



Politecnico
di Bari

Repository Istituzionale dei Prodotti della Ricerca del Politecnico di Bari

Computing the matrix Mittag-Leffler function with applications to fractional calculus

This is a post print of the following article

Original Citation:

Computing the matrix Mittag-Leffler function with applications to fractional calculus / Garrappa, R., Popolizio, M.. - In: JOURNAL OF SCIENTIFIC COMPUTING. - ISSN 0885-7474. - STAMPA. - 77:1(2018), pp. 129-153. [10.1007/s10915-018-0699-5]

Availability:

This version is available at <http://hdl.handle.net/11589/169250> since: 2022-06-07

Published version

DOI:10.1007/s10915-018-0699-5

Publisher:

Terms of use:

(Article begins on next page)

Computing the matrix Mittag–Leffler function with applications to fractional calculus *

Roberto Garrappa

Università degli Studi di Bari, Dipartimento di Matematica, Bari, Italy

Member of the INdAM Research group GNCS

`roberto.garrappa@uniba.it`

Marina Popolizio

Università del Salento, Dipartimento di Matematica e Fisica “Ennio De Giorgi”. Lecce, Italy

Member of the INdAM Research group GNCS

`marina.popolizio@unisalento.it`

Submitted on June 5th, 2017

Abstract

The computation of the Mittag-Leffler (ML) function with matrix arguments, and some applications in fractional calculus, are discussed. In general the evaluation of a scalar function in matrix arguments may require the computation of derivatives of possible high order depending on the matrix spectrum. Regarding the ML function, the numerical computation of its derivatives of arbitrary order is a completely unexplored topic; in this paper we address this issue and three different methods are tailored and investigated. The methods are combined together with an original derivatives balancing technique in order to devise an algorithm capable of providing high accuracy. The conditioning of the evaluation of matrix ML functions is also studied. The numerical experiments presented in the paper show that the proposed algorithm provides high accuracy, very often close to the machine precision.

Keywords: Mittag–Leffler function, Matrix function, Derivatives of the Mittag–Leffler function, Fractional calculus, Conditioning.

1 Introduction

When the Swedish mathematician Magnus Gustaf Mittag-Leffler introduced, at the beginning of the twentieth century, the function that successively would have inherited his name [42, 43], he perhaps ignored the importance it would have gained several years later; indeed, although the introduction of the Mittag-Leffler (ML) function was motivated just by the analysis of divergent series, this function has nowadays a fundamental role in the

*This is the post-print of the paper: R.Garrappa, M.Popolizio, “Computing the matrix Mittag–Leffler function with applications to fractional calculus”, *Journal of Scientific Computing* (Springer), October 2018, Volume 77, Issue 1, pp 129-153, doi: 10.1007/s10915-018-0699-5, available at <https://doi.org/10.1007/s10915-018-0699-5>. This work is supported under the GNCS-INdAM 2017 project “Analisi numerica per modelli descritti da operatori frazionari”.

theory of operators of fractional (i.e., non integer) order and it has become so important in this context as to be defined the “Queen function of fractional calculus” [26, 28, 39].

For complex parameters α and β , with $\Re(\alpha) > 0$, the ML function is defined by means of the series

$$E_{\alpha,\beta}(z) = \sum_{j=0}^{\infty} \frac{z^j}{\Gamma(\alpha j + \beta)}, \quad z \in \mathbb{C}, \quad (1)$$

where $\Gamma(\cdot)$ is the Euler’s gamma function. $E_{\alpha,\beta}(z)$ is an entire function of order $\rho = 1/\Re(\alpha)$ and type $\sigma = 1$ and it is clearly a generalization of the exponential function to which it reduces when $\alpha = \beta = 1$ since for $j \in \mathbb{N}$ it is $\Gamma(j + 1) = j!$. Throughout the paper we consider real values for α and β which is the case of interest for common applications.

Despite the great interest in fractional calculus, few works have so far concerned the accurate evaluation of the ML function. Indeed this is in and of itself challenging and expensive and few efficient algorithms for this task have been devised only recently (e.g., see [18, 23, 27, 54, 56]).

The ML function with matrix arguments is just as valuable as its scalar version and it can be successfully employed in several applications: for the efficient and stable solution of systems of fractional differential equations (FDEs), to determine the solution of certain multiterm FDEs, in control theory and in other related fields (e.g., see [20, 21, 22, 44, 46, 48, 55]).

The aim of this paper is to discuss numerical techniques for the evaluation of the ML function with matrix arguments; in particular we are interested in methods working with high accuracy, if possible very close to the machine precision, in order to provide results which can be considered virtually exact in finite precision arithmetic.

Incidentally, and under a more general perspective, we cannot get out of mentioning the recent interest in the numerical approximation of matrix functions for applications in a wide range of areas, for which we refer the reader to the textbook by Higham [33] and the readable papers [7, 15, 29, 34].

For general matrix functions the Schur–Parlett algorithm [6] represents the most powerful method. It is based on the Schur decomposition of the matrix argument combined with the Parlett recurrence to evaluate the matrix function on the triangular factor [25]. Additionally, for the diagonal part, the Taylor series of the scalar function is used. This last task requires the knowledge of the derivatives of the underlying scalar function up to an order depending on the eigenvalues of the matrix argument. In particular, high orders are needed in the presence of multiple or highly clustered eigenvalues. This is a crucial issue which complicates the matrix case with respect to the scalar one.

Facing the evaluation of derivatives of the ML function, as requested by the Schur–Parlett algorithm, is a demanding task other than a rather unexplored topic (except, as far as we know, for one work dealing with just the first order derivative [27]). A great portion of this paper is therefore devoted to discuss and deeply investigate three different methods for the evaluation of derivatives of the ML function; the rationale for introducing different methods is that each of them properly works in limited regions of the complex plane and for different parameter ranges. Thanks to our investigation, we are therefore able to tune a combined algorithm which applies, in an accurate way, to matrices with any eigenvalues location.

Besides, we analyze the conditioning of the computation of matrix ML functions to understand the sensitivity of the matrix function to perturbations in the data. We thus give a first contribution on this topic which we think can be of interest for readers interested in numerical applications of the matrix ML function, where rounding errors and

perturbations are inevitable.

This paper is organized as follows: in Section 2 we survey some of the most common applications of the ML function with matrix arguments. In Section 3 we discuss the generalization of a scalar function to matrix arguments and we review the Schur-Parlett algorithm for its numerical computation. Some methods for the accurate and efficient evaluation of derivatives of the ML function, as requested by the Schur-Parlett algorithm, are hence described in Section 4: we study in detail each method in order to identify strengths and weaknesses and provide some criteria for the selection of the most suitable method in each situation. The conditioning of the ML function is studied in Section 5 and, finally, in Section 6 we present the results of some numerical experiments.

2 Applications of matrix Mittag-Leffler functions

The evaluation of matrix ML functions is not just a curiosity-driven problem; several practical applications can indeed benefit from calculating the value of $E_{\alpha,\beta}$ in matrix arguments.

This section presents some important applications involving the numerical evaluation of matrix ML functions. The aim is to give just a flavor of the assorted fields in which these objects are required, while for their thorough discussion we refer the reader to the existing literature (e.g., see [26, 30, 49]).

2.1 Solution of systems of FDEs

The matrix ML function is crucial to explicitly represent the solution of a linear system of FDEs of order $\alpha > 0$

$$D_t^\alpha Y(t) = AY(t), \quad Y^{(\ell)}(0) = Y_0^\ell, \quad \ell = 0, \dots, m-1, \quad (2)$$

where $A \in \mathbb{R}^{n \times n}$, $Y(t) : [0, T] \rightarrow \mathbb{R}^n$, $m = \lceil \alpha \rceil$ is the smallest integer greater or equal to α and D_t^α is the Caputo's fractional derivative

$$D_t^\alpha y(t) \equiv \frac{1}{\Gamma(m-\alpha)} \int_0^t \frac{y^{(m)}(u)}{(t-u)^{\alpha+1-m}} du,$$

with $y^{(m)}$ denoting the standard integer-order derivative.

It is immediate to verify that the exact solution of the system (2) is

$$Y(t) = \sum_{\ell=0}^{m-1} t^\ell E_{\alpha, \ell+1}(t^\alpha A) Y_0^\ell \quad (3)$$

and, once some tool for the computation of the matrix ML function is available, the solution of (2) can be evaluated, possibly with high accuracy, directly at any time $t > 0$. Conversely, the commonly used step-by-step methods usually involve considerable computational costs because of the persistent memory of fractional operators [9].

2.2 Solution of time-fractional partial differential equations

Given a linear time-fractional partial differential equation

$$D_t^\alpha u(t, x) = \nabla_x^2 u(t, x) + f(t, x), \quad (4)$$

subject to some initial and boundary conditions, a preliminary discretization along the spatial variables allows to recast (4) in terms of a semi-linear system of FDEs in the form

$$D_t^\alpha U(t) = AU(t) + F(t), \quad (5)$$

with $U(t) = (u_1(t), \dots, u_{N_x}(t))^T$, being $u_i(t)$ an approximation of $u(x_i, t)$ on a partition $x_1, \dots, x_{N_x}$ of the spatial domain, and $F(t)$ obtained from the source term and the boundary conditions. By assuming for simplicity $0 < \alpha < 1$, the exact solution of the semi-discretized system (5) can be formulated as

$$U(t) = E_{\alpha,1}(t^\alpha A)U_0 + \int_0^t (t-\tau)^{\alpha-1} E_{\alpha,\alpha}((t-\tau)^\alpha A)F(\tau) d\tau. \quad (6)$$

A common approach to solve (5) relies on product-integration rules which actually approximate, in the standard integral formulation of (5), the vector-field $AU(t) + F(t)$ by piecewise interpolating polynomials. This method, however, has severe limitations for convergence due to the non-smooth behavior of $U(t)$ at the origin [13]. However, the same kind of approximation when used in (6), where just $F(t)$ is replaced by polynomials, does not suffer from the same limitations [17] and it is possible to obtain high order methods under the reasonable assumption of a sufficiently smooth source term $F(t)$. In addition, solving (4) by using matrix ML functions in (6) allows also to overcome stability issues since the usual stiffness of the linear part of (5) is solved in a virtually exact way.

2.3 Solution of linear multiterm FDEs

Let us consider a linear multiterm FDE of *commensurate* order $0 < \alpha < 1$

$$\sum_{k=0}^n a_k D_t^{k\alpha} y(t) = f(t), \quad (7)$$

with $D_t^{k\alpha}$ derivatives of Caputo type and associated initial conditions

$$y(0) = b_0, y'(0) = b_1, \dots, y^{(m-1)}(0) = b_{m-1}, \quad m = \lceil n\alpha \rceil.$$

When $n\alpha < 1$, or when $\alpha \in \mathbb{Q}$, the FDE (7) can be reformulated as a linear system of FDEs (see [10, 11] or [9, Theorems 8.1-8.2]). For instance, if $\alpha = p/q$, with $p, q \in \mathbb{N}$, one puts $y_1(t) = y(t)$, $y_k(t) = D_t^{1/q} y_{k-1}(t)$ and, by introducing the vector notation $Y(t) = (y_1(t), y_2(t), \dots, y_N(t))^T$, it is possible to rewrite (7) as

$$D_t^{1/q} Y(t) = AY(t) + e_N f(t), \quad Y(0) = Y_0 \quad (8)$$

where $e_N = (0, 0, \dots, 0, 1)^T \in \mathbb{R}^N$, $Y_0 \in \mathbb{R}^N$ is composed in a suitable way on the basis of the initial values b_j , the coefficient matrix $A \in \mathbb{R}^{N \times N}$ is a companion matrix and the size of problem (8) is $N = np$.

Also in this case it is possible to express the exact solution by means of matrix ML functions as

$$Y(t) = E_{\alpha,1}(t^\alpha A)Y_0 + \int_0^t (t-\tau)^{\alpha-1} E_{\alpha,\alpha}((t-\tau)^\alpha A)e_N f(\tau) d\tau \quad (9)$$

and hence approximate the above integral by some well-established technique. As a special case, with a polynomial source term $f(t) = c_0 + c_1 t + \dots c_s t^s$ it is possible to explicitly represent the exact solution of (8) as

$$Y(t) = E_{\alpha,1}(t^\alpha A)Y_0 + \sum_{\ell=0}^s \ell! c_\ell t^{\alpha+\ell} E_{\alpha,\alpha+\ell+1}(t^\alpha A) e_N. \quad (10)$$

We have also to mention that, as deeply investigated in [38], the solution of the linear multiterm FDE (7) can be expressed as the convolution of the given input function $f(t)$ with a special multivariate ML function of scalar type, an alternative formulation which can be also exploited for numerically solving (7).

It is beyond the scope of this work to compare the approach proposed in [38] with numerical approximations based on the matrix formulation (9); it is however worthwhile to highlight the possibility of recasting the same problem in different, but analytically equivalent, formulations: one involving the evaluation of scalar multivariate ML functions and the other based on the computation of standard ML functions with matrix arguments; clearly the two approaches present different computational features whose advantages and disadvantages could be better investigated in the future.

2.4 Controllability and observability of fractional linear systems

For a continuous-time control system the complete controllability describes the possibility of finding some input signal $u(t)$ such that the system can be driven, in a finite time, to any final state x_T starting from the fixed initial state x_0 . The dual concept of observability instead describes the possibility of determining, at any time t , the state $x_0 = x(t_0)$ of the system from the subsequent history of the input signal $u(t)$ and of the output $y(t)$.

Given a linear time-invariant system of fractional order $\alpha > 0$

$$\begin{cases} D_t^\alpha x(t) = Ax(t) + Bu(t) \\ y(t) = Cx(t) + Du(t) \end{cases} \quad (11)$$

controllability and observability are related to the corresponding Gramian matrices, defined respectively as [2, 40]

$$\mathcal{C}_\alpha(t) := \int_0^t E_{\alpha,\alpha}((t-\tau)^\alpha A) B B^T E_{\alpha,\alpha}((t-\tau)^\alpha A^T) d\tau$$

and

$$\mathcal{O}_\alpha(t) = \int_0^t E_{\alpha,1}(\tau^\alpha A^T) C^T C E_{\alpha,1}((t-\tau)^\alpha A) d\tau$$

(see also [41] for problems related to the control of fractional order systems).

In particular, system (11) is controllable (resp. observable) on $[0, T]$ if the matrix $\mathcal{C}_\alpha(t)$ (resp. the matrix $\mathcal{O}_\alpha(t)$) is positive definite.

Although alternative characterizations of controllability and observability are available in terms of properties of the state matrix A , the computation of the controllability Gramian $\mathcal{C}_\alpha(t)$ (resp. the observability Gramian $\mathcal{O}_\alpha(t)$) can be used to find an appropriate control to reach a state x_T starting from the initial state x_0 (resp. to reconstruct the state $x(t)$ from the input $u(t)$ and the output $y(t)$). Moreover, computing the Gramians is also useful for input-output transfer function models derived from linearization of nonlinear state-space models around equilibrium points that depend on working conditions of the real modeled systems (see, for instance, [36, 37]).

3 The Mittag-Leffler functions with matrix arguments: definitions

3.1 Theoretical background

Given a function f of scalar arguments and a matrix A , the problem of finding a suitable definition for $f(A)$ goes back to Cayley (1858) and it has been broadly analyzed since then. The extension from scalar to matrix arguments is straightforward for simple functions, like polynomials, since it trivially consists in substituting the scalar argument with the matrix A . This is also the case of some transcendental functions, like the ML, for which (1) becomes

$$E_{\alpha,\beta}(A) = \sum_{j=0}^{\infty} \frac{A^j}{\Gamma(\alpha j + \beta)} \quad (12)$$

(this series representation is useful for defining $E_{\alpha,\beta}(A)$ but usually it cannot be used for computation since it presents the same issues, which will be discussed later in Subsection 4.1, of the use of (1) for the scalar function, amplified by the computation increase due to the matrix argument).

In more general situations a preliminary definition is of fundamental importance to understand the meaning of evaluating a function on a matrix.

Definition 1. Let A be a $n \times n$ matrix with s distinct eigenvalues $\lambda_1, \dots, \lambda_s$ and let n_i be the *index* of λ_i , that is, the smallest integer k such that $(A - \lambda_i I)^k = 0$ with I denoting the $n \times n$ identity matrix. Then the function f is said to be defined on the spectrum of A if the values $f^{(j)}(\lambda_i)$, $j = 0, \dots, n_i - 1$, $i = 1, \dots, s$ exist.

Then for functions defined on the spectrum of A the *Jordan canonical form* can be of use to define matrix functions.

Definition 2. Let f be defined on the spectrum of A and let A have the Jordan canonical form

$$A = Z J Z^{-1} = Z \operatorname{diag}(J_1, \dots, J_p) Z^{-1}, \quad J_k = \begin{bmatrix} \lambda_k & 1 & & \\ & \lambda_k & \ddots & \\ & & \ddots & 1 \\ & & & \lambda_k \end{bmatrix} \in \mathbb{C}^{m_k \times m_k}.$$

Then

$$f(A) = Z f(J) Z^{-1} = Z \operatorname{diag}(f(J_1), \dots, f(J_p)) Z^{-1},$$

with

$$f(J_k) = \begin{bmatrix} f(\lambda_k) & f'(\lambda_k) & \dots & \frac{f^{(m_k-1)}(\lambda_k)}{(m_k-1)!} \\ & f(\lambda_k) & \ddots & \vdots \\ & & \ddots & f'(\lambda_k) \\ & & & f(\lambda_k) \end{bmatrix}.$$

This definition highlights a fundamental issue related to matrix functions: when multiple eigenvalues are present, the function derivatives need to be computed. In our context this is a serious aspect to consider since no numerical method has been considered till now for the derivatives of the (scalar) ML function, except for the first order case [27].

In practice, however, the Jordan canonical form is rarely used since the similarity transformation may be very ill conditioned and, even more seriously, it cannot be reliably computed in floating point arithmetic.

3.2 Schur-Parlett algorithm

Diagonalizable matrices are favorable arguments for simple calculations; indeed a matrix is diagonalizable if there exist a nonsingular matrix P and a diagonal matrix $D = \text{diag}(\lambda_1, \dots, \lambda_n)$ such that $A = P^{-1}DP$. In this case $f(A) = P^{-1}f(D)P$ and $f(D)$ is still a diagonal matrix with principal entries $f(\lambda_i)$. Clearly the computation involves only scalar arguments. Problems arise when the matrix P is nearly singular or in general badly conditioned, since the errors in the computation are related to the condition number of the similarity matrix P (see [33]). For these reasons well conditioned transformations are preferred in this context.

The Schur decomposition is a standard instrument in linear algebra and we refer to [25] for a complete description and for a comprehensive treatment of the related issues. It nowadays represents the starting point for general purpose algorithms to compute matrix functions thanks also to its backward stability. Once the Schur decomposition for the matrix argument A is computed, say $A = QTQ^*$ with Q unitary and T upper triangular, then $f(A) = Qf(T)Q^*$, for any well defined f . The focus thus moves to the computation of the matrix function for triangular arguments.

In general this is a delicate issue which can be the cause of severe errors if not cleverly accomplished. Here we describe the enhanced algorithm due to Davies and Higham [6]: once the triangular matrix T is computed, it is reordered and blocked so to get a matrix $\tilde{T}$ such that each block $\tilde{T}_{ij}$ has clustered eigenvalues and distinct diagonal blocks have far enough eigenvalues.

To evaluate $f(\tilde{T}_{ii})$ a Taylor series is considered about the mean σ of the eigenvalues of $\tilde{T}_{ii}$. So, if m is the dimension of $\tilde{T}_{ii}$ then $\sigma = \text{trace}(\tilde{T}_{ii})/m$ and

$$f(\tilde{T}_{ii}) = \sum_{k=0}^{\infty} \frac{f^{(k)}(\sigma)}{k!} M^k$$

with M such that $\tilde{T}_{ii} = \sigma I + M$. The powers of M will decay quickly after the $(m-1)$ -st since, by construction, the eigenvalues of $\tilde{T}_{ii}$ are “sufficiently” close.

Once the diagonal blocks $f(\tilde{T}_{ii})$ are evaluated, the rest of $f(\tilde{T})$ is computed by means of the block form of the Parlett recurrence [33]. Finally the inverse similarity transformations and the reordering are applied to get $f(A)$.

4 Computation of derivatives of the ML function

The Schur-Parlett algorithm substantially relies on the knowledge of the derivatives of the scalar function in order to compute the corresponding matrix function. Since in the case of the ML function $E_{\alpha,\beta}(z)$ the derivatives are not immediately at disposal, it is necessary to devise efficient and reliable numerical methods for their computation for arguments located in any region of the complex plane and up to any possible order, in dependence of the spectrum of A .

To this purpose we present and discuss three different approaches, namely

1. series expansion,
2. numerical inversion of the Laplace transform,
3. summation formulas.

Although the three approaches are based on formulas which are equivalent from an analytical point of view, their behavior when employed for numerical computation differs in a substantial way.

We therefore investigate in details weaknesses and strengths in order to identify the range of arguments for which each of them provides more accurate results and we tune, at the end of this Section, an efficient algorithm implementing, in a combined way, the different methods.

4.1 Series expansion

After a term-by-term derivation, it is immediate to construct from (1) the series representation of the derivatives of the ML function

$$\frac{d^k}{dz^k} E_{\alpha,\beta}(z) = \sum_{j=k}^{\infty} \frac{(j)_k}{\Gamma(\alpha j + \beta)} z^{j-k}, \quad k \in \mathbb{N}, \quad (13)$$

with $(x)_k$ denoting the falling factorial

$$(x)_k = x(x-1) \cdots (x-k+1).$$

At a glance, the truncation of (13) to a finite number $J \in \mathbb{N}$ of terms, namely

$$\frac{d^k}{dz^k} E_{\alpha,\beta}(z) \approx \sum_{j=k}^J \frac{(j)_k}{\Gamma(\alpha j + \beta)} z^{j-k}, \quad (14)$$

may appear as the most straightforward method to compute derivatives of the ML function. Anyway the computation of (14) presents some not negligible issues entailing serious consequences if not properly addressed.

We have first to mention that in IEEE-754 double precision arithmetic the Gamma function can be evaluated only for arguments not greater than 171.624, otherwise overflow occurs; depending on α and β the upper bound for the number of terms in (14) is hence

$$J_{\max} := \left\lfloor \frac{171.624 - \beta}{\alpha} \right\rfloor,$$

with $\lfloor x \rfloor$ being the largest integer less than or equal to x ; as a consequence the truncated series (14) can be used only when the $J_{\max}$ -th term is small enough to disregard the remaining terms in the original series (13). By assuming a target accuracy τ as the threshold for truncating the summation, a first bound on the range of admissible values of z is then obtained by requiring that

$$|z| \leq \left(\tau \frac{\Gamma(\alpha J_{\max} + \beta)}{(J_{\max})_k} \right)^{\frac{1}{J_{\max}-k}}.$$

The effects of round-off errors (related to the finite-precision arithmetic used for the computation) may however further reduce, in a noticeable way, the range of admissible arguments z for which the series expansion can be employed. When $|z| > 1$ and $|\arg(z)| > \alpha\pi/2$, some of the terms in (14) are indeed in modulus much more larger than the modulus of $E_{\alpha,\beta}(z)$; actually, small values of $E_{\alpha,\beta}(z)$ result as the summation of large terms with alternating signs, a well-known source of heavy numerical cancellation (e.g., see [32]).

To decide whether or not to accept the results of the computation, it is necessary to devise a reliable estimate of the round-off error. We inherently assume that in finite-precision arithmetic the sum $S_1 = c_0 + c_1$ of any two terms c_0 and c_1 leads to $\hat{S}_1 = (c_0 + c_1)(1 + \delta_1)$, with $|\delta_1| < \epsilon$ and $\epsilon > 0$ the machine precision. Then, for $J \geq 2$, the general summation $S_J = c_0 + c_1 + \dots + c_J$ of $J + 1$ terms is actually computed as

$$\hat{S}_J = S_J + (c_0 + c_1) \sum_{j=1}^J \delta_j + c_2 \sum_{j=2}^J \delta_j + c_3 \sum_{j=3}^J \delta_j + \dots + c_J \delta_J \quad (15)$$

where $|\delta_1|, |\delta_2|, \dots, |\delta_J| < \epsilon$ and terms proportional to $\mathcal{O}(\epsilon^2)$ have been discarded. It is thus immediate to derive the following bound for the round-off error

$$|S_J - \hat{S}_J| \leq \epsilon \left(J|c_0| + \sum_{j=1}^J (J - j + 1)|c_j| \right). \quad (16)$$

The order by which the terms c_j are summed is relevant especially for the reliability of the above estimator; clearly, an ascending sorting of $|c_j|$ makes the estimate (16) more conservative and hence more useful for practical use. It is therefore advisable, especially in the more compelling cases (namely, for arguments z with large modulus and $|\arg(z)| > \alpha\pi/2$) to perform a preliminary sorting of the terms in (14) with respect to their modulus.

An alternative estimate of the round-off error can be obtained after reformulating (15) as

$$\hat{S}_J = S_J + \sum_{j=1}^J S_j \delta_j$$

from which one can easily derive

$$|S_J - \hat{S}_J| \leq \epsilon \sum_{j=1}^J |S_j|. \quad (17)$$

The bound (17) is surely more sharp than (16) and it does not involve a large amount of extra computation since it is possible to evaluate S_J by storing all the partial sums S_j ; anyway (17) is based on the exact values of the partial sums S_j which are actually not available; since just the computed values $\hat{S}_j$ can be employed, formula (17) can underestimate the round-off error. We find useful to use a mean value between the bounds provided by (16) and (17).

Unfortunately, only “a posteriori” estimates of the round-off error are feasible and it seems not possible to prevent from performing the computation before establishing whether or not to accept the results. Obviously, on the basis of the above discussion, it is possible to reduce the range of arguments for which the evaluation by (14) is attempted and, anyway, if well organized, the corresponding algorithm runs in reasonable time.

4.2 Numerical inversion of the Laplace transform

After a change of the summation index, it is straightforward from (13) to derive the following representation of the derivatives of the ML function

$$\frac{d^k}{dz^k} E_{\alpha, \beta}(z) = \sum_{j=0}^{\infty} \frac{(j+k)_k z^j}{\Gamma(\alpha j + \alpha k + \beta)} = k! E_{\alpha, \alpha k + \beta}^{k+1}(z), \quad (18)$$

where

$$E_{\alpha,\beta}^{\gamma}(z) = \frac{1}{\Gamma(\gamma)} \sum_{j=0}^{\infty} \frac{\Gamma(j+\gamma)z^j}{j!\Gamma(\alpha j + \beta)} \quad (19)$$

is a three parameter ML function which is recognized in literature as the Prabhakar function [47] (the relationship between the derivatives of the ML function and the Prabhakar function has been recently highlighted also in [16, 45]).

This function, which is clearly a generalization of the ML function (1), since $E_{\alpha,\beta}^1(z) = E_{\alpha,\beta}(z)$, is recently attracting an increasing attention in view of its applications in the description of relaxation properties in a wide range of anomalous phenomena (see, for instance [5, 19, 24, 35, 50]). For this reason some authors recently investigated the problem of its numerical computation [18, 51].

The approach proposed in [18] is based on the numerical inversion of the Laplace transform and extends a method previously proposed in [23, 55]. Since it allows to perform the computation with an accuracy very close to the machine precision, the extension to derivatives of the ML function appears of particular interest.

For any $t > 0$ and $z \in \mathbb{C}$, the Laplace transform (LT) of the function $t^{\beta-1}E_{\alpha,\beta}^{\gamma}(t^{\alpha}z)$ is given by

$$\mathcal{L}(t^{\beta-1}E_{\alpha,\beta}^{\gamma}(t^{\alpha}z); s) = \frac{s^{\alpha\gamma-\beta}}{(s^{\alpha}-z)^{\gamma}}, \quad \Re(s) > 0, \quad |zs^{-\alpha}| < 1. \quad (20)$$

After selecting $t = 1$ and denoting

$$H_k(s; z) = \frac{s^{\alpha-\beta}}{(s^{\alpha}-z)^{k+1}},$$

with a branch-cut imposed on the negative real semi axis to make $H_k(s; z)$ single-valued, the derivatives of the ML function can be evaluated by inverting the LT (20), namely by recasting (18) in form of the integral

$$\frac{d^k}{dz^k}E_{\alpha,\beta}(z) = \frac{k!}{2\pi i} \int_{\mathcal{C}} e^s H_k(s; z) ds =: I_k(z), \quad (21)$$

over a contour $\mathcal{C}$ in the complex plane encompassing at the left all the singularities of $H_k(s; z)$.

The trapezoidal rule has been proved to possess excellent properties for the numerical quadrature of contour integrals of analytic functions [4, 53] but satisfactory enough results can be obtained also in the presence of singularities if the contour is selected in a suitable way.

In particular, in order to truncate, after a reasonable small number of terms, the infinite sum resulting from the application of the trapezoidal rule to (21) it is necessary that the integrand decays in a fast way; therefore contours beginning and ending in the left half of the complex plane must be selected. Our preference is for parabolic contours

$$\mathcal{C} : \sigma(u) = \mu(iu + 1)^2, \quad -\infty < u < \infty$$

whose simplicity allows accurate estimates of the errors and a fine tuning of the main parameters involved in the integration.

Given a grid $u_j = jh$, $j = -N, \dots, N$, with constant step-size $h > 0$, the truncated trapezoidal rule applied to (21) reads as

$$I_k^{[N]}(z) = \frac{k!h}{2\pi i} \sum_{j=-N}^N e^{\sigma(u_j)} H_k(\sigma(u_j); z) \sigma'(u_j).$$

For choosing the integration parameters μ , h and N the two main components of the error are considered, that is the discretization and the truncation errors. The balancing of these errors, according to the procedure described first in [55], and hence applied to the ML function in [18, 23], allows to select optimal parameters in order to achieve a prescribed tolerance $\tau > 0$.

The error analysis is performed on the basis of the distance between the contour and the integrand singularities. As observed in [18], in addition to the branch-point singularity at $s = 0$, the set of the non zero poles s^* of $H_k(s; z)$ in the main Riemann sheet, i.e. such that $-\pi < \text{Arg}(s^*) \leq \pi$, is given by

$$S^* = \left\{ |s|^{1/\alpha} e^{i \frac{\text{Arg}(s) + 2j\pi}{\alpha}} \mid -\frac{\alpha}{2} - \frac{\text{Arg}(s)}{2\pi} < j \leq \frac{\alpha}{2} - \frac{\text{Arg}(s)}{2\pi}, j \in \mathbb{Z} \right\}. \quad (22)$$

Depending on α and z , one can therefore expect singularities in any number and in any region of the complex plane but the selection of contours so wide to encompass all the singularities usually leads to numerical instability. In order to avoid that round-off errors dominate discretization and truncation errors, thus preventing from achieving the target accuracy, it is indeed necessary to select contours which do not extend too far in the right half-part of the complex plane; in particular, it has been found [18] that the main parameter μ of the parabolic contour must satisfy

$$\mu < \log \tau - \log \epsilon, \quad (23)$$

with ϵ the machine precision and $\tau > \epsilon$ the target accuracy.

Since in many cases it is not possible to leave all the singularities at the left of the contour and simultaneously satisfy (23), some of the singularities must be removed from (21) to gain more freedom in the selection of the contour $\mathcal{C}$; the corresponding residues are therefore subtracted in view of the Cauchy's residue theorem and hence

$$\frac{d^k}{dz^k} E_{\alpha, \beta}(z) = k! \sum_{s^* \in S_{\mathcal{C}}^*} \text{Res}(e^s H_k(s; z), s^*) + \frac{k!}{2\pi i} \int_{\mathcal{C}} e^s H_k(s; z) ds, \quad (24)$$

where $S_{\mathcal{C}}^* \subseteq S^*$ is the set of the singularities of $H_k(s; z)$ laying to the right of $\mathcal{C}$ (which now is not constrained to encompass all the singularities) and $\text{Res}(e^s H_k(s; z), s^*)$ is the residue of $e^s H_k(s; z)$ at s^* ; since the selected branch-cut, the singularity at the origin is always kept at the left of $\mathcal{C}$.

To make equation (24) applicable for computation we give an explicit representation of the residues in terms of elementary functions by means of the following result.

Proposition 1. *Let $\alpha > 0$, $\beta \in \mathbb{R}$, $k \in \mathbb{N}$, $z \neq 0$ and $s_{\star} \in S_{\star}$ one of the poles of $H_k(s; z)$. Then*

$$\text{Res}(e^s H_k(s; z), s_{\star}) = \frac{1}{\alpha^{k+1}} e^{s_{\star}} s_{\star}^{1-\alpha k - \beta} P_k(s_{\star})$$

where $P_k(x) = p_0^{(k)} + p_1^{(k)}x + \dots + p_k^{(k)}x^k$ is the k -th degree polynomial whose coefficients are

$$p_j^{(k)} = \frac{1}{j!} \sum_{\ell=0}^{k-j} \frac{(\alpha - \beta)_{\ell}}{\ell!} H_{k-j-\ell}^{(k)}, \quad j = 0, 1, \dots, k$$

with the coefficients $H_j^{(k)}$ evaluated recursively as

$$H_0^{(k)} = 1, \quad H_j^{(k)} = -\frac{1}{\alpha} \sum_{\ell=1}^j \binom{\alpha}{\ell+1} \left(\frac{k\ell}{j} + 1 \right) H_{j-\ell}^{(k)}, \quad j = 1, 2, \dots, k$$

and the generalized binomial coefficients are defined as

$$\binom{\alpha}{j} = \frac{\alpha(\alpha-1)(\alpha-2)\cdots(\alpha-j+1)}{j!} = (-1)^j \frac{\Gamma(j-\alpha)}{\Gamma(-\alpha)\Gamma(j+1)}.$$

Proof. Since s_* is a pole of order $k+1$ of $e^s H_k(s; z)$, the corresponding residue can be evaluated by differentiation

$$\text{Res}(e^s H_k(s; z), s_*) = \frac{1}{k!} \lim_{s \rightarrow s_*} \frac{d^k}{ds^k} e^s H_k(s; z) (s - s_*)^{k+1}. \quad (25)$$

To study the term $H_k(s; z)(s - s_*)^{k+1}$ we observe that for s sufficiently close to s_* we can write

$$s^\alpha = s_*^\alpha \left(\frac{s - s_*}{s_*} + 1 \right)^\alpha = s_*^\alpha + s_*^\alpha \sum_{j=1}^{\infty} \binom{\alpha}{j} \frac{(s - s_*)^j}{s_*^j},$$

and, hence, since $s_*^\alpha = z$ we have

$$s^\alpha - z = s_*^\alpha \sum_{j=1}^{\infty} \binom{\alpha}{j} \frac{(s - s_*)^j}{s_*^j}.$$

After observing that $\binom{\alpha}{1} = \alpha$, it is a direct computation to provide the expansion

$$\begin{aligned} H_k(s; z)(s - s_*)^{k+1} &= s^{\alpha-\beta} \left(\frac{s^\alpha - z}{s - s_*} \right)^{-(k+1)} \\ &= \frac{s_*^{-(\alpha-1)(k+1)}}{\alpha^{k+1}} s^{\alpha-\beta} \left(1 + \frac{1}{\alpha} \sum_{j=1}^{\infty} \binom{\alpha}{j+1} \frac{(s - s_*)^j}{s_*^j} \right)^{-(k+1)}. \end{aligned}$$

The reciprocal of the $(k+1)$ -th power of the unitary formal power series can be evaluated by applying the Miller's formula [31, Theorem 1.6c] thanks to which it is possible to provide the expansion

$$H_k(s; z)(s - s_*)^{k+1} = \frac{s_*^{-(\alpha-1)(k+1)}}{\alpha^{k+1}} s^{\alpha-\beta} \sum_{j=0}^{\infty} H_j^{(k)} \frac{(s - s_*)^j}{s_*^j}.$$

To evaluate (25) we now introduce, for convenience, the functions

$$F_1(s) = e^s, \quad F_2(s) = s^{\alpha-\beta}, \quad F_3(s) = \sum_{j=0}^{\infty} H_j^{(k)} \frac{(s - s_*)^j}{s_*^j},$$

and hence, thanks to (25) we are able to write

$$\frac{d^k}{ds^k} (e^s H_k(s; z)(s - s_*)^{k+1}) = \frac{s_*^{-(\alpha-1)(k+1)}}{\alpha^{k+1}} \sum_{k_1=0}^k \sum_{k_2=0}^{k-k_1} \frac{k! F_1^{(k_1)}(s) F_2^{(k_2)}(s) F_3^{(k-k_1-k_2)}(s)}{k_1! k_2! (k - k_1 - k_2)!}.$$

It is immediate to evaluate the derivatives of the functions F_1 , F_2 and F_3 and their limit as $s \rightarrow s_*$

$$\lim_{s \rightarrow s_*} F_1^{(j)}(s) = e^{s_*}, \quad \lim_{s \rightarrow s_*} F_2^{(j)}(s) = (\alpha - \beta)_j s_*^{\alpha-\beta-j}, \quad \lim_{s \rightarrow s_*} F_3^{(j)}(s) = \frac{j! H_j^{(k)}}{s_*^j},$$

and hence we are able to compute

$$\text{Res}(e^s H_k(s; z), s_*) = \frac{s_*^{-(\alpha-1)(k+1)}}{\alpha^{k+1}} \sum_{k_1=0}^k \sum_{k_2=0}^{k-k_1} \frac{e^{s_*} (\alpha - \beta)_{k_2} s_*^{\alpha-\beta-k_2} H_{k-k_1-k_2}^{(k)}}{k_1! k_2! s_*^{k-k_1-k_2}}$$

from which we obtain

$$\text{Res}(e^s H_k(s; z), s_*) = \frac{s_*^{1-\alpha k-\beta}}{\alpha^{k+1}} e^{s_*} \sum_{k_1=0}^k \frac{s_*^{k_1}}{k_1!} \sum_{k_2=0}^{k-k_1} \frac{(\alpha - \beta)_{k_2} H_{k-k_1-k_2}^{(k)}}{k_2!}$$

and the proof immediately follows. $\square$ $\square$

We observe that, with respect to the formula presented in [52], the result of Proposition 1 is slightly more general and easier to be used for computation. The coefficients of polynomials $P_k(x)$ can be evaluated without any particular difficulty by means of a simple algorithm. For ease of presentation we show here their first few instances

$$\begin{aligned} P_0(x) &= 1 \\ P_1(x) &= (\alpha - \beta + 1) + x \\ P_2(x) &= \frac{1}{2} + \left(\frac{3\alpha}{2} - \beta + \frac{3}{2} \right) x + \left(\alpha^2 - \frac{3\alpha\beta}{2} + \frac{3\alpha}{2} + \frac{\beta^2}{2} - \beta + \frac{1}{2} \right) x^2 \end{aligned}$$

We have to consider that the algorithm devised in [18] for the inversion of the LT of the ML function needs some adjustment to properly evaluate the derivatives. As k increases a loss of accuracy can indeed be expected for very small arguments also on the negative real semi axis (namely the branch-cut). This is a consequence of the fact that the origin actually behaves as a singularity, of order $k + 1$, and hence its effects must be included in the balance of the error.

Moreover, for high order derivatives, round-off errors must be expected especially when a singularity of $H_k(s; z)$ lies not so close to the origin (thus to make very difficult the choice of a contour encompassing all of them) but at the same time not so far from the origin (thus, making difficult to select contour in the region between the origin and the singularity). To safely deal with situations of this kind we will introduce, later in Subsection 4.4, a technique which keeps as low as possible the order of the derivatives to be evaluated by the numerical inversion of the Laplace transform.

4.3 Summation formulas

The Djrbashian's formula [14] allows to express the first derivative of the ML function in terms of two instances of the same function, according to

$$\frac{d}{dz} E_{\alpha, \beta}(z) = \frac{E_{\alpha, \beta-1}(z) + (1 - \beta) E_{\alpha, \beta}(z)}{\alpha z}, \quad (26)$$

and, obviously, (26) holds only for $z \neq 0$; anyway, from (13) it is immediate to verify that

$$\left. \frac{d^k}{dz^k} E_{\alpha, \beta}(z) \right|_{z=0} = \frac{k!}{\Gamma(\alpha k + \beta)}. \quad (27)$$

Equation (26) provides an exact formulation for the first-order derivative of the ML function but we are able to generalize it also to higher-order derivatives by means of the following result.

	$k = 0$	$k = 1$	$k = 2$	$k = 3$
$c_0^{(k)}$	1	$1 - \beta$	$(1 - \beta)(1 - \beta - \alpha)$	$(1 - \beta)(1 - \beta - \alpha)(1 - \beta - 2\alpha)$
$c_1^{(k)}$		1	$3 - 2\beta - \alpha$	$(1 - \beta)(1 - \beta - \alpha) + (3 - 2\beta - \alpha)(2 - \beta - 2\alpha)$
$c_2^{(k)}$			1	$6 - 3\beta - 3\alpha$
$c_3^{(k)}$				1

Table 1: First few coefficients (29) of the summation formulas (28) and (32).

Proposition 2 (Summation formula of Djrbashian type). *Let $\alpha > 0, \beta \in \mathbb{R}$ and $z \neq 0$. For any $k \in \mathbb{N}$ it is*

$$\frac{d^k}{dz^k} E_{\alpha, \beta}(z) = \frac{1}{\alpha^k z^k} \sum_{j=0}^k c_j^{(k)} E_{\alpha, \beta-j}(z), \quad (28)$$

where $c_0^{(0)} = 1$ and the remaining coefficients $c_j^{(k)}$, $j = 0, 1, \dots, k$, are recursively evaluated as

$$c_j^{(k)} = \begin{cases} (1 - \beta - \alpha(k-1))c_0^{(k-1)} & j = 0 \\ c_{j-1}^{(k-1)} + (1 - \beta - \alpha(k-1) + j)c_j^{(k-1)} & j = 1, \dots, k-1 \\ 1 & j = k \end{cases} \quad (29)$$

Proof. We proceed by induction on k . For $k = 0$ the proof is obvious and for $k = 1$ it is a consequence of (26). Assume now that (28) holds for $k-1$; then, the application of standard derivative rules, together with the application of (26), allows us to write

$$\begin{aligned} \frac{d^k}{dz^k} E_{\alpha, \beta}(z) &= \frac{1}{\alpha^{k-1} z^{k-1}} \sum_{j=0}^{k-1} \frac{c_j^{(k-1)}}{\alpha z} \left(E_{\alpha, \beta-j-1}(z) - (\beta - j - 1) E_{\alpha, \beta-j}(z) \right) \\ &\quad - \frac{(k-1)z^{k-2}}{\alpha^{k-1} z^{2k-2}} \sum_{j=0}^{k-1} c_j^{(k-1)} E_{\alpha, \beta-j}(z) \end{aligned}$$

and the proof follows after defining the coefficients $c_j^{(k)}$ as in (29). $\square$ $\square$

The Djrbashian summation formula (SF) of Proposition 2 allows to represent the derivatives of the ML function in terms of a linear combination of values of the same function; clearly it can be employed in actual computation once a reliable procedure for the evaluation of $E_{\alpha, \beta}(z)$ for any parameter α and β is available, as it is the case of the method described in [18].

Unfortunately, it seems not possible to provide an explicit closed form for the coefficients $c_j^{(k)}$ in the Djrbashian SF (28); anyway the recursive relationship (29) is very simple to compute and the first few coefficients, holding up to the derivative of third order, are shown in Table 1.

An alternative approach to compute the coefficients $c_j^{(k)}$ can be derived by observing, from (27), that the derivatives of the ML function have finite values at $z = 0$; hence, since $E_{\alpha, \beta}$ is entire, all possible negative powers of z in (28) must necessarily vanish and, since $c_k^{(k)} = 1$, it is immediate to verify that the coefficients $c_j^{(k)}$ are solution of the linear system

$$\sum_{j=0}^{k-1} \frac{1}{\Gamma(\alpha\ell + \beta - j)} c_j^{(k)} = -\frac{1}{\Gamma(\alpha\ell + \beta - k)}, \quad \ell = 0, \dots, k-1.$$

To study the stability of the Djrbashian SF (28), we observe that since the derivatives of the ML function have finite values at $z = 0$, see Eq. (27), from (28) it follows

$$\sum_{j=0}^k c_j^{(k)} E_{\alpha, \beta-j}(z) = \mathcal{O}(z^k), \quad |z| \rightarrow 0. \quad (30)$$

As a consequence, differences of almost equal values are involved by (28) when $|z|$ is small; the unavoidable numerical cancellation, further magnified since the division by z^k , makes this summation formula not reliable for values of z close to the origin, especially for derivatives of high order.

To overcome these difficulties we can exploit again the relationship (18) between the ML derivatives and the Prabhakar function together with the formula (see [47])

$$E_{\alpha, \beta}^{k+1}(z) = \frac{E_{\alpha, \beta-1}^k(z) + (1 - \beta + \alpha k) E_{\alpha, \beta}^k(z)}{\alpha k}, \quad (31)$$

expressing the Prabhakar function as difference of two instances of the same function but with a smaller value of the third parameter. We are hence able to provide the following alternative SF of Prabhakar type.

Proposition 3 (Summation formula of Prabhakar type). *Let $\alpha > 0$ and $\beta \in \mathbb{R}$. For any $k \in \mathbb{N}$ it is*

$$\frac{d^k}{dz^k} E_{\alpha, \beta}(z) = \frac{1}{\alpha^k} \sum_{j=0}^k c_j^{(k)} E_{\alpha, \alpha k + \beta - j}(z), \quad (32)$$

where $c_j^{(k)}$, $j = 0, 1, \dots, k$, are the coefficients (29) of Proposition 2.

Proof. The claim is obvious for $k = 0$; for $k = 1$ it is a consequence of (31) and then one can easily prove it by induction with algebraic transformations quite similar to the ones used to prove (28). $\square$ $\square$

For small values of $|z|$, the Prabhakar SF (32) is expected to be less affected by round-off errors than (28). Indeed not only the sum in (32) returns $\mathcal{O}(1)$ when $|z| \rightarrow 0$ but possible round-off errors are not amplified by division by z^k as in (28).

In Figures 1 and 2 we compare the errors given by the SF (28) of Djrbashian type with those provided by the SF (32) of Prabhakar type; here and in the following tests, the reference values for the derivatives of $E_{\alpha, \beta}(z)$ are obtained by computing the series expansion (13) in high-precision floating-point arithmetic with 2000 digits by using Maple 15. As we can clearly see, when $|z|$ is small the Prabhakar SF (32) is surely more accurate whilst for larger values of $|z|$ their performances are almost similar.

4.4 Combining algorithms with derivatives balancing

The mathematical considerations in the previous subsection, supported by several experiments we conducted on a wide range of parameters α and β , and for arguments z in different sectors of the complex plane, indicate that the Prabhakar SF (32) is capable of providing accurate and satisfactory results in the majority of the cases.

Anyway we have to consider that larger errors should be expected as the degree of the derivative increases; this a consequence of the large number of operations involved by (32) and, in particular, of the accumulation of round-off errors. Unfortunately, also the numerical inversion of the LT can show a loss of accuracy for high order derivatives since

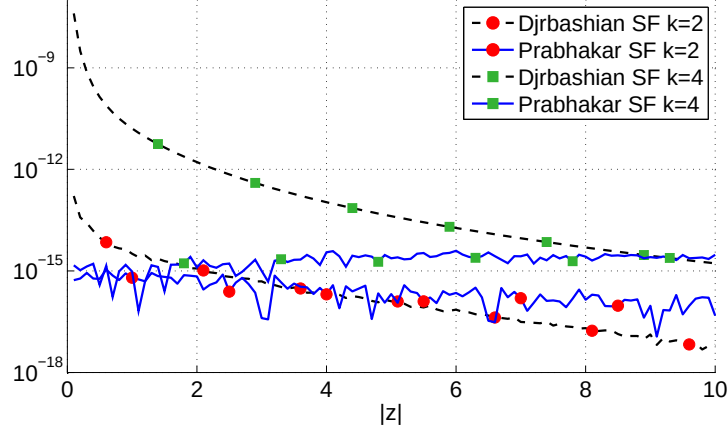


Figure 1: Comparison of errors provided by Djrbashian and Prabhakar SFs for $\alpha = 0.6$, $\beta = 1.0$ and $\arg(z) = \pi$.

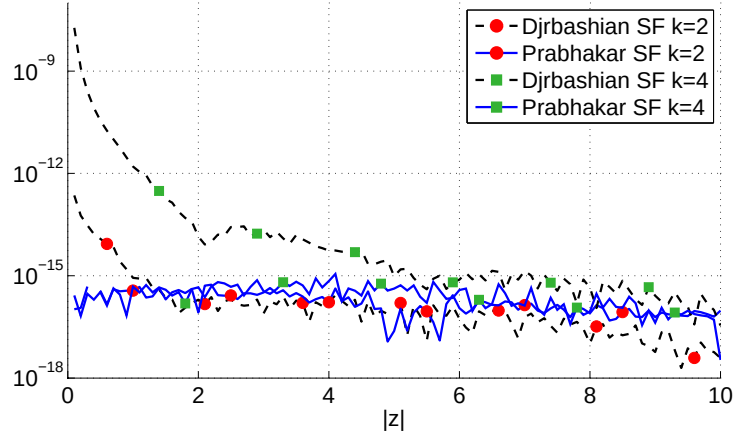


Figure 2: Comparison of errors provided by Djrbashian and Prabhakar SFs for $\alpha = 0.8$, $\beta = 1.2$ and $\arg(z) = \pi/2$.

the strong singularity in $H_k(s; z)$ may lead to large round-off errors in the integration process.

In critical situations it is advisable to adopt mixed strategies by combining different algorithms. For instance, it is natural to use the truncated series expansion (14) for small arguments and switch to other methods as $|z|$ becomes larger; the error estimates discussed in Subsection 4.1 can be of help for implementing an automatic procedure performing this switching.

In addition, since all methods reduce their accuracy for high order derivatives, it could be necessary to adopt a strategy operating a balance of the derivatives thus to bring the problem to the evaluation of lower order derivatives. To this purpose, Proposition 3 can be generalized in order to express k -th order derivatives of the ML function in terms of lower order derivatives instead of the ML function itself.

Proposition 4. *Let $\alpha > 0, \beta \in \mathbb{R}$. For any $k \in \mathbb{N}$ and $p \leq k$ it is*

$$\frac{d^k}{dz^k} E_{\alpha, \beta}(z) = \frac{1}{\alpha^{k-p}} \sum_{j=0}^{k-p} c_j^{(k-p)} \frac{d^p}{dz^p} E_{\alpha, (k-p)\alpha + \beta - j}(z), \quad (33)$$

where $c_j^{(k)}$, $j = 0, 1, \dots, k$, are the coefficients (29) of Proposition 2.

The proof of the above result is similar to the proof of Proposition 3 and is omitted for shortness.

Thanks to Proposition 4 it is possible to combine the algorithms introduced in the previous subsections and employ the SF (33) with underlying derivatives evaluated by the numerical inversion of the LT. Clearly, the aim is to avoid working with high-order derivatives in the numerical inversion of the LT and, at the same time, reduce the number of terms in (32).

In Figure 3 we illustrate the positive effects on accuracy of performing the derivatives balancing. As we can see, the SF (32) tends to lose accuracy for higher order derivatives (here a 5-th order has been considered); the numerical inversion of the LT is instead more accurate but when the integration contour is constrained to stay close to the singularity the error can rapidly increase (see the peak in the interval $[0, 2]$). The derivatives balancing (operated here for $p = 1$) allows to keep the error at a moderate level on the whole interval, without any worrisome peak.

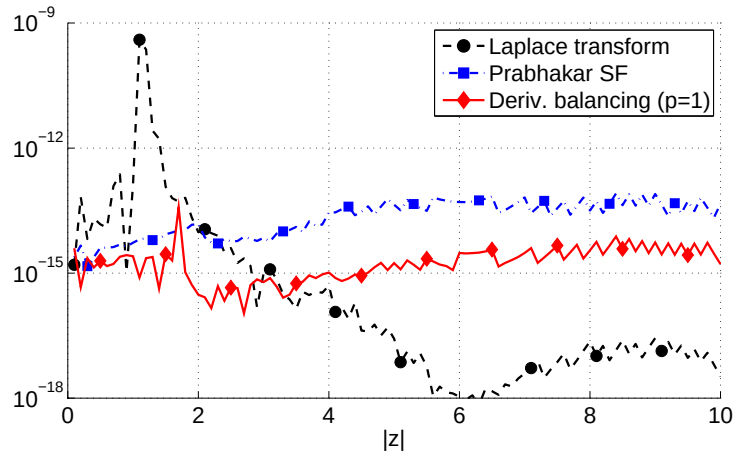


Figure 3: Comparison of errors for 5-th derivative of the ML function when $\alpha = 0.6$, $\beta = 1.0$ and $\arg(z) = \pi/2$.

Finally, we present a couple of tests in which the different techniques described in the paper are combined to provide accurate results for the evaluation of the derivatives of the ML function. Roughly speaking, the algorithm uses the truncated series for small arguments and hence, on the basis of the error estimate, it switches to the Prabhakar SF (32) using the numerical inversion of the Laplace transform according to a derivative balancing approach.

In several cases we have observed that the computation of lower order derivatives turns out to be more accurate (see Figure 4) although it is not possible to infer a general trend as we can observe from Figure 5. Anyway, in all cases (we have performed experiments, not reported here for brevity, on a wide set of arguments and parameters) the errors remain in a range $10^{-13} \sim 10^{-15}$ which can be surely considered satisfactory for applications. The plotted error is $|E - \tilde{E}|/(1 + |E|)$ where E is a reference value and $\tilde{E}$ the evaluated derivative.

5 Conditioning of the matrix ML function

To hone the analysis of the matrix ML function it is of fundamental importance to measure its sensitivity to small perturbations in the matrix argument; indeed, even when dealing

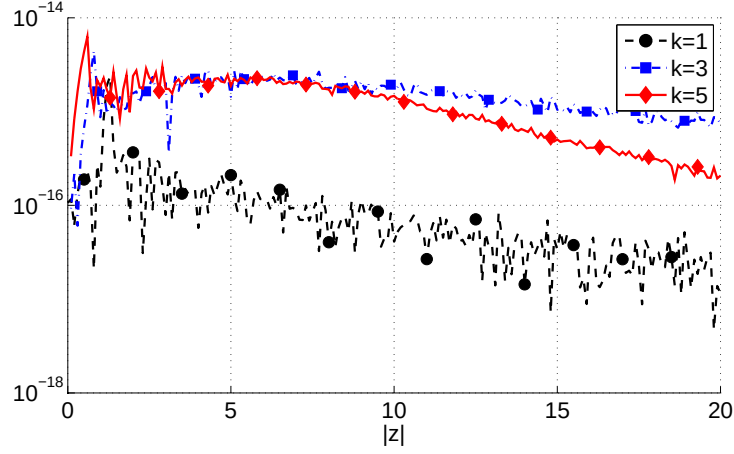


Figure 4: Errors in the computation of k -th derivatives of the ML function when $\alpha = 0.6$, $\beta = 0.6$ and $\arg(z) = 0.8\pi$.

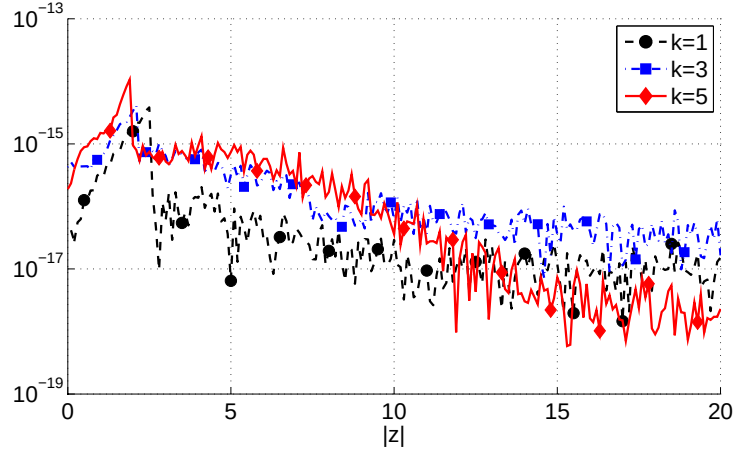


Figure 5: Errors in the computation of k -th derivative of the ML function when $\alpha = 0.8$, $\beta = 1.2$ and $\arg(z) = \pi/2$.

with exact data, the rounding errors can severely affect the computation, exactly as in the scalar case. A useful tool for this analysis is the *condition number* whose definition can be readily extended to matrix functions [33].

Definition 3. For a general matrix function f , if $\|\cdot\|$ denotes any matrix norm, the *absolute condition number* is defined as

$$\kappa_a(f, A) = \lim_{\varepsilon \rightarrow 0} \sup_{\|E\| \leq \varepsilon} \frac{\|f(A + E) - f(A)\|}{\varepsilon}$$

while the *relative condition number* is defined as

$$\kappa_r(f, A) = \lim_{\varepsilon \rightarrow 0} \sup_{\|E\| \leq \varepsilon \|A\|} \frac{\|f(A + E) - f(A)\|}{\varepsilon \|f(A)\|}$$

being the matrix E a perturbation to the original data A .

For the actual computation of these quantities the Fréchet derivative of the matrix function f is of use. We recall that given a matrix function f defined in $\mathbb{C}^{n \times n}$, its Fréchet derivative at a point $A \in \mathbb{C}^{n \times n}$ is a linear function $L(A, E)$ such that

$$f(A + E) - f(A) - L(A, E) = o(\|E\|).$$

The norm of the Fréchet derivative is defined as

$$\|L(A)\| := \max_{Z \neq 0} \frac{\|L(A, Z)\|}{\|Z\|}$$

and it turns out to be a useful tool for the sensitivity analysis of the function itself thanks to the following result.

Theorem 4. [33] *The absolute and relative condition numbers are given by*

$$\begin{aligned}\kappa_a(f, A) &= \|L(A)\|, \\ \kappa_r(f, A) &= \frac{\|L(A)\|}{\|f(A)\|} \|A\|.\end{aligned}$$

The computation of the Fréchet derivative is thus essential to derive the function conditioning. Nowadays only specific functions have effective algorithms to compute it, like the exponential, the logarithm and fractional powers [8, 33]. For general matrix functions the reference algorithm for the Fréchet derivative is due to Al-Mohy and Higham [1, 34].

To give a flavor of the conditioning of the ML function we consider 25 test matrices of dimension 10×10 taken from the Matlab `gallery`. Figure 6 reports the 1-norm relative condition number $\kappa_r(E_{\alpha,1}, A)$ computed by the `funm_condest1` routine from the Matrix Function Toolbox [33]. Amazingly, for all but the chebspec matrix (the one plotted in red line with small circles), the conditioning is almost independent on α ; this shows that the ML function has almost the same conditioning of the exponential function for any α . For most of the test matrices the ML function conditioning is large; in particular, for small values of α the ML conditioning for the chebspec Matlab matrix is huge. An explanation could be that the Schur-Parlett algorithm does not work properly with this matrix, as stressed in [33]; indeed, this matrix is similar to a Jordan block with eigenvalue 0 while its computed eigenvalues lie roughly on a circle with centre 0 and radius 0.2. The well-known ill-conditioning of Pascal matrices (for $n = 10$ it is $\kappa(A) \approx 8.13 \times 10^9$) clearly leads to a severe ill-conditioning of the matrix function too (see the straight blue line).

In general, more insights can derive from a comparison of this conditioning with the 1-norm condition of A , that we denote with $\kappa(A)$. To this purpose, Figure 7 displays the ratio $\kappa_r(E_{\alpha,1}, A)/\kappa(A)$. From this plot we can appreciate that in most of the cases the ML function is better conditioned than the matrix argument itself, while in the remaining cases the conditioning is not too much amplified.

Moreover in Figure 9 of the next section we show that $\kappa_r(E_{\alpha,1}, A)\epsilon$ well bounds the error between the computed and the reference values of $E_{\alpha,1}(A)$; this property seems to reveal that our technique to compute $E_{\alpha,1}(A)$ is *normwise forward stable*, according to the definition given in [32].

6 Numerical experiments

In addition to the numerical tests already presented in the previous sections, in this section we verify the accuracy obtained for some test matrices. All the experiments have been carried out in Matlab ver. 8.3.0.532 (R2014a) and, in order to plot the errors, also in these experiments reference values have been evaluated by means of Maple 15 with a floating-point precision of 2000 digits.

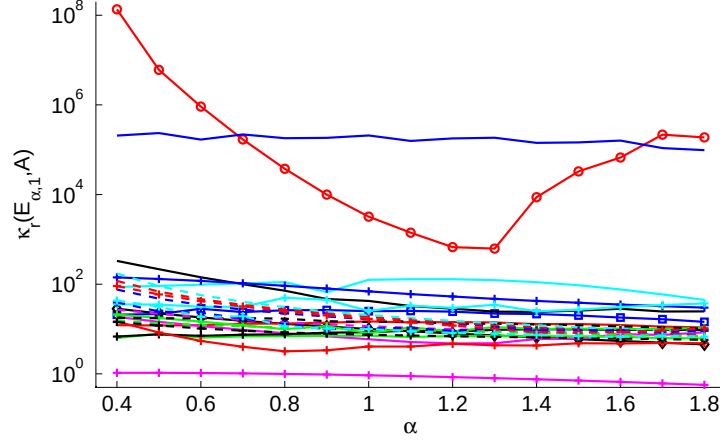


Figure 6: Conditioning $\kappa_r(E_{\alpha,1}, A)$ for 25 matrices of dimension 10×10 from the Matlab gallery.

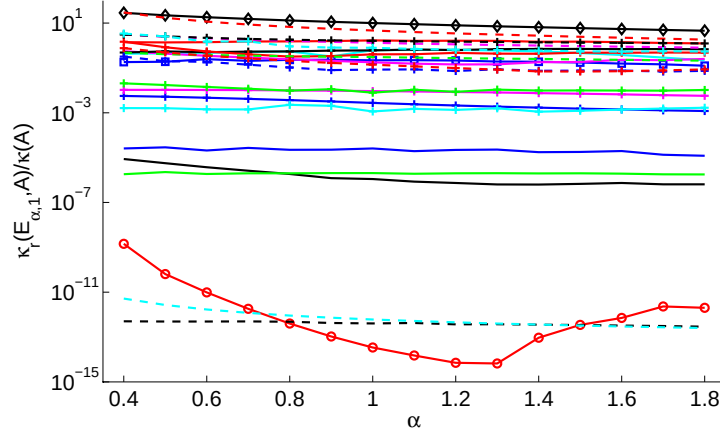


Figure 7: Conditioning ratio $\kappa_r(E_{\alpha,1}, A)/\kappa(A)$ for 25 matrices of dimension 10×10 from the Matlab gallery.

For the Schur-Parlett algorithm we use the Matlab `funm` function, whilst the derivatives of the ML function are evaluated by means of the combined algorithm described in Subsection 4.4 and making use of the derivatives balancing. The corresponding Matlab code, which will be used in all the subsequent experiments, is available in the file exchange service of the Mathworks website¹. The plotted errors between reference values E and approximations $\tilde{E}$ are evaluated as $\|E - \tilde{E}\|/(1 + \|E\|)$, where $\|\cdot\|$ is the usual Frobenius norm.

Example 1: the Redheffer matrix

For the first test we use the Redheffer matrix from the matrix gallery of Matlab. This is a quite simple binary matrix whose elements $a_{i,j}$ are equal to 1 when $j = 1$ or i divides j , otherwise they are all equal to 0; however, the eigenvalues are highly clustered: for a $n \times n$ Redheffer matrix the number of eigenvalues equal to 1 is exactly $n - \lfloor \log_2 n \rfloor - 1$ (see [3]). As a consequence, computing $E_{\alpha,\beta}(-A)$ demands the evaluation of high order derivatives (up to the 24-th order in our experiments) thus making this matrix appealing

¹www.mathworks.com/matlabcentral/fileexchange/66272-mittag-leffler-function-with-matrix-arguments

for test purposes.

In Figure 8 we observe the error for some values of α with $n \times n$ Redheffer matrices of increasing size (n is selected in the interval $[4, 20]$ and $\beta = 1$ is always used). As we can clearly see, the algorithm is able to provide an excellent approximation, with an error very close to machine precision, also for matrices of notable size.

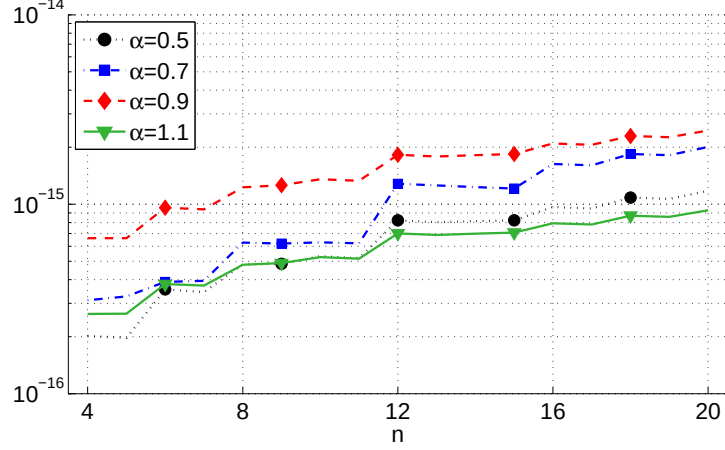


Figure 8: Errors in the computation of $E_{\alpha,1}(-A)$ with A the Redheffer matrix as its dimension changes

Example 2: matrices with selected eigenvalues

The second experiment involves 4 different matrices of size 40×40 suitably built in order to have eigenvalues with moderately high multiplicities. In practice we fix some values and we consider diagonal matrices having them as principal entries, repeated according to the multiplicities we want (as listed in Table 2). Then, by similarity transformations, we get the full matrices with the desired spectrum. To reduce rounding errors for the similarity transformations we use orthogonal matrices provided by the `gallery` Matlab function.

	Eigenvalues (multiplicities)					Max derivative order
Matrix 1	$\pm 1.0(5)$	$\pm 1.0001(4)$	$\pm 1.001(4)$	$\pm 1.01(4)$	$\pm 1.1(3)$	15
Matrix 2	$\pm 1.0(8)$	$2(8)$	$-5(8)$	$-10(8)$		3
Matrix 3	$-1(2)$	$-5(2)$	$1 \pm 10i(6)$	$-4 \pm 1.5i(6)$	$\pm 5i(6)$	3
Matrix 4	$1(4)$	$1.0001(4)$	$1.001(4)$	$1 \pm 10i(7)$	$-4 \pm 1.5i(7)$	7

Table 2: Eigenvalues (with multiplicities) for the 40×40 matrices of the Example 2.

The maximum order of the derivatives required by the computation is indicated in the last column of Table 2. As we can observe from Figure 9 (solid lines), despite the large size of the matrices and the high clustering of the eigenvalues, it is possible to evaluate $E_{\alpha,\beta}(A)$ with a reasonable accuracy. In the same plot we can also appreciate (dotted lines) how the bounds $\kappa_r(E_{\alpha,1}, A)\epsilon$ (see Section 5) give a reasonable estimate for the error.

Example 3: solution of a multiterm FDE

In the third experiment we consider an application to the multiterm FDE

$$2y(t) + 6D_0^\alpha y(t) + 7D_0^{2\alpha} y(t) + 4D_0^{3\alpha} y(t) + D_0^{4\alpha} y(t) = f(t) \quad (34)$$

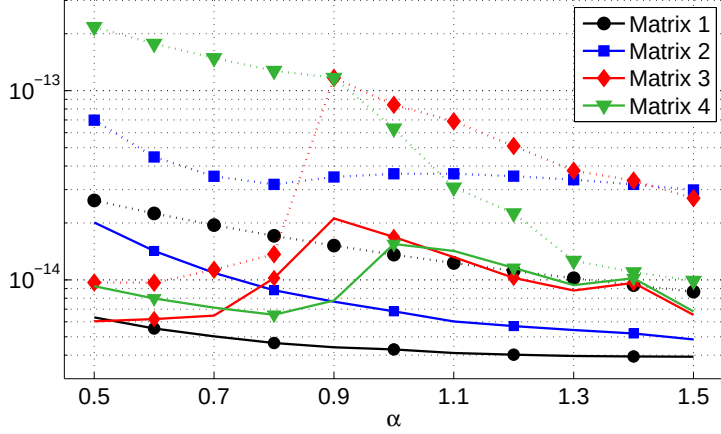


Figure 9: Relative errors (solid lines) for $E_{\alpha,1}(A)$ for the four test matrices $A \in \mathbb{R}^{40 \times 40}$ of Example 2 and conditioning $\kappa_r(E_{\alpha,1}, A)\epsilon$ (dotted lines).

with homogeneous initial conditions, which is first reformulated in terms of the linear system (8) with a companion coefficient matrix A , of size 16×16 , with both simple and double eigenvalues. The reference solution, for $\alpha = 0.8$ and an external source term $f(t) = -t^2/2 + 2t$, is shown in Figure 10 on the interval $[0, 6]$.

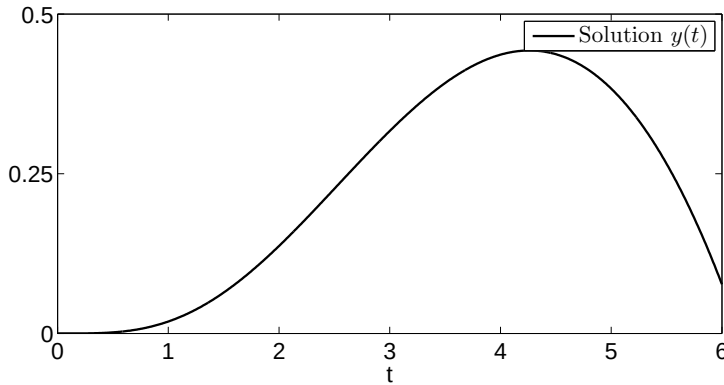


Figure 10: Solution of the multiterm FDE (34) for $\alpha = 0.8$.

In Figure 11 we compare the error of the solution evaluated by using the matrix ML function in (10) with those obtained by a step-by-step trapezoidal product integration (PI) rule generalized to multiterm FDEs according to the ideas discussed in [12]; in the plot we simply denote this approach as “Trapez. PI” with indicated the step-size h used for the computation.

Also in this case we observe the excellent results obtained thanks to the use of matrix ML functions. We must also note that in order to achieve a comparable accuracy, the Trapez. PI would require very small step-sizes, with a remarkable computational cost especially for integration on intervals of large size; the use of matrix ML functions instead allows to directly evaluate the solution at any time t .

7 Concluding remarks

In this paper we have discussed the evaluation of the ML function with matrix arguments and illustrated some remarkable applications. Since one of the most efficient algorithms for

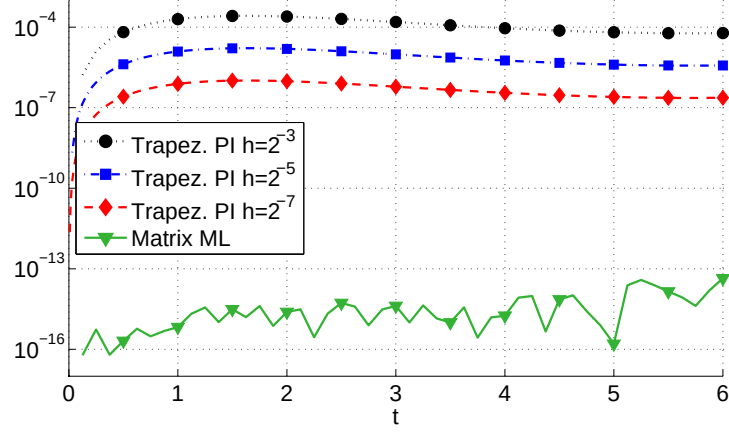


Figure 11: Errors in the solution of the multiterm FDE (34) for $\alpha = 0.8$.

the evaluation of matrix functions requires the evaluation of the derivatives of the original scalar function, a large portion of the paper has been devoted to present different methods for the accurate evaluation of derivatives of the ML function, a subject which, as far as we know, has not been faced before except for first derivatives. We have also discussed some techniques for combining the different methods in an efficient way with the aim of devising an algorithm capable of achieving high accuracy with matrices having any kind of spectrum. The analysis on the conditioning has shown that it is possible to keep errors under control when evaluating matrix ML functions. Finally, the numerical experiments presented at the end of the paper have highlighted the possibility of evaluating the matrix ML function with great accuracy also in the presence of not simple matrix arguments.

References

- [1] Al-Mohy, A.H., Higham, N.J.: The complex step approximation to the Fréchet derivative of a matrix function. *Numer. Algorithms* **53**(1), 113–148 (2010)
- [2] Balachandran, K., Govindaraj, V., Ortigueira, M., Rivero, M., Trujillo, J.: Observability and controllability of fractional linear dynamical systems. *IFAC Proceedings Volumes* **46**(1), 893 – 898 (2013)
- [3] Barrett, W.W., Jarvis, T.J.: Spectral properties of a matrix of Redheffer. *Linear Algebra Appl.* **162/164**, 673–683 (1992). *Directions in matrix theory* (Auburn, AL, 1990)
- [4] Bornemann, F., Laurie, D., Wagon, S., Waldvogel, J.: The SIAM 100-digit challenge. Society for Industrial and Applied Mathematics (SIAM), Philadelphia, PA (2004)
- [5] Colombaro, I., Giusti, A., Vitali, S.: Storage and dissipation of energy in Prabhakar viscoelasticity. *Mathematics* **6**(2), 15 (2018). DOI 10.3390/math6020015
- [6] Davies, P.I., Higham, N.J.: A Schur-Parlett algorithm for computing matrix functions. *SIAM J. Matrix Anal. Appl.* **25**(2), 464–485 (electronic) (2003)
- [7] Del Buono, N., Lopez, L., Politi, T.: Computation of functions of Hamiltonian and skew-symmetric matrices. *Math. Comp. Simul.* **79**(4), 1284–1297 (2008)

- [8] Dieci, L., Papini, A.: Conditioning and Padé approximation of the logarithm of a matrix. *SIAM J. Matrix Analysis Applications* **21**(3), 913–930 (2000)
- [9] Diethelm, K.: The analysis of fractional differential equations, *Lecture Notes in Mathematics*, vol. 2004. Springer-Verlag, Berlin (2010)
- [10] Diethelm, K., Ford, N.J.: Numerical solution of the Bagley-Torvik equation. *BIT* **42**(3), 490–507 (2002)
- [11] Diethelm, K., Ford, N.J.: Multi-order fractional differential equations and their numerical solution. *Appl. Math. Comput.* **154**(3), 621–640 (2004)
- [12] Diethelm, K., Luchko, Y.: Numerical solution of linear multi-term initial value problems of fractional order. *J. Comput. Anal. Appl.* **6**(3), 243–263 (2004)
- [13] Dixon, J.: On the order of the error in discretization methods for weakly singular second kind Volterra integral equations with nonsmooth solutions. *BIT* **25**(4), 624–634 (1985)
- [14] Džrbašjan [Djrbashian], M.M.: Harmonic analysis and boundary value problems in the complex domain, *Operator Theory: Advances and Applications*, vol. 65. Birkhäuser Verlag, Basel (1993). Translated from the manuscript by H. M. Jerbashian and A. M. Jerbashian [A. M. Dzhrbashyan]
- [15] Frommer, A., Simoncini, V.: Matrix functions. In: Model order reduction: theory, research aspects and applications, *Math. Ind.*, vol. 13, pp. 275–303. Springer, Berlin (2008)
- [16] Garra, R., Garrappa, R.: The Prabhakar or three parameter Mittag-Leffler function: theory and application. *Commun. Nonlinear Sci. Numer. Simul.* **56**, 314–329 (2018)
- [17] Garrappa, R.: Exponential integrators for time-fractional partial differential equations. *Eur. Phys. J. Spec. Top.* **222**(8), 1915–1927 (2013)
- [18] Garrappa, R.: Numerical evaluation of two and three parameter Mittag-Leffler functions. *SIAM J. Numer. Anal.* **53**(3), 1350–1369 (2015)
- [19] Garrappa, R., Mainardi, F., Maione, G.: Models of dielectric relaxation based on completely monotone functions. *Fract. Calc. Appl. Anal.* **19**(5), 1105–1160 (2016)
- [20] Garrappa, R., Moret, I., Papolizio, M.: Solving the time-fractional Schrödinger equation by Krylov projection methods. *J. Comput. Phys.* **293**, 115–134 (2015)
- [21] Garrappa, R., Moret, I., Papolizio, M.: On the time-fractional Schrödinger equation: theoretical analysis and numerical solution by matrix Mittag-Leffler functions. *Comput. Math. Appl.* **74**(5), 977–992 (2017)
- [22] Garrappa, R., Papolizio, M.: On the use of matrix functions for fractional partial differential equations. *Math. Comput. Simulation* **81**(5), 1045–1056 (2011)
- [23] Garrappa, R., Papolizio, M.: Evaluation of generalized Mittag-Leffler functions on the real line. *Adv. Comput. Math.* **39**(1), 205–225 (2013)
- [24] Giusti, A., Colombaro, I.: Prabhakar-like fractional viscoelasticity. *Commun. Nonlinear Sci. Numer. Simul.* **56**, 138–143 (2018)

- [25] Golub, G.H., Van Loan, C.F.: Matrix computations, third edn. Johns Hopkins Studies in the Mathematical Sciences. Johns Hopkins University Press, Baltimore, MD (1996)
- [26] Gorenflo, R., Kilbas, A.A., Mainardi, F., Rogosin, S.: Mittag-Leffler functions. Theory and Applications. Springer Monographs in Mathematics. Springer, Berlin (2014)
- [27] Gorenflo, R., Loutchko, J., Luchko, Y.: Computation of the Mittag-Leffler function $E_{\alpha,\beta}(z)$ and its derivative. *Fract. Calc. Appl. Anal.* **5**(4), 491–518 (2002)
- [28] Gorenflo, R., Mainardi, F.: Fractional calculus: integral and differential equations of fractional order. In: *Fractals and fractional calculus in continuum mechanics* (Udine, 1996), *CISM Courses and Lect.*, vol. 378, pp. 223–276. Springer, Vienna (1997)
- [29] Hale, N., Higham, N.J., Trefethen, L.N.: Computing $\mathbf{A}^\alpha$, $\log(\mathbf{A})$, and related matrix functions by contour integrals. *SIAM J. Numer. Anal.* **46**(5), 2505–2523 (2008)
- [30] Haubold, H.J., Mathai, A.M., Saxena, R.K.: Mittag-Leffler functions and their applications. *J. Appl. Math.* pp. Art. ID 298,628, 51 (2011)
- [31] Henrici, P.: Applied and computational complex analysis, vol. 1. John Wiley & Sons, New York-London-Sydney (1974)
- [32] Higham, N.J.: Accuracy and stability of numerical algorithms, second edn. Society for Industrial and Applied Mathematics (SIAM), Philadelphia, PA (2002)
- [33] Higham, N.J.: Functions of matrices. Society for Industrial and Applied Mathematics (SIAM), Philadelphia, PA (2008)
- [34] Higham, N.J., Al-Mohy, A.H.: Computing matrix functions. *Acta Numer.* **19**, 159–208 (2010)
- [35] Liemert, A., Sandev, T., Kantz, H.: Generalized Langevin equation with tempered memory kernel. *Phys. A* **466**, 356–369 (2017)
- [36] Lino, P., Maione, G.: Design and simulation of fractional-order controllers of injection in CNG engines. *IFAC Proceedings Volumes (IFAC-PapersOnline)*, 582–587 (2013)
- [37] Lino, P., Maione, G.: Fractional order control of the injection system in a CNG engine. 2013 European Control Conference, ECC 2013, 3997–4002 (2013)
- [38] Luchko, Y., Gorenflo, R.: An operational method for solving fractional differential equations with the Caputo derivatives. *Acta Math. Vietnam.* **24**(2), 207–233 (1999)
- [39] Mainardi, F., Mura, A., Pagnini, G.: The M -Wright function in time-fractional diffusion processes: a tutorial survey. *Int. J. Differ. Equ.* pp. Art. ID 104,505, 29 (2010)
- [40] Matignon, D., d’Andréa Novel, B.: Some results on controllability and observability of finite-dimensional fractional differential systems. In: *Computational Engineering in Systems Applications, Proceedings of the IMACS, IEEE SMC Conference, Lille, France*, pp. 952–956 (1996)
- [41] Matychyn, I., Onyshchenko, V.: Time-optimal control of fractional-order linear systems. *Fract. Calc. Appl. Anal.* **18**(3), 687–696 (2015)

- [42] Mittag-Leffler, M.G.: Sopra la funzione $E_\alpha(x)$. Rend. Accad. Lincei **13**(5), 3–5 (1904)
- [43] Mittag-Leffler, M.G.: Sur la représentation analytique d’une branche uniforme d’une fonction monogène - cinquième note. Acta Mathematica **29**(1), 101–181 (1905)
- [44] Moret, I., Novati, P.: On the convergence of Krylov subspace methods for matrix Mittag-Leffler functions. SIAM J. Numer. Anal. **49**(5), 2144–2164 (2011)
- [45] de Oliveira, D.S., Capelas de Oliveira, E., Deif, S.: On a sum with a three-parameter Mittag-Leffler function. Integral Transforms Spec. Funct. **27**(8), 639–652 (2016)
- [46] Popolizio, M.: Numerical solution of multiterm fractional differential equations using the matrix Mittag-Leffler functions. Mathematics **1**(6), 7 (2018)
- [47] Prabhakar, T.R.: A singular integral equation with a generalized Mittag-Leffler function in the kernel. Yokohama Math. J. **19**(1), 7–15 (1971)
- [48] Rodrigo, M.R.: On fractional matrix exponentials and their explicit calculation. J. Differential Equations **261**(7), 4223–4243 (2016)
- [49] Rogosin, S.: The role of the Mittag-Leffler function in fractional modeling. Mathematics **3**(2), 368–381 (2015)
- [50] Sandev, T.: Generalized Langevin equation and the Prabhakar derivative. Mathematics **5**(4), 66 (2017)
- [51] Stanislavsky, A., Weron, K.: Numerical scheme for calculating of the fractional two-power relaxation laws in time-domain of measurements. Comput. Phys. Commun. **183**(2), 320–323 (2012)
- [52] Tomovski, Ž., Pogány, T.K., Srivastava, H.M.: Laplace type integral expressions for a certain three-parameter family of generalized Mittag-Leffler functions with applications involving complete monotonicity. J. Franklin Inst. **351**(12), 5437–5454 (2014)
- [53] Trefethen, L.N., Weideman, J.A.C.: The exponentially convergent trapezoidal rule. SIAM Rev. **56**(3), 385–458 (2014)
- [54] Valério, D., Tenreiro Machado, J.: On the numerical computation of the Mittag-Leffler function. Commun. Nonlinear Sci. Numer. Simul. **19**(10), 3419–3424 (2014)
- [55] Weideman, J.A.C., Trefethen, L.N.: Parabolic and hyperbolic contours for computing the Bromwich integral. Math. Comp. **76**(259), 1341–1356 (2007)
- [56] Zeng, C., Chen, Y.: Global Padé approximations of the generalized Mittag-Leffler function and its inverse. Fract. Calc. Appl. Anal. **18**(6), 1492–1506 (2015)