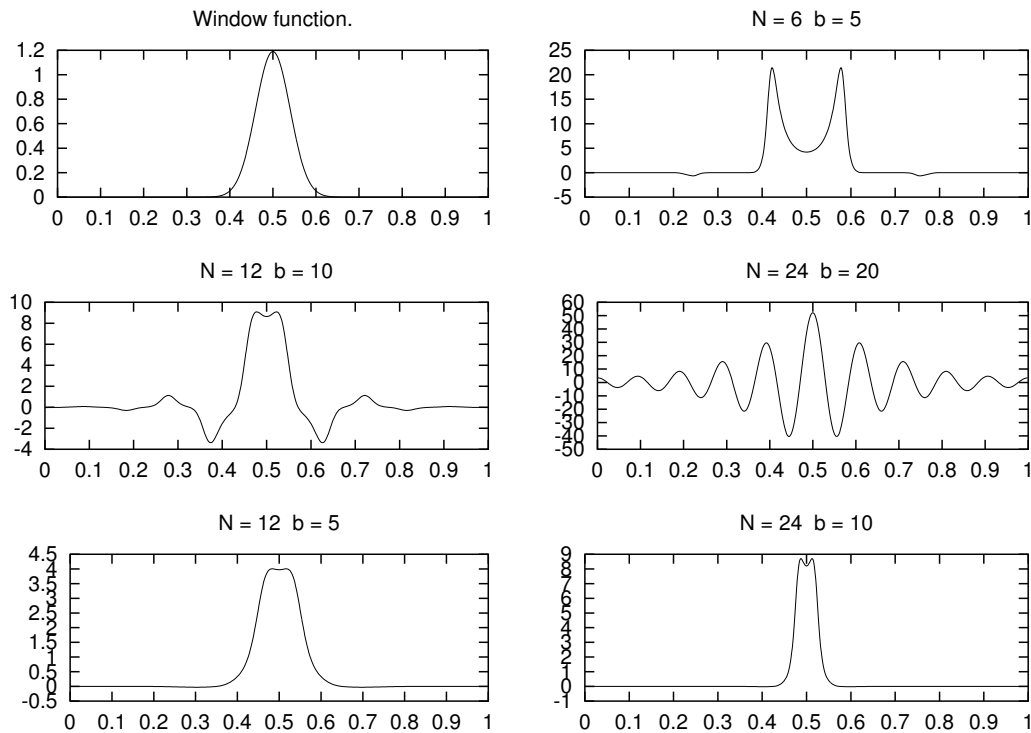


FIGURE 7.4. The figure shows a window function and the corresponding canonical dual window function for different values of N and b . The window function is the periodized Gaussian $\mathcal{P}\varphi_W$ with $W = 10$.

point on the rise of a graph. They have different width and decay properties, but the overall shape is the same.

The rise seems to be dictated by the parameter b and the fall by the parameter N . I have no explanation why the rise seems to be related to b , but the fall of the graphs are quite easy to explain: When W is big compared to N , the translations of g do not seem to cover the whole axis, see figure 7.5.

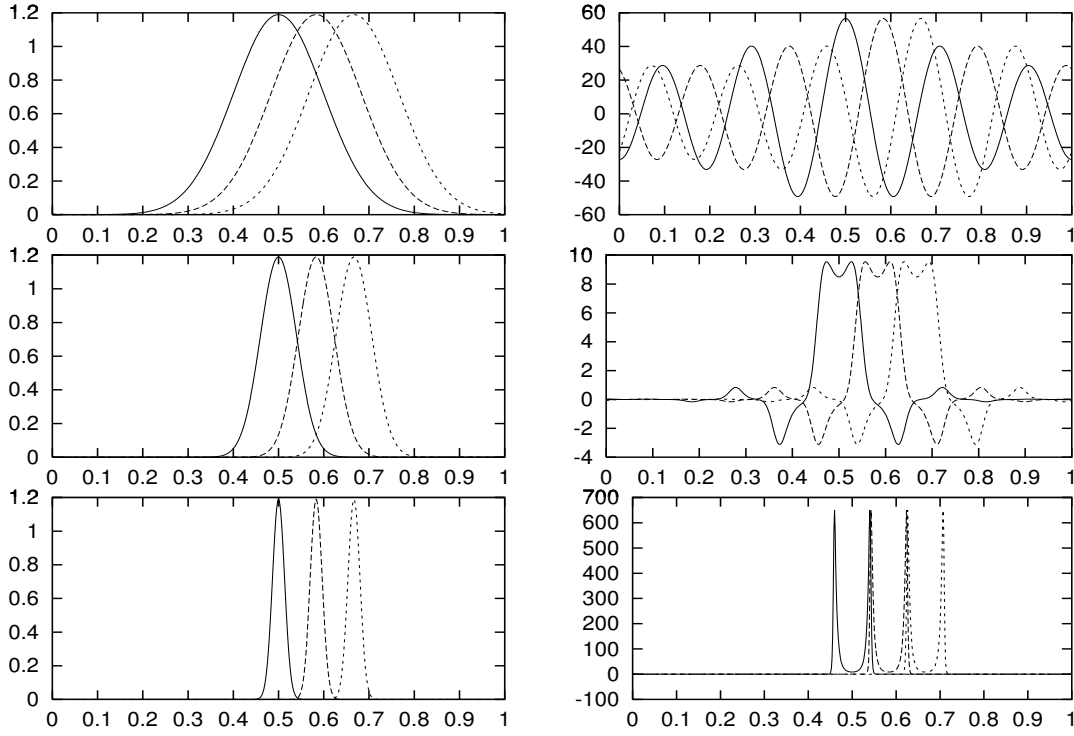


FIGURE 7.5. The figure shows three translations of the windows and dual windows on figure 7.1.

The figure shows how the translations cover the axis. When the window is wide, there is much overlap as seen in the first row. When the window is small, as seen in the last row, there seems to be almost no overlap, and the high peaks of the dual window occurs exactly in the area where the translations of the window functions do not seem to overlap.

Since the Gaussians decay exponentially, they never decay fully to zero. Therefore the window functions are supported on all of $\mathbb{T}$. Therefore the translations overlap no matter how small the width of the function is, and the dual window grows as the size of the overlap shrinks.

Two of the graphs on figure 7.3 correspond to highly redundant frames with densities $\frac{b}{N} = 0.43$. Two dual windows corresponding to points on these graphs can be seen in the last row on figure 7.4. They seem much softer than the other dual windows: They have no oscillations or no sharp peaks. This could be a benefit for applications.

The graphs on figure 7.3 seem identical for identical densities $\frac{b}{N}$ of the frames, just shifted. This also holds true when looking at the dual windows. This effect is displayed on figure 7.6 .

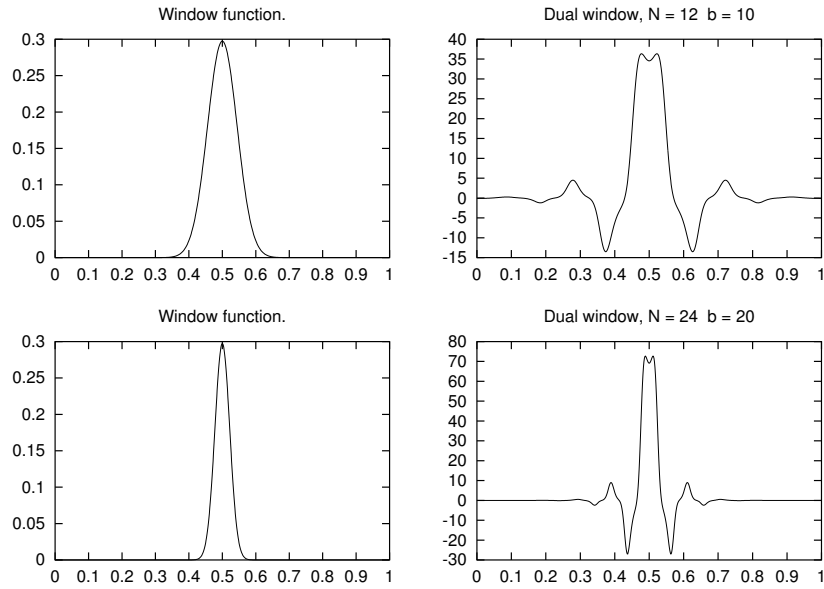


FIGURE 7.6. The figure shows the window and canonical dual window for two Gabor frames, where N , b and the width of the window are scaled in the same manner. Notice that the second dual window seems identical to the first except that it is half as narrow and twice as high. The window function are periodized Gaussians $\mathcal{P}\varphi_W$ with $W = 10$ and $W = 20$.

The figure shows that if W and b, N are scaled the same way, the dual windows look like dilated versions of each other. Since the functions live on the circle, $\gamma \in L^2(\mathbb{T})$, then true dilations are not possible, but the effect seems real because the region where the functions differ consist of values close to zero.

This effect shows that when working with two Gabor frames $\mathfrak{G}$ and $\mathfrak{H}$ as defined in theorem 5.27, there is two main ways of increasing the sampling rate L :

- Keeping b , N and the window width W fixed while increasing L . This will retain the window g and dual window γ of $\mathfrak{G}$, and these windows will be better resolved as they become sampled more finely. Since $L = bM$ this corresponds to increasing the number of modulations in the system $\mathfrak{H}$.
- Letting b , N and the window width W scale with L . This will narrow the window g and the dual window γ , but the number of samples across the main part of the windows remains constant. Because $L = bM$ this corresponds to keeping the number of modulations in the system $\mathfrak{H}$ constant and increasing the number of translations.

Since theorem 6.9 shows that the precision of the Gabor derivative increases with the number of modulations, the first method of scaling a Gabor system will be used for testing. The second method of scaling is well suited for very long signals, like musical signals, where all oscillations have a relatively short duration compared to the length of the signal.

The figure 7.7 shows some completely different windows and their corresponding dual windows. The windows have been constructed from indicator functions of different width. None of the widths divides any of the numbers a, b, M, N or L . The graphs have been computed with very many points (14400 points along the x-axis), so the structures of the dual windows

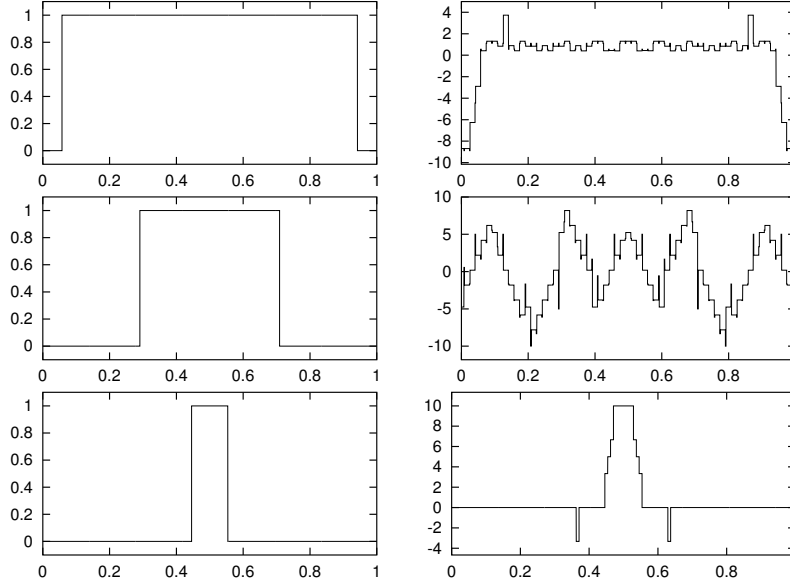


FIGURE 7.7. The figure shows three windows constructed from the indicator function, and their corresponding canonical dual windows. The windows have width 0.88, 0.42 and 0.11, and $N = 12$ and $b = 10$.

are well resolved, and not a result of insufficient computations. The figure just serves an illustrative purpose. I will not use these windows for any computations, or make any statements of whether they really do generate Gabor frames.

7.2. Differentiation. Theorem 6.9 provides an order of accuracy for the Gabor derivative as a function of the smoothness of f , if $f \in C^1(\mathbb{T})$ and $f^{(2)}$ is piecewise monotone and if the canonical dual window is smooth. From looking at the dual windows on figure 7.1 and figure 7.4 some of them look smooth. Though I have not found a mathematical proof to justify that this holds true. Some results of smoothness of the canonical dual window for Gabor frames for $L^2(\mathbb{R}^d)$ are presented in [Groch01, chapter 13.] and in [Jans95b], but I have chosen not to transfer these results to the $L^2(\mathbb{T})$ setting.

A comparison of the accuracy of various numerical methods for differentiation can be seen on figure 7.8. It compares second order polynomial differentiation, fourth order polynomial differentiation, spectral differentiation and differentiation using Gabor frames. The smooth test function is $f(x) = e^{\sin(2\pi x)}$. It is shown on figure 7.9 along with its analytical derivative. The accuracy of the second and fourth order methods are clearly visible. The spectral method using a DFT converges very quickly in the beginning, and then grows slowly. The reason for the slow growth of the spectral method is an accumulation of numerical round-off errors. The Gabor derivative starts out terrible but still converges with faster than polynomial growth, until it is stopped by numerical round-off errors at the same level as the spectral method. The faster than polynomial growth is predicted by theorem 6.9, so this serves as an illustration of that.

Figure 7.10 shows how the accuracy of the Gabor derivative change when the width of the periodized Gaussian is varied. For small values of W the method is very accurate in the beginning, but fails to reach the full precision, and starts to diverge. For larger values the

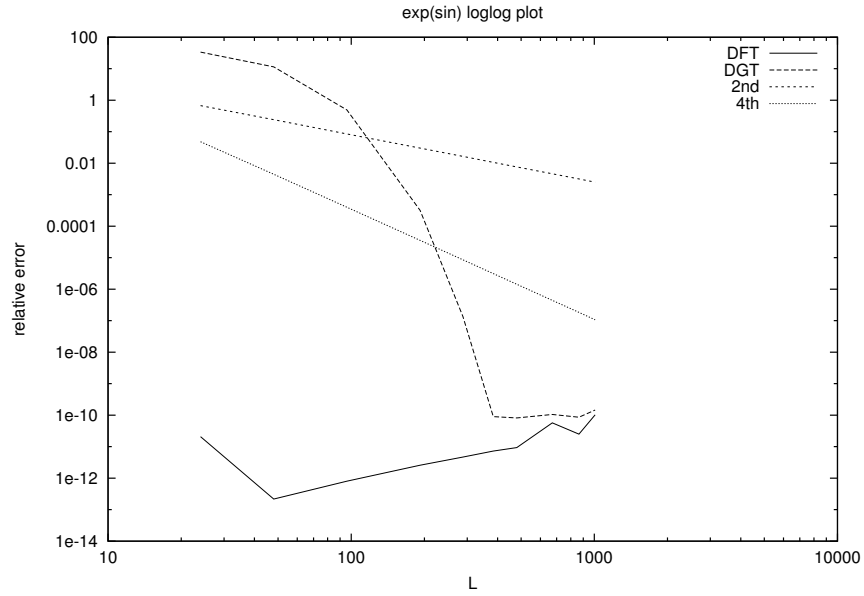


FIGURE 7.8. Comparison of the accuracy of numerically differentiating the function $f(x) = e^{\sin(2\pi x)}$, $x \in [0; 1]$. The Gabor frames used has the parameters $N = 8$, $b = 6$ and the window is a periodized Gaussian with $W = 4$. The sampling rate L is placed on the x -axis.

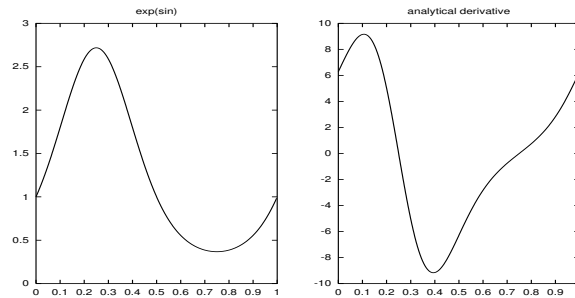


FIGURE 7.9. The figure shows the function $f(x) = e^{\sin(2\pi x)}$, $x \in [0; 1]$ and its analytical derivative.

method starts out much worse, but reaches the full precision at some point. The canonical dual windows corresponding to the first four lines can be seen on figure 7.11. Perhaps a little surprising the parameters N, b, W of the best working method does not correspond to a maximum point on the graph of figure 7.3, but instead to a point a little to the left. This is also shown on figure 7.11: The dual window corresponding to the best working method (number three on the plot) has oscillations. This certainly makes the dual Gabor atoms resemble the complex exponentials as much as possible, and since the spectral method is always the most accurate for smooth functions, it probably explains why this particular dual window seems to work best.

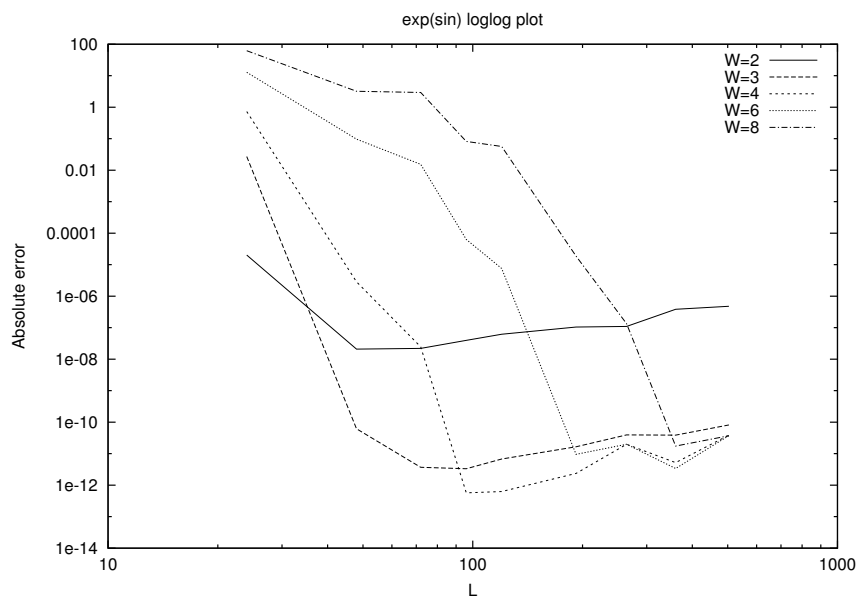


FIGURE 7.10. The figure shows the accuracy of the Gabor derivative for different choices of window functions. The window functions are periodized Gaussians of width W , as shown in the plot. The testing function is $f(x) = e^{\sin(2\pi x)}$, $x \in [0; 1]$.

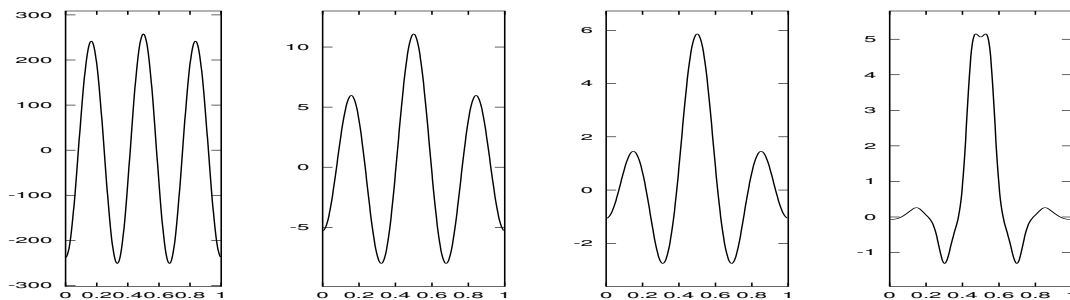


FIGURE 7.11. The figure shows the canonical dual windows corresponding to the first four lines on figure 7.10. The Gabor frames have parameters $N = 8$, $b = 6$ and W equal to 2, 3, 4 and 6.

The only test function considered so far has been smooth. Some non-smooth test functions are displayed on figure 7.12.

The first test function is discontinuous. The second is not differentiable, and the third is differentiable, but not two times differentiable. The third function are two parabolas combined. The rate of convergence is strongly dependent on the smoothness of the function for all three functions, and the Gabor derivative is no better or worse than the other methods. Since the two first functions are not even differentiable, it does not make much sense to calculate the global error, because it will be dominated by what happens in the discontinuity, which is not properly defined anyway.

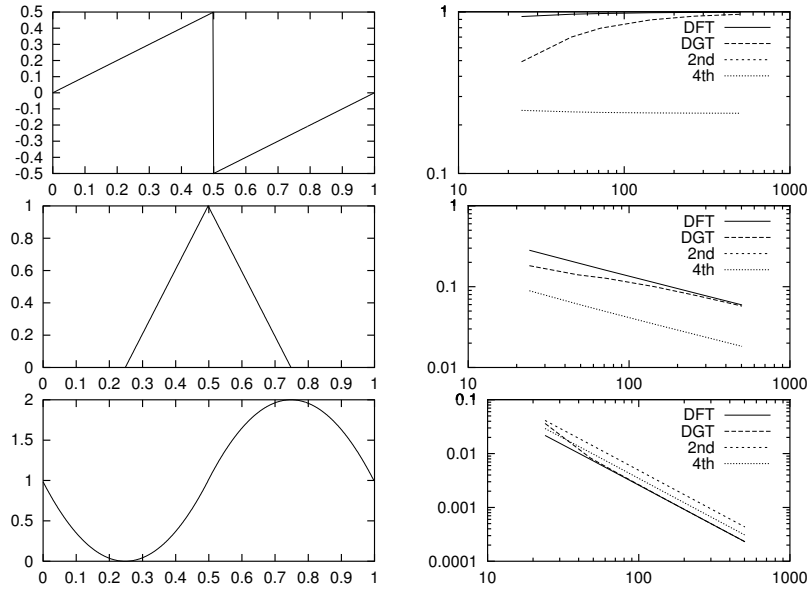


FIGURE 7.12. The left column shows three non-smooth test functions, and the right column shows the relative, global error of using the different numerical methods on that particular test function. The Gabor frames used are the same as in figure 7.8.

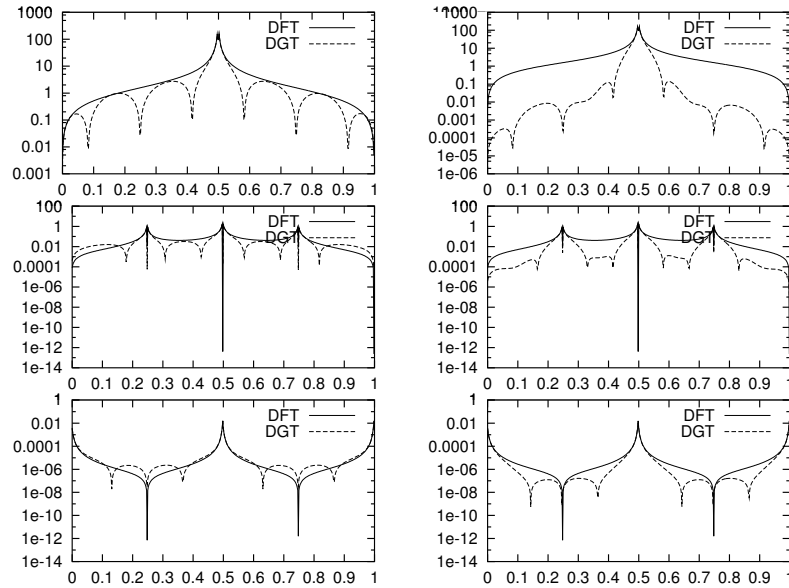


FIGURE 7.13. The figure shows the pointwise error of the numerical differentiation over the interval $[0;1]$. The test functions in each row are the same as on figure 7.12. In the left column a Gabor frame with parameters $N = 8$, $b = 6$ and $W = 4$ are used, and for the right column $W = 10$ instead. The sampling rate is $L = 504$.

Figure 7.13 shows the pointwise error over the interval $[0;1]$. At the discontinuities both methods make exactly the same error. The interesting thing is what happens away from the discontinuities. Away from the discontinuity the spectral method has slowly decaying errors. This corresponds to oscillations around the correct solution - an effect of the so-called Gibbs phenomena. The Gabor method does not do much better on the plots in the left column. The error decays much faster, but returns again to the level of the spectral method. This is probably due to the wide support of the dual window, see figure 7.14. Since the dual window never decay to zero, then all translations of the dual window overlap the discontinuity.

On the plots in the right hand column, the window of the Gabor frame have much more narrow support, and this changes the situation. For the first two examples the Gabor method is significantly better at reducing the oscillations away from the discontinuities. For the third example the spectral method starts to catch up.

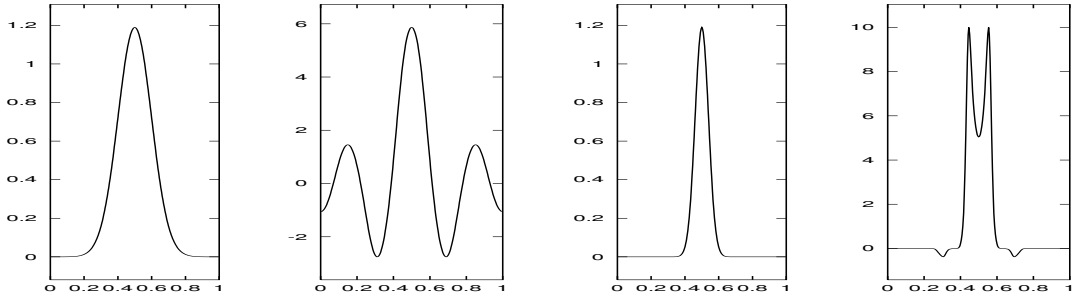


FIGURE 7.14. Two pairs of windows and dual windows corresponding to the one used in figure 7.13. The first pair has parameters $N = 8$, $b = 6$, and $W = 4$. The second has parameters $N = 8$, $b = 6$ and $W = 10$.

This example shows that for the same parameters N, b , the optimal window seems to be dependent on the test function. It illustrates that the accuracy of the Gabor derivative is dependent on N, b , the window function and on the test function.

This was just a short introduction to the Gabor derivative. A deeper understanding of the dependency of the accuracy on N, b , the window function and on the test function requires a much more detailed study.

8. CONCLUDING REMARKS.

This project started out dealing only with differentiation, and ended up dealing also with Gabor frames for periodic spaces.

The spaces $L^2(\mathbb{T})$ and $\mathbb{C}_p^L$ seems to be a very good setting for Gabor frame theory. The structure of Gabor frames for $\mathbb{C}_p^L$ allows for efficient computations, and Janssens sampling theorem shows that they can be used for approximation of Gabor frames for $L^2(\mathbb{T})$. The space $L^2(\mathbb{T})$ is infinite dimensional, but compared to $L^2(\mathbb{R})$ it is much simpler. It is easy to find windows that satisfy condition A, and it is also relatively easy to verify that a Gabor system is a frame. These aspects are much more delicate in $L^2(\mathbb{R})$. I find the method of considering eigenvalues to show that a Gabor system is a frame, especially interesting. If it can be shown for a class of nice functions that the minimum eigenvalue can be found for a particular and explicitly known x , the method would be both efficient and accurate.

Furthermore the two types of Gabor frames resemble Fourier series and the discrete Fourier transform. In this thesis this relationship was used to replace the DFT in the computation of the spectral derivative, and to there by get a new method of differentiating. There might be other areas where the same substitution could be done.

The Gabor derivative has the same convergence properties as the spectral derivative, but it generally requires a higher sampling rate to converge to the same margin of error. I cannot say whether it can be useful for some applications. It is never more accurate than the spectral derivative. It does handle discontinuities better, but the error at the discontinuity is the exact same. It is far more flexible, because there are so many parameters to vary, but this also means that without some detailed study one might chose some very suboptimal parameters.

However, a new numerical method for differentiation with speed and accuracy close to that of spectral differentiation is here by founded.

Peter Lempel Søndergaard.

REFERENCES

- [BuNe71] Butzer, P.L, Nessel, R.J. *Fourier Analysis and Approximation, vol 1*. Birkhäuser Verlag Basel, Stuttgart 1971.
- [Briggs95] Briggs, William L. *The DFT, An Owner's Manual for the Discrete Fourier Transform*. SIAM Philadelphia 1995.
- [Chr02] Christensen, Ole. *An Introduction to Frames and Riesz Bases*. to appear.
- [Groech98] Gröchenig, Karlheinz. Aspects of Gabor Analysis on Locally Compact Abelian Groups. In *Gabor Analysis and Algorithms*, page 211 - 231. Birkhäuser 1998.
- [Groch01] Gröchenig, Karlheinz. *Foundations of Time-Frequency Analysis*. Birkhäuser 2001.
- [Orr93] Orr, Richard s. *The Order of Computation or Finite Discrete Gabor Transforms*. IEEE Transactions on Signal Processing. vol 41. no 1. 1993
- [Jans95a] Janssen, A.J.E.M. *On Rationally Oversampled Weyl-Heisenberg Frames*. Signal Processing. vol. 47, page 239 - 245. 1995
- [Jans95b] Janssen, A.J.E.M. Duality and biorthogonality for Weyl-Heisenberg frames. J. Fourier Anal. Appl. 1(4), page 403-436. 1995
- [Jans97] Janssen, A.J.E.M. *From Continuous to Discrete Weyl-Heisenberg Frames Through Sampling*. The Journal of Fourier Analysis and Applications, volume 3 number 5, 1997.
- [Jens92] Jensen, Helge Elbrønd *Matematisk Analyse, bind 4*. DTH 1992
- [Qiu98] Qiu, Sigang. *Discrete Gabor Transforms: The Gabor-Gram Matrix Approach*. The Journal of Fourier Analysis and Applications. vol 4, issue 1, 1998
- [Rud66] Rudin, Walter. *Real and complex analysis*. McGraw-Hill 1966.
- [Soen02] Søndergaard, Peter *Gabor Frames and Pseudo-Differential Operators*. Forprojekt, MAT, DTU 2002.
- [Stroh98] Strohmer, Thomas. Numerical Algorithms for Discrete Gabor Expansions. in *Gabor Analysis and Algorithms*, page 267 - 294. Birkhäuser 1998.
- [Tref] Trefethen, Lloyd N. *Spectral Methods in Matlab*. SIAM 1998.
- [vL92] van Loan, Charles. *Computational Frameworks for The Fast Fourier Transform*. SIAM, Philadelphia 1992.