

AD-A249 029



DTIC
ELECTE
APR 23 1992
S C D

①



On the Rapid Evaluation of Trigonometric Series

A. Dutt and V. Rokhlin

Research Report YALEU/DCS/RR-893
March 1992

by
D1
A

D13

A

92-10445



EXEMPTED FROM GDS

Approved for public release;
Distribution Unlimited

YALE UNIVERSITY
DEPARTMENT OF COMPUTER SCIENCE

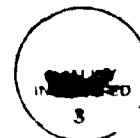
92 4 92 147

A group of algorithms is presented generalizing the Fast Fourier Transform to the case of non-integer frequencies and non-equispaced nodes on the interval $[-\pi, \pi]$. The schemes of this paper are based on a combination of certain analytical considerations with the classical Fast Fourier Transform, and generalize both the forward and backward FFTs. Each of the algorithms requires $O(N \cdot \log N + N \cdot \log(1/\varepsilon))$ arithmetic operations, where ε is the precision of computations and N is the number of nodes. The efficiency of the approach is illustrated by several numerical examples.

On the Rapid Evaluation of Trigonometric Series

A. Dutt and V. Rokhlin
Research Report YALEU/DCS/RR-893
March 1992

Approved For		☒
V. D. Dutt		
V. D. Dutt		
Approved		
Justification		
By		
Distribution/		
Availability Codes		
Dist	Avail and/or Special	
A-1		



The authors were supported in part by the Office of Naval Research under Grant N00014-89-J-1527 and in part by the National Science Foundation under Grant DMS9012751.

Approved for public release: distribution is unlimited.

Keywords: FFT, Trigonometric Series, Fourier Analysis, Interpolation, Approximation Theory

1 Introduction

Fourier techniques have been a popular analytical tool in the study of physics and engineering for more than two centuries. A reason for the usefulness of such techniques is that the trigonometric functions $e^{i\omega x}$ are eigenfunctions of the differentiation operator and can be effectively used to model solutions of differential equations which arise in the fields mentioned above.

More recently, the arrival of digital computers and the development of the Fast Fourier Transform (FFT) algorithm in the 1960s have established Fourier analysis as a powerful and practical numerical tool. The FFT, which computes discrete Fourier transforms (DFTs), is now central to many areas, most notably spectral analysis and signal processing. In some applications however, the input data is not uniformly spaced, a condition which is required for the FFT. In this paper we present a set of algorithms for computing more efficiently some generalizations of the DFT, namely the forward and inverse transformations described by the equations

$$f_j = \sum_{k=0}^N \alpha_k \cdot e^{i\omega_k x_j} \quad (1)$$

for $j = 0, \dots, N$, where $f_j \in \mathbb{C}$, $\alpha_k \in \mathbb{C}$, $\omega_k \in [-N/2, N/2]$ and $x_j \in [-\pi, \pi]$. Each algorithm requires a number of arithmetic operations proportional to

$$N \cdot \log N + N \cdot \log \left(\frac{1}{\varepsilon} \right) \quad (2)$$

where ε is the desired accuracy, compared with $O(N^2)$ operations required for the direct application and $O(N^3)$ for the direct inversion.

Remark 1.1 The DFT in "unaliased" form, as described by the equations

$$f_j = \sum_{k=-N/2}^{N/2-1} \alpha_k \cdot e^{2\pi i k j / N}, \quad (3)$$

is clearly a special case of (1); the FFT algorithm reduces the number of operations for the DFT from $O(N^2)$ to $O(N \cdot \log N)$ by a sequence of algebraic manipulations. In the more general case of (1), the structure of the linear transformation is also exploitable, and the algorithms of this paper combine certain analytical results with the existing FFT.

The plan of the paper is as follows. We start in Section 2 with some results from analysis and approximation theory which are used in the design of the algorithms. An exact statement of the problem in Section 3 is then followed by informal descriptions of the algorithms in Section 4. In Section 5 we introduce some notation which is used in a set of more detailed algorithm descriptions in Section 6. Six numerical examples are presented in Section 7 to illustrate the performance of the schemes. Finally, Section 8 lists some generalizations and conclusions.

2 Mathematical and Numerical Preliminaries

2.1 Elementary Analytical Tools

In this subsection we summarize some well-known results to be used in the remainder of the paper. Lemmas 2.1 and 2.2 are obvious, and Lemmas 2.3 and 2.4 can be found, for example, in [3].

Lemma 2.1 *For any real c ,*

$$\int_{-\pi}^{\pi} e^{icx} dx = \frac{2}{c} \sin(c\pi). \quad (4)$$

Lemma 2.2 *For any integer k ,*

$$\frac{1}{2\pi} \int_{-\pi}^{\pi} e^{ikx} dx = \begin{cases} 1 & \text{if } k = 0 \\ 0 & \text{otherwise} \end{cases} \quad (5)$$

Lemma 2.3 *For any real $b > 0$ and complex z ,*

$$\int_{-\infty}^{\infty} e^{-bx^2} \cdot e^{zx} dx = \sqrt{\frac{\pi}{b}} \cdot e^{z^2/4b}. \quad (6)$$

Lemma 2.4 *For any real $b > 0$ and $a > 0$,*

$$\int_a^{\infty} e^{-bx^2} dx < \frac{e^{-ba^2}}{2ba}. \quad (7)$$

Proof.

$$\int_a^{\infty} e^{-bx^2} dx = \int_0^{\infty} e^{-b(x+a)^2} dx < e^{-ba^2} \int_0^{\infty} e^{-2bax} dx = \frac{e^{-ba^2}}{2ba}. \quad (8)$$

□

2.2 Relevant Facts from Approximation Theory

The principal tool of this paper is a somewhat detailed analysis of Fourier series of functions $\phi : [-\pi, \pi] \rightarrow \mathbb{C}$ given by the formula

$$\phi(x) = e^{-bx^2} \cdot e^{icx} \quad (9)$$

where $b > \frac{1}{2}$ and c are real numbers. We present this analysis in the Lemmas and Theorems of this subsection, numbered 2.5–2.10.

Lemmas 2.5 and 2.6, provide two inequalities which are used in Theorem 2.7. Theorems 2.7–2.9 are intermediate results leading to Theorem 2.10, which explains how to approximate functions of the form e^{icx} using a small number of terms, and is the principal result of this section. For all approximations we derive error bounds which allow us to perform numerical computations to any specified accuracy.

Lemma 2.5 For any real $b > \frac{1}{2}$, c and any integer k ,

$$\left| 2 \int_{\pi}^{\infty} e^{-bx^2} \cos((c-k)x) dx + e^{-b\pi^2} \cdot \int_{-\pi}^{\pi} e^{i(c-k)x} dx \right| < 2\pi e^{-b\pi^2} \cdot \left(1 + \frac{1}{\pi^2}\right). \quad (10)$$

Proof. Using the triangle inequality and Lemma 2.4 we have

$$\begin{aligned} \left| 2 \int_{\pi}^{\infty} e^{-bx^2} \cdot \cos((c-k)x) dx + e^{-b\pi^2} \cdot \int_{-\pi}^{\pi} e^{i(c-k)x} dx \right| &\leq 2 \int_{\pi}^{\infty} e^{-bx^2} dx + 2\pi e^{-b\pi^2} \\ &< 2\pi e^{-b\pi^2} \cdot \left(\frac{1}{2b\pi^2} + 1\right) \\ &< 2\pi e^{-b\pi^2} \cdot \left(\frac{1}{\pi^2} + 1\right). \end{aligned} \quad (11)$$

□

Lemma 2.6 For any real $b > \frac{1}{2}$, c and any integer k ,

$$\left| 2 \int_{\pi}^{\infty} e^{-bx^2} \cos((c-k)x) dx + e^{-b\pi^2} \cdot \int_{-\pi}^{\pi} e^{i(c-k)x} dx \right| < \frac{8b\pi e^{-b\pi^2}}{(c-k)^2} \cdot \left(1 + \frac{1}{\pi^2}\right). \quad (12)$$

Proof. Integrating by parts we have

$$\begin{aligned} &2 \int_{\pi}^{\infty} e^{-bx^2} \cdot \cos((c-k)x) dx \\ &= \frac{2}{c-k} \left[e^{-bx^2} \sin((c-k)x) \right]_{\pi}^{\infty} + \frac{4b}{c-k} \int_{\pi}^{\infty} x e^{-bx^2} \sin((c-k)x) dx \\ &= -\frac{2}{c-k} e^{-b\pi^2} \sin((c-k)\pi) + \frac{4b}{c-k} \int_{\pi}^{\infty} x e^{-bx^2} \sin((c-k)x) dx. \end{aligned} \quad (13)$$

After rearranging the terms in (13) and integrating by parts again we obtain

$$\begin{aligned} &\left| 2 \int_{\pi}^{\infty} e^{-bx^2} \cos((c-k)x) dx + \frac{2e^{-b\pi^2}}{c-k} \sin((c-k)\pi) \right| \\ &= \left| \frac{4b}{c-k} \int_{\pi}^{\infty} x e^{-bx^2} \sin((c-k)x) dx \right| \\ &= \left| -\frac{4b}{(c-k)^2} \left(\left[x e^{-bx^2} \cos((c-k)x) \right]_{\pi}^{\infty} - \int_{\pi}^{\infty} (1-2bx^2) e^{-bx^2} \cos((c-k)x) dx \right) \right| \\ &\leq \frac{4b}{(c-k)^2} \left(\pi e^{-b\pi^2} + \int_{\pi}^{\infty} e^{-bx^2} dx + \int_{\pi}^{\infty} x \cdot 2bx e^{-bx^2} dx \right) \\ &< \frac{4b}{(c-k)^2} \left(\pi e^{-b\pi^2} + \int_{\pi}^{\infty} e^{-bx^2} dx + \left[-x e^{-bx^2} \right]_{\pi}^{\infty} + \int_{\pi}^{\infty} e^{-bx^2} dx \right). \end{aligned} \quad (14)$$

Finally, due to (14) and Lemmas 2.1 and 2.4 we have

$$\begin{aligned} \left| 2 \int_{\pi}^{\infty} e^{-bx^2} \cos((c-k)x) dx + e^{-b\pi^2} \cdot \int_{-\pi}^{\pi} e^{i(c-k)x} dx \right| &< \frac{4be^{-b\pi^2}}{(c-k)^2} \cdot \left(2\pi + \frac{2}{2b\pi} \right) \\ &< \frac{8b\pi e^{-b\pi^2}}{(c-k)^2} \cdot \left(1 + \frac{1}{\pi^2} \right). \end{aligned} \quad (15)$$

□

The following theorem provides an explicit expression for the coefficients of a Fourier series which approximates functions of the form (9).

Theorem 2.7 Let $\phi(x) = e^{-bx^2} e^{icx}$ for any real $b > \frac{1}{2}, c$. Then, for any $x \in (-\pi, \pi)$,

$$\left| \phi(x) - \sum_{k=-\infty}^{\infty} \rho_k e^{ikx} \right| < e^{-b\pi^2} \cdot \left(4b + \frac{70}{9} \right), \quad (16)$$

where

$$\rho_k = \frac{1}{2\sqrt{b\pi}} e^{-(c-k)^2/4b} \quad (17)$$

for $k = -\infty, \dots, \infty$.

Proof. We denote by σ_k the k -th Fourier coefficient for ϕ , so that for $x \in (-\pi, \pi)$,

$$\phi(x) = \sum_{k=-\infty}^{\infty} \sigma_k e^{ikx}, \quad (18)$$

and due to Lemma 2.3 and equation (17) we have

$$\begin{aligned} \sigma_k &= \frac{1}{2\pi} \int_{-\pi}^{\pi} \phi(x) e^{-ikx} dx \\ &= \frac{1}{2\pi} \left(\int_{-\infty}^{\infty} e^{-bx^2} e^{icx} e^{-ikx} dx - \int_{-\infty}^{-\pi} e^{-bx^2} e^{icx} e^{-ikx} dx - \int_{\pi}^{\infty} e^{-bx^2} e^{icx} e^{-ikx} dx \right) \\ &= \frac{1}{2\pi} \left(\sqrt{\frac{\pi}{b}} \cdot e^{-(c-k)^2/4b} + \int_{\infty}^{\pi} e^{-bx^2 - icx + ikx} dx - \int_{\pi}^{\infty} e^{-bx^2 + icx - ikx} dx \right) \\ &= \rho_k - \frac{1}{\pi} \int_{\pi}^{\infty} e^{-bx^2} \cos((c-k)x) dx. \end{aligned} \quad (19)$$

Rearranging equation (19) and applying Lemmas 2.5 and 2.6 we obtain the inequalities

$$\left| \sigma_k - \rho_k - \frac{e^{-b\pi^2}}{2\pi} \int_{-\pi}^{\pi} e^{icx} e^{-ikx} dx \right| < e^{-b\pi^2} \cdot \left(1 + \frac{1}{\pi^2} \right), \quad (20)$$

$$\left| \sigma_k - \rho_k - \frac{e^{-b\pi^2}}{2\pi} \int_{-\pi}^{\pi} e^{icx} e^{-ikx} dx \right| < \frac{4be^{-b\pi^2}}{(c-k)^2} \cdot \left(1 + \frac{1}{\pi^2} \right), \quad (21)$$

and it now follows from the combination of (18), (20) and (21) that, for any $x \in (-\pi, \pi)$,

$$\begin{aligned}
& \left| \phi(x) - \sum_{k=-\infty}^{\infty} \rho_k e^{ikx} - e^{-b\pi^2} \cdot e^{icx} \right| \\
&= \left| \sum_{k=-\infty}^{\infty} e^{ikx} \cdot \left(\sigma_k - \rho_k - \frac{e^{-b\pi^2}}{2\pi} \int_{-\pi}^{\pi} e^{icx} e^{-ikx} dx \right) \right| \\
&< \sum_{k, |c-k| \geq 3} \frac{4be^{-b\pi^2}}{(c-k)^2} \cdot \left(1 + \frac{1}{\pi^2} \right) + \sum_{k, |c-k| < 3} e^{-b\pi^2} \cdot \left(1 + \frac{1}{\pi^2} \right) \\
&< 4be^{-b\pi^2} \cdot \frac{9}{8} \cdot 2 \cdot \sum_{k=3}^{\infty} \frac{1}{k^2} + 6e^{-b\pi^2} \cdot \frac{10}{9}.
\end{aligned} \tag{22}$$

Some elementary analysis yields

$$\sum_{k=3}^{\infty} \frac{1}{k^2} < \frac{1}{9} + \int_3^{\infty} \frac{dx}{x^2} = \frac{1}{9} + \frac{1}{3} = \frac{4}{9} \tag{23}$$

and substituting (23) into (22) we have

$$\left| \phi(x) - \sum_{k=-\infty}^{\infty} \rho_k e^{ikx} - e^{-b\pi^2} \cdot e^{icx} \right| < e^{-b\pi^2} \cdot \left(4b + \frac{60}{9} \right). \tag{24}$$

To complete the proof we make use of the triangle inequality and (24) to obtain

$$\begin{aligned}
\left| \phi(x) - \sum_{k=-\infty}^{\infty} \rho_k e^{ikx} \right| &\leq \left| \phi(x) - \sum_{k=-\infty}^{\infty} \rho_k e^{ikx} - e^{-b\pi^2} \cdot e^{icx} \right| + \left| e^{-b\pi^2} \cdot e^{icx} \right| \\
&< e^{-b\pi^2} \cdot \left(4b + \frac{70}{9} \right).
\end{aligned} \tag{25}$$

□

Remark 2.1 According to Theorem 2.7, functions of the form $e^{-bx^2} e^{icx}$ can be approximated by a Fourier series whose coefficients are given analytically, and the error of the approximation decreases exponentially as b increases.

Remark 2.2 The coefficients ρ_k as defined by (17) have a peak at $k = [c]$, the nearest integer to c , and decay exponentially as $k \rightarrow \pm\infty$. We keep only the $q+1$ largest coefficients, where the integer q is chosen such that

$$q \geq 4b\pi, \tag{26}$$

so as to satisfy the inequality

$$e^{-(q/2)^2/4b} \leq e^{-b\pi^2}. \tag{27}$$

The following theorem estimates the truncation error under the conditions of Remark 2.2 and thus provides a way of approximating functions of the form (9) by a q -term series.

Theorem 2.8 Let $\phi(x) = e^{-bx^2} e^{icx}$ for any real $b > \frac{1}{2}$, c , and let q be an even integer such that $q \geq 4b\pi$. Then, for any $x \in (-\pi, \pi)$,

$$\left| \phi(x) - \sum_{k=[c]-q/2}^{[c]+q/2} \rho_k e^{ikx} \right| < e^{-b\pi^2} \cdot (4b + 9), \quad (28)$$

where $\{\rho_k\}$ are defined by (17).

Proof. For any $x \in (-\pi, \pi)$,

$$\left| \phi(x) - \sum_{k=[c]-q/2}^{[c]+q/2} \rho_k e^{ikx} \right| \leq \left| \phi(x) - \sum_{k=-\infty}^{\infty} \rho_k e^{ikx} \right| + \left| \sum_{k>[c]+q/2} \rho_k e^{ikx} \right| + \left| \sum_{k<[c]-q/2} \rho_k e^{ikx} \right|. \quad (29)$$

Due to (17) and the triangle inequality we have the inequalities

$$\left| \sum_{k>[c]+q/2} \rho_k e^{ikx} \right| \leq \sum_{k=[c]+q/2+1}^{\infty} \frac{e^{-(c-k)^2/4b}}{2\sqrt{b\pi}} < \sum_{k=q/2}^{\infty} \frac{e^{-k^2/4b}}{\sqrt{2\pi}}, \quad (30)$$

$$\left| \sum_{k<[c]-q/2} \rho_k e^{ikx} \right| \leq \sum_{k=-\infty}^{[c]-q/2-1} \frac{e^{-(c-k)^2/4b}}{2\sqrt{b\pi}} < \sum_{k=q/2}^{\infty} \frac{e^{-k^2/4b}}{\sqrt{2\pi}}. \quad (31)$$

Some elementary analysis and an application of Lemma 2.4 yields

$$\sum_{k=q/2}^{\infty} e^{-k^2/4b} < e^{-(q/2)^2/4b} + \int_{q/2}^{\infty} e^{-x^2/4b} dx < e^{-(q/2)^2/4b} \cdot \left(1 + \frac{4b}{2q/2}\right), \quad (32)$$

and it follows from the combination of (26), (27) and (32) that

$$\sum_{k=q/2}^{\infty} e^{-k^2/4b} < e^{-b\pi^2} \cdot \left(1 + \frac{1}{\pi}\right). \quad (33)$$

Substituting (33) into (30) and (31) we have

$$\left| \sum_{k>[c]+q/2} \rho_k e^{ikx} \right| + \left| \sum_{k<[c]-q/2} \rho_k e^{ikx} \right| < \frac{2e^{-b\pi^2}}{\sqrt{2\pi}} \cdot \left(1 + \frac{1}{\pi}\right) < e^{-b\pi^2} \cdot \frac{10}{9}, \quad (34)$$

and finally, substituting (25) and (34) into (29), we obtain

$$\left| \phi(x) - \sum_{k=[c]-q/2}^{[c]+q/2} \rho_k e^{ikx} \right| < e^{-b\pi^2} \cdot \left(4b + \frac{70}{9} + \frac{10}{9}\right) < e^{-b\pi^2} \cdot (4b + 9). \quad (35)$$

□

The following corollary describes a scheme for approximating e^{icx} using a series of q terms.

Corollary 2.9 Suppose that $m \geq 2$ is an integer and that the conditions of Theorem 2.8 are satisfied. Then, multiplying both sides of (28) by e^{bx^2} , we obtain

$$\begin{aligned} \left| e^{icx} - e^{bx^2} \cdot \sum_{k=[c]-q/2}^{[c]+q/2} \rho_k e^{ikx} \right| &< e^{bx^2} \cdot e^{-b\pi^2} \cdot (4b+9) \\ &< e^{b\pi^2/m^2} \cdot e^{-b\pi^2} \cdot (4b+9) \end{aligned} \quad (36)$$

for any $x \in [-\frac{\pi}{m}, \frac{\pi}{m}]$.

Finally, Theorem 2.10 makes use of a simple linear scaling to generalize the inequality (36) from $[-\frac{\pi}{m}, \frac{\pi}{m}]$ to any interval $[-d, d]$.

Theorem 2.10 Let $b > \frac{1}{2}$, $c, d > 0$ be real numbers, and let $m \geq 2, q \geq 4b\pi$ be integers. Then, for any $x \in [-d, d]$,

$$\left| e^{icx} - e^{b(x\pi/md)^2} \cdot \sum_{k=[cmd/\pi]-q/2}^{[cmd/\pi]+q/2} \rho_k e^{ikx\pi/md} \right| < e^{-b\pi^2(1-1/m^2)} \cdot (4b+9) \quad (37)$$

where $\{\rho_k\}$ are defined by (17).

Remark 2.3 The estimated error bounds obtained in the above theorems are rather pessimistic. Numerical estimates for the actual errors can be found in Appendix A to this paper.

3 Exact Statement of the Problem

In the remainder of this paper we will operate under the following assumptions:

1. $\omega = \{\omega_0, \dots, \omega_N\}$ and $x = \{x_0, \dots, x_N\}$ are finite sequences of real numbers.
2. $\omega_k \in [-N/2, N/2]$ for $k = 0, \dots, N$.
3. $x_j \in [-\pi, \pi]$ for $j = 0, \dots, N$.
4. $\alpha = \{\alpha_0, \dots, \alpha_N\}$, $f = \{f_{-N/2}, \dots, f_{N/2}\}$, $\beta = \{\beta_{-N/2}, \dots, \beta_{N/2}\}$, $g = \{g_0, \dots, g_N\}$, $\gamma = \{\gamma_0, \dots, \gamma_N\}$ and $h = \{h_0, \dots, h_N\}$ are finite sequences of complex numbers.

We will consider the problems of applying and inverting the matrix of the Fourier kernel and its transpose, i.e. we are interested in the transformations $F, G : \mathbb{C}^{N+1} \rightarrow \mathbb{C}^{N+1}$ and their inverses defined by the formulae

$$f_j = F(\alpha)_j = \sum_{k=0}^N \alpha_k \cdot e^{i\omega_k \cdot 2\pi j/N} \quad (38)$$

$$g_j = G(\beta)_j = \sum_{k=-N/2}^{N/2} \beta_k \cdot e^{ikx_j} \quad (39)$$

Remark 3.1 If $x_j = -\omega_j \cdot 2\pi/N$, then $G = F^*$.

We will also consider the more general transformation $H : \mathbb{C}^{N+1} \rightarrow \mathbb{C}^{N+1}$ defined by the formula

$$h_j = H(\gamma)_j = \sum_{k=0}^N \gamma_k \cdot e^{i\omega_k x_j} \quad (40)$$

More formally, we consider the following problems

- **Problem 1**
Given α , find $f = F(\alpha)$.
- **Problem 2**
Given β , find $g = G(\beta)$.
- **Problem 3**
Given γ , find $h = H(\gamma)$.
- **Problem 4**
Given f , find $\alpha = F^{-1}(f)$.
- **Problem 5**
Given g , find $\beta = G^{-1}(g)$.

Remark 3.2 We wish to perform all calculations with a fixed relative accuracy $\varepsilon > 0$. In the case of Problem 1, for instance, we are looking for a vector $\tilde{f} = \{\tilde{f}_{-N/2}, \dots, \tilde{f}_{N/2}\}$ such that

$$\frac{\|\tilde{f} - f\|}{\|f\|} \leq \varepsilon. \quad (41)$$

In this sense, all algorithms described in this paper are approximate ones.

4 Informal Descriptions of the Algorithms

4.1 Algorithms 1, 2 and 3 for Problems 1, 2 and 3

Observation 4.1 *According to Theorem 2.10, any function of the form e^{icx} can be accurately represented on any finite interval on the real line using a small number of terms of the form $e^{bx^2} \cdot e^{ikx}$, and this number of terms, q , is independent of the value of c .*

The FFT algorithm applies the matrix of the Fourier kernel to arbitrary complex vectors in $O(N \log N)$ operations when $\{\omega_k\}$ are integers and $\{x_j\}$ are equally spaced in $[-\pi, \pi]$. For the efficient application of the transformations described by (38), (39) and (40), we relate these more general cases to the equispaced case of the FFT. Observation 4.1 is used in two ways to achieve this:

- To approximate each $e^{i\omega_k x}$ in terms of a q -term Fourier series.
- To approximate the value of a Fourier series at each x_j in terms of the values of this series at the nearest q equispaced nodes.

This interpolation between equispaced and non-equispaced sets of points can thus be performed in $O(Nq)$ operations. The overall complexity of each of Algorithms 1, 2 and 3 will therefore be $O(N \log N + Nq)$ operations.

4.2 Algorithms 4 and 5 for Problems 4 and 5

Here we are interested in applying the complex matrices A^{-1} and $(A^*)^{-1}$ to arbitrary complex vectors where the elements of A are defined by

$$A_{jk} = e^{ikx_j} \quad (42)$$

for $j = 0, \dots, N$ and $k = -N/2, \dots, N/2$.

We make use of the following two simple observations.

Observation 4.2 *The matrix AA^* is Toeplitz, and furthermore, its $2N + 1$ distinct elements can be computed using Algorithm 1.*

Proof. It is obvious from (42) that

$$(AA^*)_{jk} = \sum_{l=0}^N e^{ijx_l} \cdot e^{-ikx_l} = \sum_{l=0}^N e^{i(j-k)x_l} \quad (43)$$

which is a function only of $(j - k)$, and is of the same form as (38), the description for Problem 1. $\square$

Observation 4.3 *From elementary matrix identities we see that*

$$A^{-1} \equiv A^*(AA^*)^{-1} \quad (44)$$

$$(A^*)^{-1} \equiv (AA^*)^{-1}A. \quad (45)$$

The Toeplitz matrix AA^* can be applied to arbitrary vectors in $O(N \log N)$ operations using an FFT-based discrete convolution. $(AA^*)^{-1}$ can therefore be applied to a vector in $O(\kappa(A) \cdot N \log N)$ operations using the conjugate gradient method where $\kappa(A)$ is the condition number of A .

Observation 4.4 *As application of A^* and A are $O(N \log N + Nq)$ procedures using Algorithms 1 and 2, A^{-1} and $(A^*)^{-1}$ can be applied in $O(\kappa(A) \cdot N \log N + Nq)$ operations due to Observation 4.3.*

Remark 4.5 It is well known that the condition number of A is 1 if the points $\{x_j\}$ are equally spaced. While the condition number deteriorates as the distribution of points becomes more non-uniform, in many cases of practical interest the points will be fairly uniformly spaced, so the condition number will not be very large.

4.3 Algorithm 6 for a Variant of Problem 5

The following lemma describes a way of computing the coefficients of an $\frac{N}{2}$ -term Fourier series which is tabulated at N points.

Lemma 4.1 Suppose that the $N + 1$ function values $g_0, \dots, g_N$ are given by the formula

$$g_j = \sum_{k=-N/4}^{N/4} \beta_k \cdot e^{ikx_j}, \quad (46)$$

and the vector $\xi = \{\xi_0, \dots, \xi_N\}$ is the unique solution of the linear system described by the equations

$$\sum_{j=0}^N \xi_j \cdot e^{ikx_j} = \begin{cases} 1 & \text{if } k = 0 \\ 0 & \text{otherwise} \end{cases} \quad (47)$$

for $k = -N/2, \dots, N/2$. Then, for $k = -N/4, \dots, N/4$,

$$\beta_k = \sum_{j=0}^N \xi_j \cdot g_j \cdot e^{-ikx_j}. \quad (48)$$

Proof. Substituting for g_j from (46) we have for $k = -N/4, \dots, N/4$

$$\sum_{j=0}^N \xi_j \cdot g_j \cdot e^{-ikx_j} = \sum_{j=0}^N \xi_j \cdot e^{-ikx_j} \cdot \sum_{l=-N/4}^{N/4} \beta_l \cdot e^{ilx_j} \quad (49)$$

$$= \sum_{l=-N/4}^{N/4} \beta_l \cdot \sum_{j=0}^N \xi_j \cdot e^{i(l-k)x_j} \quad (50)$$

$$= \beta_k. \quad (51)$$

□

Remark 4.6 According to (47) and Lemma 2.2,

$$\sum_{j=0}^N \xi_j \cdot e^{ikx_j} = \frac{1}{2\pi} \int_{-\pi}^{\pi} e^{ikx} dx \quad (52)$$

for $k = -N/2, \dots, N/2$. Thus, the set of numbers $\{\xi_j\}$ can be considered as weights which integrate exactly all N -th order trigonometric polynomials at the nodes $\{x_j\}$.

Observation 4.7 Rewriting (47) in matrix notation we see that

$$A^T \xi = (0, \dots, 0, 1, 0, \dots, 0)^T = A^* \xi \quad (53)$$

where $A_{jk} = e^{ikx_j}$, so the vector ξ is real, and can be computed using the algorithm for Problem 4 as introduced in Section 4.2.

Observation 4.8 Equation (48) is of the same form as (38). Thus, provided the vector ξ is known, the vector β can be computed in $O(N \log N + Nq)$ operations using the algorithm for Problem 1 as introduced in Section 4.1.

Remark 4.9 According to Observation 4.8, if a function described by an $N/2$ -term Fourier series is tabulated at N arbitrary nodes, the $N/2$ coefficients can be obtained in $O(N \log N + Nq)$ operations.

Also, due to Observations 4.7 and 4.4, the precomputation of the numbers $\{\xi_j\}$ needed for this algorithm requires $O(\kappa(A) \cdot N \log N + Nq)$ operations.

5 Notation

In this section we introduce the notation to be used in the detailed descriptions of the algorithms in the next section.

For an integer $m \geq 2$ and a real number $b > 0$, we will define a real number $\varepsilon > 0$ by

$$\varepsilon = e^{-b\pi^2(1-1/m^2)} \cdot (4b + 9), \quad (54)$$

and we will denote by q the smallest even natural number such that

$$q \geq 4b\pi. \quad (55)$$

For an integer m and a set of real numbers $\{\omega_k\}$ we will denote by μ_k the nearest integer to $m\omega_k$ for $k = 0, \dots, N$, and by $\{P_{jk}\}$ a set of real numbers defined by the formula

$$P_{jk} = \frac{1}{2\sqrt{b\pi}} \cdot e^{-(m\omega_k - (\mu_k + j))^2 / 4b} \quad (56)$$

for $k = 0, \dots, N$ and $j = -q/2, \dots, q/2$.

Observation 5.1 Setting $d = \pi$ in Theorem 2.10 we see that

$$\left| e^{i\omega_k x} - e^{b(x/m)^2} \cdot \sum_{j=-q/2}^{q/2} P_{jk} \cdot e^{i(\mu_k + j)x/m} \right| < \varepsilon \quad (57)$$

for any $k = 0, \dots, N$ and any $x \in [-\pi, \pi]$, where ε is defined by (54).

For a given set of complex numbers $\{\alpha_k\}$, we will denote by $\{\tau_j\}$ the unique set of complex coefficients such that

$$\sum_{k=1}^N \alpha_k \cdot \sum_{j=-q/2}^{q/2} P_{jk} \cdot e^{i(\mu_k + j)x/m} = \sum_{j=-mN/2}^{mN/2-1} \tau_j e^{ijx/m}, \quad (58)$$

so that

$$\tau_l = \sum_{j,k,\mu_k+j=l} \alpha_k \cdot P_{jk}. \quad (59)$$

We will denote by $\{T_j\}$ a set of complex numbers defined by the formula

$$T_j = \sum_{k=-mN/2}^{mN/2-1} \tau_k \cdot e^{2\pi i k j / mN} \quad (60)$$

for $j = -mN/2, \dots, mN/2 - 1$.

Further, we will denote by $\{\tilde{f}_j\}$ another set of complex numbers defined by the formula

$$\tilde{f}_j = e^{b(2\pi j / mN)^2} \cdot T_j \quad (61)$$

for $j = -N/2, \dots, N/2$.

Observation 5.2 Combining (57) – (61) with the triangle inequality, we see that

$$|f_j - \tilde{f}_j| < \varepsilon \cdot \sum_{k=0}^N |\alpha_k| \quad (62)$$

for $j = -N/2, \dots, N/2$, where $\{f_j = F(\alpha)_j\}$ are defined by (38).

For an integer m and a set of real numbers $\{x_j\}$ we will denote by ν_j the nearest integer to $x_j mN/2\pi$ for $j = 0, \dots, N$, and by $\{Q_{jk}\}$ a set of real numbers defined by the formula

$$Q_{jk} = \frac{1}{2\sqrt{b\pi}} \cdot e^{-(x_j mN/2\pi - (\nu_j + k))^2 / 4b} \quad (63)$$

for $j = 0, \dots, N$ and $k = -q/2, \dots, q/2$.

Observation 5.3 Setting $d = N/2$ in Theorem 2.10 we see that

$$\left| e^{ikx_j} - e^{b(2\pi k / mN)^2} \cdot \sum_{l=-q/2}^{q/2} Q_{jk} \cdot e^{i(\nu_j + l)2\pi k / mN} \right| < \varepsilon \quad (64)$$

for any $j = 0, \dots, N$ and any $k \in [-N/2, N/2]$, where ε is defined by (54).

For a given set of complex numbers $\{\beta_k\}$, we will denote by $\{u_k\}$ a set of complex numbers defined by the formula

$$u_k = \beta_k \cdot e^{b(2\pi k / mN)^2} \quad (65)$$

for $k = -N/2, \dots, N/2$, and by $\{U_l\}$ a set of complex numbers defined by the formula

$$U_l = \sum_{k=-N/2}^{N/2} u_k \cdot e^{2\pi i k l / mN} \quad (66)$$

for $l = -mN/2, \dots, mN/2 - 1$.

Further, we will denote by $\{\tilde{g}_j\}$ another set of complex numbers defined by the formula

$$\tilde{g}_j = \sum_{l=-q/2}^{q/2} Q_{jl} \cdot U_{\nu_j+l} \quad (67)$$

for $j = 0, \dots, N$.

Observation 5.4 Combining (64) - (67) with the triangle inequality, we see that

$$|g_j - \tilde{g}_j| < \varepsilon \cdot \sum_{k=0}^N |\beta_k| \quad (68)$$

for $j = 0, \dots, N$, where $\{g_j = G(\beta)_j\}$ are defined by (39).

For a set of real numbers $\{x_j\}$ we will denote by η_j the nearest integer to $x_j N/2\pi$ for $j = 0, \dots, N$, and by $\{R_{jk}\}$ a set of real numbers defined by the formula

$$R_{jk} = \frac{1}{2\sqrt{b\pi}} \cdot e^{-(x_j N/2\pi - (\eta_j + k))^2 / 4b} \quad (69)$$

for $j = 0, \dots, N$ and $k = -q/2, \dots, q/2$.

Observation 5.5 Setting $d = N/2$ in Theorem 2.10 we see that

$$\left| e^{ikx_j/m} - e^{b(2\pi k/mN)^2} \cdot \sum_{l=-q/2}^{q/2} R_{jk} \cdot e^{i(\eta_j+l)2\pi k/mN} \right| < \varepsilon \quad (70)$$

for any $j = 0, \dots, N$ and any $k \in [-N/2, N/2]$ where ε is defined by (54).

For a given set of complex numbers $\{\gamma_k\}$, we will denote by $\{v_j\}$ the unique set of complex coefficients such that

$$\sum_{k=0}^N \gamma_k \cdot \sum_{j=-q/2}^{q/2} P_{jk} \cdot e^{i(\mu_k+j)x/m} = \sum_{j=-mN/2}^{mN/2} v_j \cdot e^{ijx/m}, \quad (71)$$

so that

$$v_l = \sum_{j,k, \eta_k+j=l} \gamma_k \cdot P_{jk}. \quad (72)$$

We denote by $\{V_l\}$ a set of complex numbers defined by the formula

$$V_l = \sum_{k=-mN/2}^{mN/2} v_k \cdot e^{b(2\pi k/m^2 N)^2} \cdot e^{2\pi ikl/m^2 N} \quad (73)$$

for $l = -m^2 N/2, \dots, m^2 N/2 - 1$.

Further, we will denote by $\{\tilde{h}_j\}$ another set of complex numbers defined by the formula

$$\tilde{h}_j = e^{b(x_j/m)^2} \cdot \sum_{l=-q/2}^{q/2} R_{jl} \cdot V_{\eta_j+l} \quad (74)$$

for $j = 0, \dots, N$.

Observation 5.6 *Combining (57) and (70) – (74) with the triangle inequality, we see that*

$$|h_j - \tilde{h}_j| < \delta \cdot \sum_{k=0}^N |\gamma_k| \quad (75)$$

for $j = 0, \dots, N$, where $\{h_j = H(\gamma)_j\}$ are defined by (40), and

$$\delta = 2e^{-b\pi^2(1-2/m^2)} \cdot (4b + 9). \quad (76)$$

For a set of real numbers $\{x_j\}$, A will denote a complex matrix whose elements are given by

$$A_{jk} = e^{ikx_j} \quad (77)$$

for $k = -N/2, \dots, N/2$ and $j = 0, \dots, N$, and $a = \{a_{-N}, \dots, a_N\}$ will denote a set of complex numbers defined by the formula

$$a_k = \sum_{j=0}^N e^{ikx_j}. \quad (78)$$

Finally, $\xi = \{\xi_0, \dots, \xi_N\}$ will denote a real vector defined by

$$\xi = (A^*)^{-1}(0, \dots, 0, 1, 0, \dots, 0)^T. \quad (79)$$

Remark 5.7 It is clear from Observation 4.2 that

$$(AA^*)_{jk} = a_{j-k}. \quad (80)$$

6 Detailed Descriptions of the Algorithms

This section contains step by step descriptions and operation counts for the six algorithms of this paper. In the tables below we will make use of the facts that $q \sim \log(\frac{1}{\epsilon})$ and $m^2 \ll N$.

Algorithm 1.

Step	Complexity	Description
Init	$O(Nq)$	Comment [Input parameters are the vector $\{\omega_0, \dots, \omega_N\}$ and a real number $\epsilon > 0$.] Choose b and q in terms of ϵ do $k = 0, N$ Determine μ_k, the nearest integer to $m\omega_k$ do $j = -q/2, q/2$ Compute P_{jk} according to (56) end do end do do $j = -N/2, N/2$ Compute $e^{b(2\pi j/mN)^2}$ end do
1	$O(Nq)$	Comment [Input parameter is the vector $\{\alpha_0, \dots, \alpha_N\}$.] Comment [Compute Fourier coefficients τ_j.] do $k = 0, N$ do $j = -q/2, q/2$ $\tau_{\mu_k+j} \leftarrow \tau_{\mu_k+j} + P_{jk} \cdot \alpha_k$ end do end do
2	$O(mN \log N)$	Comment [Evaluate this Fourier Series at equispaced points in $[-m\pi, m\pi]$ using inverse FFT.] Compute $T_j = \sum_{k=-mN/2}^{mN/2-1} \tau_k \cdot e^{2\pi i k j / mN}$ for $j = -mN/2, \dots, mN/2 - 1$.
3	$O(N)$	Comment [Scale the values at those points which lie in $[-\pi, \pi]$.] do $j = -N/2, N/2$ $\tilde{f}_j = T_j \cdot e^{b(2\pi j/mN)^2}$ end do
Total	$O(N \cdot \log(\frac{1}{\epsilon}) + mN \cdot \log N)$	

Algorithm 2.

Step	Complexity	Description
Init	$O(Nq)$	Comment [Input parameters are the vector $\{x_0, \dots, x_N\}$ and a real number $\varepsilon > 0$.] Choose b and q in terms of ε do $j = 0, N$ Determine ν_j, the nearest integer to $x_j mN/2\pi$ do $k = -q/2, q/2$ Compute Q_{jk} according to (63) end do end do do $k = -N/2, N/2$ Compute $e^{b(2\pi k/mN)^2}$ end do
1	$O(N)$	Comment [Input parameter is the vector $\{\beta_{-N/2}, \dots, \beta_{N/2}\}$.] Comment [Compute new, scaled Fourier coefficients.] do $k = -N/2, N/2$ $u_k = \beta_k \cdot e^{b(2\pi k/mN)^2}$ end do
2	$O(mN \log N)$	Comment [Evaluate this Fourier Series at equispaced points in $[-\pi, \pi]$ using inverse FFT.] Compute $U_j = \sum_{k=-N/2}^{N/2} u_k \cdot e^{2\pi i k j / mN}$ for $j = -mN/2, \dots, mN/2 - 1$.
3	$O(Nq)$	Comment [Compute approximate values at desired points in terms of the values at equispaced points.] do $j = 0, N$ do $k = -q/2, q/2$ $\tilde{g}_j \leftarrow \tilde{g}_j + Q_{jk} \cdot U_{\nu_j+k}$ end do end do
Total	$O(mN \cdot \log N + N \cdot \log(\frac{1}{\varepsilon}))$	

Algorithm 3.

Step Complexity

Description

Init $O(Nq)$

Comment [Input parameters are the vectors $\{\omega_0, \dots, \omega_N\}$ and $\{x_0, \dots, x_N\}$ and a real number $\varepsilon > 0$.]

Choose b and q in terms of ε

do $k = 0, N$

Determine μ_k , the nearest integer to $m\omega_k$

do $j = -q/2, q/2$

Compute P_{jk} according to (56)

end do

end do

do $k = -mN/2, mN/2$

Compute $e^{b(2\pi k/m^2 N)^2}$

end do

do $j = 0, N$

Determine η_j , the nearest integer to $x_j N/2\pi$

do $k = -q/2, q/2$

Compute R_{jk} according to (69)

end do

Compute $e^{b(x_j/m)^2}$

end do

1 $O(Nq)$

Comment [Input parameter is the vector $\{\gamma_0, \dots, \gamma_N\}$.]

Comment [Compute Fourier coefficients v_j .]

do $k = 0, N$

do $j = -q/2, q/2$

$v_{\mu_k+j} \leftarrow v_{\mu_k+j} + P_{jk} \cdot \gamma_k$

end do

end do

2 $O(mN)$

Comment [Scale the coefficients.]

do $k = -mN/2, mN/2$

$v_k \leftarrow v_k \cdot e^{b(2\pi k/m^2 N)^2}$

end do

3 $O(m^2 N \log N)$

Comment [Evaluate this Fourier Series at equispaced points in $[-m\pi, m\pi]$ using inverse FFT.]

Compute $V_j = \sum_{k=-mN/2}^{mN/2} v_k \cdot e^{2\pi i k j / m^2 N}$ for $j = -m^2 N/2, \dots, m^2 N/2 - 1$.

4	$O(Nq)$	Comment [Compute approximate values at desired points in terms of the values at equispaced points.] do $j = 0, N$ do $k = -q/2, q/2$ $\tilde{h}_j \leftarrow \tilde{h}_j + R_{jk} \cdot V_{\eta_j+k}$ end do end do
5	$O(N)$	Comment [Scale the values.] do $j = 0, N$ $\tilde{h}_j \leftarrow \tilde{h}_j \cdot e^{b(x_j/m)^2}$ end do
Total $O(m^2 N \cdot \log N + N \cdot \log(\frac{1}{\epsilon}))$		

Algorithm 4.

Step	Complexity	Description
Init	$O(N \log N + Nq)$	Initialization for Algorithm 1. Initialization for Algorithm 2. Compute elements $\{a_k\}$ of Toeplitz matrix $(AA^*)^{-1}$ as defined by (78) using Algorithm 1.
1	$O(N \log N + Nq)$	Compute Af using Algorithm 2.
2	$O(\kappa(A) \cdot N \log N)$	Compute $\tilde{\alpha} = (AA^*)^{-1}(Af)$ using Conjugate Gradient algorithm.
Total $O(\kappa(A) \cdot N \cdot \log N + N \cdot \log(\frac{1}{\epsilon}))$		

Algorithm 5.

Step	Complexity	Description
Init	$O(N \log N + Nq)$	Initialization for Algorithm 1. Compute elements $\{a_k\}$ of Toeplitz matrix $(AA^*)^{-1}$ as defined by (78) using Algorithm 1.
1	$O(\kappa(A) \cdot N \log N)$	Compute $(AA^*)^{-1}g$ using Conjugate Gradient algorithm.
2	$O(N \log N + Nq)$	Compute $\tilde{\beta} = A^*((AA^*)^{-1}g)$ using Algorithm 1.
Total $O(\kappa(A) \cdot N \cdot \log N + N \cdot \log(\frac{1}{\epsilon}))$		

Algorithm 6.

Step	Complexity	Description
Init	$O(\kappa(A) \cdot N \log N + Nq)$	Initialization for Algorithm 5. Compute ξ as defined by (79) using Algorithm 5.
1	$O(N)$	Compute $\hat{g}_j = \xi_j g_j$ for $j = 0, \dots, N$.
2	$O(N \log N + Nq)$	Compute $\tilde{\beta} = A^* \hat{g}$ using Algorithm 1.
Total	$O(N \cdot \log N + N \cdot \log(\frac{1}{\epsilon}))$	

The storage requirements of an algorithm are also an important characteristic. From the above descriptions for the initialization steps the asymptotic storage requirements for each algorithm are of the form

$$\lambda \cdot N \cdot q \quad (81)$$

where the coefficient λ is software- and hardware-dependent.

7 Numerical Results

We have written FORTRAN implementations of the six algorithms of this paper and have tested these programs on the Sun SPARCstation 1 for a variety of input data. Six such examples are presented in this section, one for each algorithm.

Tables 1–6 contain accuracies and CPU timings for the algorithms with computations performed in both single and double precision arithmetic, and the input size N varying between 64 and 4096. In addition, each table contains the CPU times required to solve the same set of problems via a direct calculation, and Tables 1–3 include timings for an FFT of the same size. Tables 1–3 also contain the accuracies of the direct single precision calculations.

Two measures of accuracy were chosen for each example. In Examples 1, 2 and 3, these are defined by the formulae

$$E_\infty = \max_{0 \leq j \leq N} |\tilde{f}_j - f_j| / \sum_{j=0}^N |\alpha_j|, \quad (82)$$

and

$$E_2 = \sqrt{\sum_{j=0}^N |\tilde{f}_j - f_j|^2 / \sum_{j=0}^N |f_j|^2}, \quad (83)$$

where α is the input vector, f is the result of a direct computation in double precision arithmetic and $\tilde{f}$ is the result of the computation being considered.

In Examples 4, 5 and 6, they are given by

$$E_{\infty} = \max_{0 \leq j \leq N} |\tilde{\alpha}_j - \alpha_j| / \max_{0 \leq j \leq N} |\alpha_j|, \quad (84)$$

and

$$E_2 = \sqrt{\sum_{j=0}^N |\tilde{\alpha}_j - \alpha_j|^2 / \sum_{j=0}^N |\alpha_j|^2}, \quad (85)$$

where α is the input for a direct double precision computation, and $\tilde{\alpha}$ is the result of applying the algorithm to the result of this computation.

Remark 7.1 The formulae (82) – (85) measure fairly accurately the errors of all single precision computations. However, they can only provide rough estimates of the errors produced by the double precision versions of the algorithms.

Several technical details of our implementations appear to be worth mentioning here:

1. Each implementation consists of two main subroutines: the first is an initialization stage in which the matrix operators of the algorithm are precomputed and stored, and the second is an evaluation stage in which these operators are applied. Successive application of the linear transformations to multiple vectors requires the initialization to be performed only once.
2. For the single precision versions of each algorithm our choice of parameters was $m = 2$, $b = 0.5993$ and $q = 10$. For double precision we chose $m = 2$, $b = 1.5629$ and $q = 28$.
3. The algorithms as described in this paper all require an FFT of size proportional to N , and thus will perform efficiently whenever the FFT does. This restriction on the FFT size can be removed by means of a simple modification, thereby ensuring that the algorithms perform efficiently for any choice of N . In our implementations, these changes were made.
4. In the direct methods for Problems 1 and 2, we used the fact that $e^{ikx_j} = (e^{ix_j})^k$ to reduce the number of exponential computations from N^2 to N .
5. N^2 exponentials are required in the direct method for Problem 3, and for larger N , the available memory on the machine is insufficient for the precomputation and storage of these numbers. The direct implementation we used for this problem computes each exponential when it is needed.
6. Standard LINPACK Gaussian Elimination subroutines were used as the direct methods for comparing timings in Examples 4, 5 and 6. Estimated timings are presented for larger N , where this computation became impractical.

Example 1.

Here we consider the transformation $F : \mathbb{C}^{N+1} \rightarrow \mathbb{C}^{N+1}$ of Problem 1 as defined by the formula

$$F(\alpha)_j = \sum_{k=0}^N \alpha_k \cdot e^{i\omega_k \cdot 2\pi j/N} \quad (86)$$

for $j = -N/2, \dots, N/2$. In this example, $\{\omega_0, \dots, \omega_N\}$ were randomly distributed on the interval $[-N/2, N/2]$ and $\{\alpha_0, \dots, \alpha_N\}$ were generated randomly on the unit square in the complex plane defined by the formulae

$$0 \leq \text{Re}(z) \leq 1, 0 \leq \text{Im}(z) \leq 1. \quad (87)$$

The results of applying Algorithm 1 to this problem are presented in Tables 1(a) and 1(b).

Table 1(a)
Example 1, Single Precision Computations

N	Errors				Timings (sec.)			
	Algorithm		Direct		Algorithm		Direct	FFT
	E_∞	E_2	E_∞	E_2	Init.	Eval.		
64	0.175 E-05	0.190 E-05	0.337 E-06	0.811 E-06	0.011	0.003	0.01	0.0008
128	0.973 E-06	0.203 E-05	0.449 E-06	0.162 E-05	0.022	0.006	0.03	0.0015
256	0.113 E-05	0.348 E-05	0.100 E-05	0.374 E-05	0.044	0.014	0.13	0.0034
512	0.138 E-05	0.602 E-05	0.168 E-05	0.769 E-05	0.088	0.030	0.49	0.0078
1024	0.280 E-05	0.128 E-04	0.193 E-05	0.133 E-04	0.174	0.067	1.90	0.0174
2048	0.267 E-05	0.243 E-04	0.424 E-05	0.268 E-04	0.348	0.141	7.47	0.0352
4096	0.551 E-05	0.452 E-04	0.532 E-05	0.565 E-04	0.701	0.323	30.24	0.0846

Table 1(b)
Example 1, Double Precision Computations

N	Errors		Timings (sec.)			
	E_∞	E_2	Alg. Init.	Alg. Eval.	Direct	FFT
64	0.495 E-14	0.634 E-14	0.036	0.008	0.02	0.001
128	0.689 E-14	0.104 E-13	0.075	0.017	0.06	0.002
256	0.717 E-14	0.119 E-13	0.148	0.038	0.22	0.005
512	0.306 E-14	0.164 E-13	0.297	0.079	0.84	0.012
1024	0.460 E-14	0.310 E-13	0.600	0.163	3.23	0.026
2048	0.694 E-14	0.625 E-13	1.204	0.342	12.55	0.059
4096	0.129 E-13	0.126 E-12	2.418	0.747	49.69	0.132

Example 2.

Here we consider the transformation $G : \mathbb{C}^{N+1} \rightarrow \mathbb{C}^{N+1}$ of Problem 2 as defined by the formula

$$G(\beta)_j = \sum_{k=-N/2}^{N/2} \beta_k \cdot e^{ikx_j} \quad (88)$$

for $j = 0, \dots, N$. In this example, $\{x_0, \dots, x_N\}$ were randomly distributed on the interval $[-\pi, \pi]$ and $\{\beta_{-N/2}, \dots, \beta_{N/2}\}$ were generated randomly on the unit square in the complex plane defined by the formulae

$$0 \leq \text{Re}(z) \leq 1, \quad 0 \leq \text{Im}(z) \leq 1. \quad (89)$$

The results of applying Algorithm 2 to this problem are presented in Tables 2(a) and 2(b).

Table 2(a)
Example 2, Single Precision Computations

N	Errors				Timings (sec.)			
	Algorithm		Direct		Algorithm		Direct	FFT
	E_∞	E_2	E_∞	E_2	Init.	Eval.		
64	0.121 E-05	0.173 E-05	0.237 E-06	0.506 E-06	0.011	0.003	0.01	0.0008
128	0.595 E-06	0.337 E-05	0.362 E-06	0.985 E-06	0.022	0.007	0.03	0.0015
256	0.133 E-05	0.620 E-05	0.233 E-06	0.158 E-05	0.043	0.015	0.12	0.0034
512	0.953 E-06	0.890 E-05	0.828 E-06	0.334 E-05	0.086	0.033	0.46	0.0078
1024	0.164 E-05	0.996 E-05	0.431 E-05	0.523 E-05	0.174	0.067	1.80	0.0174
2048	0.223 E-05	0.156 E-04	0.687 E-05	0.824 E-05	0.345	0.149	7.11	0.0352
4096	0.470 E-05	0.338 E-04	0.759 E-05	0.134 E-04	0.687	0.332	29.03	0.0846

Table 2(b)
Example 2, Double Precision Computations

N	Errors		Timings (sec.)			
	E_∞	E_2	Alg. Init.	Alg. Eval.	Direct	FFT
64	0.249 E-14	0.814 E-14	0.038	0.009	0.01	0.001
128	0.501 E-14	0.746 E-14	0.075	0.020	0.05	0.002
256	0.418 E-14	0.623 E-14	0.150	0.043	0.18	0.005
512	0.356 E-14	0.831 E-14	0.297	0.088	0.69	0.012
1024	0.793 E-14	0.192 E-13	0.598	0.182	2.72	0.026
2048	0.138 E-13	0.405 E-13	1.192	0.379	10.86	0.059
4096	0.278 E-13	0.904 E-13	2.417	0.816	43.36	0.132

Example 3.

Here we consider the transformation $H : \mathbb{C}^{N+1} \rightarrow \mathbb{C}^{N+1}$ of Problem 3 as defined by the formula

$$H(\gamma)_j = \sum_{k=0}^N \gamma_k \cdot e^{i\omega_k x_j} \quad (90)$$

for $j = 0, \dots, N$. In this example, $\{\omega_0, \dots, \omega_N\}$ were randomly distributed on the interval $[-N/2, N/2]$, $\{x_0, \dots, x_N\}$ were randomly distributed on the interval $[-\pi, \pi]$ and $\{\gamma_0, \dots, \gamma_N\}$ were generated randomly on the unit square in the complex plane defined by the formulae

$$0 \leq \text{Re}(z) \leq 1, 0 \leq \text{Im}(z) \leq 1. \quad (91)$$

The results of applying Algorithm 3 to this problem are presented in Tables 3(a) and 3(b).

Table 3(a)
Example 3, Single Precision Computations

N	Errors				Timings (sec.)			
	Algorithm		Direct		Algorithm		Direct	FFT
	E_∞	E_2	E_∞	E_2	Init.	Eval.		
64	0.111 E-05	0.291 E-05	0.570 E-06	0.707 E-06	0.022	0.007	0.24	0.0008
128	0.204 E-05	0.405 E-05	0.593 E-06	0.145 E-05	0.044	0.015	0.63	0.0015
256	0.127 E-05	0.450 E-05	0.900 E-06	0.296 E-05	0.087	0.032	2.60	0.0034
512	0.181 E-05	0.683 E-05	0.111 E-05	0.490 E-05	0.177	0.067	10.63	0.0078
1024	0.304 E-05	0.176 E-04	0.221 E-05	0.126 E-04	0.352	0.145	43.60	0.0174
2048	0.482 E-05	0.356 E-04	0.260 E-05	0.264 E-04	0.708	0.327	176.75	0.0352
4096	0.734 E-05	0.751 E-04	0.389 E-05	0.530 E-04	1.404	0.702	714.40	0.0846

Table 3(b)
Example 3, Double Precision Computations

N	Errors		Timings (sec.)			
	E_∞	E_2	Alg. Init.	Alg. Eval.	Direct	FFT
64	0.222 E-13	0.320 E-13	0.078	0.019	0.21	0.001
128	0.247 E-13	0.370 E-13	0.153	0.041	0.85	0.002
256	0.249 E-13	0.334 E-13	0.305	0.085	3.43	0.005
512	0.145 E-13	0.232 E-13	0.605	0.180	13.62	0.012
1024	0.237 E-13	0.416 E-13	1.210	0.363	54.71	0.026
2048	0.194 E-13	0.795 E-13	2.403	0.774	219.38	0.059
4096	0.411 E-13	0.120 E-12	4.827	1.588	889.04	0.132

Example 4.

Here we consider Problem 4 of Section 3. In this example, the numbers $\{\omega_k\}$ were defined by the formula

$$\omega_k = -\frac{N}{2} + (k + 0.5 + \delta_k) \cdot \frac{N}{N+1} \quad (92)$$

for $k = 0, \dots, N$, where δ_k were randomly distributed on the interval $[-0.1, 0.1]$. In addition, the numbers $\{\alpha_0, \dots, \alpha_N\}$ were generated randomly on the unit square in the complex plane defined by the formulae

$$0 \leq \text{Re}(z) \leq 1, \quad 0 \leq \text{Im}(z) \leq 1, \quad (93)$$

and the numbers $\{f_{-N/2}, \dots, f_{N/2}\}$ were computed directly in double precision arithmetic according to the formula

$$f_j = \sum_{k=0}^N \alpha_k \cdot e^{i\omega_k \cdot 2\pi j/N}. \quad (94)$$

The vector f was then used as input for Algorithm 4. Results of this experiment are presented in Tables 4(a) and 4(b).

Table 4(a)
Example 4, Single Precision Computations

N	Errors		Timings (sec.)		
	E_∞	E_2	Alg. Init.	Alg. Eval.	Direct
64	0.492 E-05	0.339 E-05	0.02	0.05	0.36
128	0.154 E-04	0.568 E-05	0.04	0.11	2.78
256	0.424 E-04	0.128 E-04	0.08	0.20	23.0
512	0.809 E-04	0.252 E-04	0.18	0.45	184
1024	0.203 E-03	0.474 E-04	0.36	0.82	1472 (est.)
2048	0.428 E-03	0.979 E-04	0.78	1.91	11776 (est.)
4096	0.106 E-02	0.195 E-03	1.66	4.64	94208 (est.)

Table 4(b)
Example 4, Double Precision Computations

N	Errors		Timings (sec.)		
	E_∞	E_2	Alg. Init.	Alg. Eval.	Direct
64	0.143 E-13	0.109 E-13	0.07	0.17	0.37
128	0.208 E-13	0.149 E-13	0.11	0.34	2.96
256	0.493 E-13	0.256 E-13	0.20	0.75	23.6
512	0.121 E-12	0.500 E-13	0.48	1.67	189
1024	0.279 E-12	0.926 E-13	0.94	3.55	1512 (est.)
2048	0.593 E-12	0.192 E-12	1.95	7.75	12096 (est.)
4096	0.138 E-11	0.375 E-12	4.02	18.28	96768 (est.)

Example 5.

Here we consider Problem 5 of Section 3. In this example, the numbers $\{x_j\}$ were defined by the formula

$$x_j = -\pi + 2\pi \cdot \frac{j + 0.5 + \delta_j}{N + 1} \quad (95)$$

for $j = 0, \dots, N$, where δ_j were randomly distributed on the interval $[-0.1, 0.1]$. In addition, the numbers $\{\beta_{-N/2}, \dots, \beta_{N/2}\}$ were generated randomly on the unit square in the complex plane defined by the formulae

$$0 \leq \operatorname{Re}(z) \leq 1, \quad 0 \leq \operatorname{Im}(z) \leq 1, \quad (96)$$

and the numbers $\{g_0, \dots, g_N\}$ were computed directly in double precision arithmetic according to the formula

$$g_j = \sum_{k=-N/2}^{N/2} \beta_k \cdot e^{ikx_j}. \quad (97)$$

The vector g was then used as input for Algorithm 5. Results of this experiment are presented in Tables 5(a) and 5(b).

Table 5(a)
Example 5, Single Precision Computations

N	Errors		Timings (sec.)		
	E_∞	E_2	Alg. Init.	Alg. Eval.	Direct
64	0.212 E-05	0.183 E-05	0.02	0.06	0.36
128	0.117 E-04	0.523 E-05	0.04	0.10	2.78
256	0.190 E-04	0.955 E-05	0.09	0.19	23.0
512	0.287 E-04	0.202 E-04	0.19	0.41	184
1024	0.560 E-04	0.376 E-04	0.37	0.83	1472 (est.)
2048	0.106 E-03	0.752 E-04	0.78	1.90	11776 (est.)
4096	0.225 E-03	0.148 E-03	1.66	4.67	94208 (est.)

Table 5(b)
Example 5, Double Precision Computations

N	Errors		Timings (sec.)		
	E_∞	E_2	Alg. Init.	Alg. Eval.	Direct
64	0.310 E-13	0.120 E-13	0.06	0.18	0.37
128	0.389 E-13	0.146 E-14	0.10	0.36	2.96
256	0.577 E-13	0.204 E-13	0.24	0.76	23.6
512	0.673 E-13	0.325 E-13	0.47	1.61	189
1024	0.118 E-12	0.817 E-13	0.97	3.54	1512 (est.)
2048	0.190 E-12	0.134 E-12	1.86	7.73	12096 (est.)
4096	0.429 E-12	0.288 E-12	3.93	18.21	96768 (est.)

Example 6.

Here we consider the variant of Problem 5 which was described in Section 4.3. In this example, the numbers $\{x_j\}$ were defined by the formula

$$x_j = -\pi + 2\pi \cdot \frac{j + 0.5 + \delta_j}{N + 1} \quad (98)$$

for $j = 0, \dots, N$, where δ_j were randomly distributed on the interval $[-0.1, 0.1]$. In addition, the numbers $\{\beta_{-N/4}, \dots, \beta_{N/4}\}$ were generated randomly on the unit square in the complex plane defined by the formulae

$$0 \leq \text{Re}(z) \leq 1, \quad 0 \leq \text{Im}(z) \leq 1, \quad (99)$$

and the numbers $\{g_0, \dots, g_N\}$ were computed directly in double precision arithmetic according to the formula

$$g_j = \sum_{k=-N/4}^{N/4} \beta_k \cdot e^{ikx_j}. \quad (100)$$

The vector g was then used as input for Algorithm 6. Results of this experiment are presented in Tables 6(a) and 6(b).

Table 6(a)
Example 6, Single Precision Computations

N	Errors		Timings (sec.)		
	E_∞	E_2	Alg. Init.	Alg. Eval.	Direct
64	0.195 E-05	0.174 E-05	0.09	0.004	0.36
128	0.412 E-05	0.286 E-05	0.17	0.007	2.78
256	0.105 E-04	0.913 E-05	0.34	0.014	23.0
512	0.195 E-04	0.147 E-04	0.70	0.031	184
1024	0.474 E-04	0.320 E-04	1.41	0.066	1472 (est.)
2048	0.836 E-04	0.631 E-04	3.05	0.147	11776 (est.)
4096	0.173 E-03	0.135 E-03	7.22	0.330	94208 (est.)

Table 6(b)
Example 6, Double Precision Computations

N	Errors		Timings (sec.)		
	E_∞	E_2	Alg. Init.	Alg. Eval.	Direct
64	0.878 E-14	0.762 E-14	0.32	0.008	0.37
128	0.102 E-13	0.981 E-14	0.62	0.018	2.96
256	0.213 E-13	0.180 E-13	1.25	0.039	23.6
512	0.444 E-13	0.376 E-13	2.64	0.083	189
1024	0.679 E-13	0.520 E-13	5.82	0.170	1512 (est.)
2048	0.151 E-12	0.115 E-12	12.37	0.359	12096 (est.)
4096	0.308 E-12	0.232 E-12	27.61	0.766	96768 (est.)

The following observations can be made from Tables 1-6 above, and are in agreement with results of our more extensive experiments.

1. The errors produced by Algorithms 1, 2 and 3 are comparable with those produced by the corresponding direct methods.
2. The timings for Algorithms 1 and 2 are similar, which is to be expected since Problem 2 is the adjoint of Problem 1. Algorithm 3 is about twice as costly, which is in agreement with the fact that it is a synthesis of Algorithms 1 and 2.
3. In single precision computations for this particular architecture, implementation and range of N , Algorithms 1 and 2 are roughly 4 times as costly as an FFT of the same size. For double precision the ratio is roughly 6. These ratios decrease as N increases.
4. The extrapolated break-even point of Algorithms 1 and 2 is at roughly $N = 32$ if the initialization time is ignored. If the initialization time is included, the break-even point is at $N = 256$. For Algorithm 3, the break-even points are at $N = 8$ without initialization and at $N = 32$ with initialization.
5. The timings for Algorithms 4 and 5 are similar as expected, since Problem 5 is the adjoint of Problem 4.
6. The break-even points of Algorithms 4 and 5 are at roughly $N = 32$. For Algorithm 6, the break-even points are at $N = 32$ if the initialization time is taken into account, and at $N = 16$ if it is ignored.
7. Algorithms 1 and 2 tend to be slightly more accurate than their inverses, Algorithms 4 and 5.
8. The initialization for Algorithm 6 is computationally costly, but subsequent evaluations require much less CPU time than evaluations for Algorithm 5.

Remark 7.2 The CPU timings for Algorithms 1, 2 and 3 are independent of the particular distributions of ω and x , whereas the timings for Algorithms 4 and 5 are sensitive to the distributions of these vectors.

8 Generalizations and Conclusions

The results of this paper can be generalized in the following ways:

1. Simple modifications to Algorithms 1, 2 and 3 will allow the efficient application of the linear transformations $F_1, G_1, H_1 : \mathbb{C}^{N+1} \rightarrow \mathbb{C}^{M+1}$ defined by

$$F_1(\alpha)_j = \sum_{k=0}^N \alpha_k \cdot e^{i\omega_k \cdot 2\pi j/M} \quad \text{for } j = -M/2, \dots, M/2, \quad (101)$$

$$G_1(\beta)_j = \sum_{k=-N/2}^{N/2} \beta_k \cdot e^{ikx_j} \quad \text{for } j = 0, \dots, M, \quad (102)$$

$$H_1(\gamma)_j = \sum_{k=0}^N \gamma_k \cdot e^{i\omega_k x_j} \quad \text{for } j = 0, \dots, M. \quad (103)$$

These changes have been implemented.

2. The algorithms of this paper also assume that $\omega_k \in [-N/2, N/2]$ and $x_j \in [-\pi, \pi]$. Other distributions can be handled by appropriately partitioning the vectors ω and x , treating each partition separately and finally combining the results. The following observation describes translation operators which can be used for each partition, in combination with one of Algorithms 1, 2 or 3.

Observation 8.1 *Let $a, b > 0, c, d > 0$ be real numbers and suppose that $\omega_k \in [a-b, a+b]$ for $k = 0, \dots, N$ and $x_j \in [c-d, c+d]$ for $j = 0, \dots, M$. Then we can write*

$$\sum_{k=0}^N \alpha_k \cdot e^{i\omega_k x_j} = e^{iax_j} \cdot \sum_{k=0}^N \alpha_k \cdot e^{i(\omega_k - a)c} \cdot e^{i(\omega_k - a)(x_j - c)} \quad (104)$$

$$= e^{iax_j} \cdot \sum_{k=0}^N \alpha'_k \cdot e^{i\omega'_k x'_j}, \quad (105)$$

where

$$\begin{aligned} \alpha'_k &= \alpha_k e^{i(\omega_k - a)c}, \\ \omega'_k &= (\omega_k - a)d/\pi, \\ x'_j &= (x_j - c)\pi/d \in [-\pi, \pi]. \end{aligned} \quad (106)$$

Remark 8.2 Such an algorithm will perform efficiently when the points within a partition are close together and there are very few partitions, and not so efficiently if the points are widely separated and there are many partitions. Most cases likely to be encountered in practice fall in the former category.

3. A paper describing a set of algorithms based on a different interpolation technique is currently in preparation.
4. One of the more far-reaching extensions of the results of this paper is a set of algorithms for problems in higher dimensions. Investigations into this are currently in progress.
5. The Helmholtz equation in 2 dimensions is given by

$$\nabla^2 \phi + \kappa^2 \phi = 0, \quad (107)$$

and has particular solutions of the form

$$\phi(x, y) = e^{i\mu x} \cdot e^{i\nu y}, \quad (108)$$

where $\mu^2 + \nu^2 = \kappa^2$. Solutions of this equation consist of linear combinations of such functions subject to some boundary conditions, and the results of this paper admit a generalization which constitutes a fast Helmholtz solver.

In conclusion, a group of algorithms has been presented for the rapid application and inversion of matrices of the Fourier kernel. These problems can be viewed as generalizations of the Discrete Fourier Transform, and the algorithms, while making use of certain simple results from analysis, are very versatile, and have wide-ranging potential applications in many branches of mathematics, science and engineering.

Appendix A

In this section we present numerical estimates of the errors in approximating e^{icx} using q terms of the form $e^{bx^2}e^{ikx}$ for the values of m , b and q used by the algorithms. As q is independent of the choice of c , we only considered the case $c = 0$. The results are presented in Table 7.

The approximation

$$\tilde{\phi}(x) = e^{bx^2} \cdot \sum_{k=-q/2}^{q/2} \rho_k e^{ikx} \quad (109)$$

to the constant function $\phi(x) = 1$ with ρ_k defined by equation (17) was computed at $n = 1000$ equally spaced nodes x_k in $[-\frac{\pi}{m}, \frac{\pi}{m}]$ and the entries in the table are defined as follows:

- E_∞ is the maximum absolute error defined by the formula

$$E_\infty = \max_{1 \leq k \leq n} |\tilde{\phi}(x_k) - \phi(x_k)|. \quad (110)$$

- E_2 is the relative L_2 error defined by the formula

$$E_2 = \sqrt{\frac{\sum_{k=1}^n |\tilde{\phi}(x_k) - \phi(x_k)|^2}{n}}. \quad (111)$$

- E_B is the error bound of Theorem 2.10 defined by the formula

$$E_B = e^{-b\pi^2(1-1/m^2)} \cdot (4b + 9). \quad (112)$$

TABLE 7

m	b	q	E_∞	E_2	E_B
2	0.5993	10	0.825 E-05	0.176 E-05	0.135 E-00
2	1.5629	28	0.400 E-13	0.580 E-14	0.163 E-03

References

- [1] G. Dahlquist and A. Björck, *Numerical Methods*, Prentice Hall Inc., Englewood Cliffs, N.J., 1974.
- [2] A. Dutt, *A Fast Algorithm for the Evaluation of Trigonometric Series*, Technical Report 841, Yale Computer Science Department, 1991.
- [3] I.S. Gradshteyn and I.M. Ryzhik, *Table of Integrals, Series and Products*, Academic Press Inc., 1980.
- [4] J. Stoer and R. Bulirsch, *Introduction to Numerical Analysis*, Springer-Verlag, New York, 1980.
- [5] H. Joseph Weaver, *Theory of Discrete and Continuous Fourier Analysis*, Wiley, 1989.