

Numerical solution for an inverse source problem of a fractional order diffusion-wave equation

Maryam Ebrahimi*, Mashallah Matinfar, and Afshin Babaei

Department of Applied Mathematics, University of Mazandaran, Babolsar, Iran.

Abstract

In this paper, we consider an inverse source problem of a fractional order diffusion-wave equation (FDWE), in which the space-dependent source term is unknown. In order to obtain the numerical solution of the discussed problem and to find the unknown source function, a Chebyshev collocation method is proposed. Since this inverse problem is an ill-posed problem, a regularization scheme based on the mollification technique is used to find a stable problem. Subsequently, the stable problem is solved numerically by applying the collocation method. Furthermore, the convergence analysis is considered, and finally, the effectiveness of the studied algorithm is demonstrated by some test examples.

Keywords. Inverse problem, Space-dependent source function, Mollification technique, Chebyshev collocation method, Convergence analysis.

2010 Mathematics Subject Classification. 65M12, 65F22, 86A22, 97N10.

1. INTRODUCTION

New tools and techniques are needed to solve inverse source problems of FDWEs. Many authors have developed effective techniques to solve such problems [3, 5, 10, 13–16]. The FDWE is obtained by replacing the time derivative in the usual diffusion and wave equations by the derivative of fractional order. If the derivative order for the time variable is in the range of 0 and 1, we have a generalized form of the diffusion equation, and if it is in the range of 1 and 2, we have a generalized form of the wave equation.

In problems involving differential equations, we are dealing with two categories of problems, the direct and the inverse [4]. In direct problems, the boundary conditions, initial conditions, and parameters are known and the goal is to find the unknown solution to the problem. In an inverse problem, however, some of the boundary data, initial data, or parameters may not be known. In these cases, we are dealing with inverse problems that require additional conditions to solve the problem. Usually, these additional conditions are given on the basis of measured approximate data [8]. Some of these problems are very sensitive to changes in the input parameters, and a small noise in these data lead to significant changes in the solution of the problem [9].

Inverse problems are generally ill-posed due to the presence of high-frequency noise in the input data, so we need to use regularization methods. To address this issue, regularization techniques have been developed, aimed at improving stability and accuracy in the reconstruction process. Among these techniques, the Tikhonov regularization method increases the stability of the solution by adding a regularization term to the objective function. Typically, when solving the problem $Ax = b$, this method is written as:

$$\min_x \|Ax - b\|^2 + \lambda \|x\|^2,$$

where λ is the regularization parameter that establishes a balance between the fit of the data and the stability of the solution. The value of λ must be chosen correctly, as a small value does not reduce instability, and a large value leads to the loss of important details in the solution. In the Tikhonov method, the solution is therefore highly dependent on the value of λ , which makes its selection a challenge [11].

Received: 14 April 2024; Accepted: 25 June 2025.

* Corresponding author. Email: m.ebrahimi212@yahoo.com.

In contrast, the mollification method does not add an artificial penalty term, but solves the problem with a smooth approximation of the data. So, this method can provide more physically meaningful solutions for some applications. In addition, the mollification method can be computationally more efficient for certain problems because it uses a direct smoothing filter instead of solving a modified system of equations (like the Tikhonov method), which selectively removes real noise and preserves valuable information. This is particularly effective when the data contains systematic errors and outliers [2, 9].

In this work, we consider the inverse source problem related to the FDWE of the form

$$\mathcal{D}_t^\alpha u(x, t) = \frac{\partial^2 u}{\partial x^2}(x, t) + \mu_1(x) \frac{\partial u}{\partial x}(x, t) + \mu_2(x) u(x, t) + q(t) f(x), \quad (x, t) \in \Omega_L \times \Omega_T, \quad (1.1)$$

with the initial and boundary conditions

$$u(x, 0) = \xi_0(x), \quad x \in \Omega_L, \quad (1.2)$$

$$u_t(x, 0) = \xi_1(x), \quad x \in \Omega_L, \quad (1.3)$$

$$u(0, t) = \lambda_0(t), \quad t \in \Omega_T, \quad (1.4)$$

$$u(\mathcal{L}, t) = \lambda_1(t), \quad t \in \Omega_T, \quad (1.5)$$

in which the known coefficients $\mu_i(x)$, and the conditions $\xi_i(x)$, $\lambda_i(t)$, $i = 0, 1$, are assumed to be continuous functions. Also, $\Omega_L = [0, \mathcal{L}]$, $\Omega_T = [0, \mathcal{T}]$, and L, T are positive constants. The function $q(t)f(x)$ in Eq. (1.1) is the source term with the unknown factor $f(x)$. Moreover, the operator $\mathcal{D}_t^\alpha$ denotes the fractional derivative in the Caputo formation, which is defined as [6]:

$$\mathcal{D}_t^\alpha u(\cdot, t) = \frac{1}{(1-\alpha)!} \int_0^t (t-\varsigma)^{1-\alpha} \frac{\partial^2 u}{\partial \tau^2}(\cdot, \varsigma) d\varsigma, \quad \alpha \in (1, 2), \quad (1.6)$$

where $\alpha! := \Gamma(\alpha + 1)$ and $\Gamma(\cdot)$ shows the Gamma function.

To find the unknown functions u and f in Eq. (1.1), we need an additional condition. Thus, for $\hat{t} \in (0, \mathcal{T})$, suppose

$$u(x, \hat{t}) = \kappa(x), \quad x \in \Omega_L. \quad (1.7)$$

For this problem, the proof of uniqueness and stability of the solution is investigated in [3].

In this paper, we first present a regularization method based on the mollification technique to obtain a regularized problem. Then, this stabilized problem is solved using a collocation scheme based on the sixth-kind Chebyshev polynomials (SKCPs).

2. MOLLIFICATION TECHNIQUE

In practice, we only have perturbed approximations for the input function $\kappa(x)$, and its exact value is not available. Therefore, this perturbed function may affect the solution of the problem. In this paper, the mollification regularization technique is used to smooth the disordered data. Moreover, this method is suitable for achieving consistency and stability in many types of ill-posed problems [2, 12].

Suppose $B_s = \left(\int_{-s}^s \exp(-z^2) dz \right)^{-1}$, such that $v, s > 0$ and $sv < 1/2$. The v -mollification of an integrable function is based on a convolution with the Gaussian kernel

$$f_{v,s}(t) = \begin{cases} B_s v^{-1} \exp(-\frac{x^2}{v^2}), & |x| \leq sv, \\ 0, & |x| > sv. \end{cases}$$

The v -mollifier $f_{v,s} \in C^\infty(-sv, sv)$ is a non-negative function that vanishes outside of $(-sv, sv)$ and satisfies the following condition

$$\int_{-sv}^{sv} f_{v,s}(x) dx = 1. \quad (2.1)$$



Assuming that $S := \{\hat{x}_j : j \in \mathbb{Z}\} \subset [0, 1]$ and $\Delta x := \sup_{j \in \mathbb{Z}} (\hat{x}_{j+1} - \hat{x}_j)$, satisfies $\hat{x}_{j+1} - \hat{x}_j > b > 0$ where b is a positive constant. Let $\mathcal{W} := \{\kappa(\hat{x}_j) = \kappa_j : j \in \mathbb{Z}\}$ be a discrete function defined on S . We set $h_j = \frac{1}{2}(\hat{x}_j + \hat{x}_{j+1})$, $j \in \mathbb{Z}$. The discrete v -mollification of $\mathcal{W}$ is now defined as:

$$\mathcal{J}_v \mathcal{W}(x) = \sum_{j=-\infty}^{\infty} \left(\int_{h_{j-1}}^{h_j} f_v(x-h) dh \right) \kappa_j.$$

Considering (2.1), we obtain

$$\sum_{j=-\infty}^{\infty} \left(\int_{h_{j-1}}^{h_j} f_v(x-h) dh \right) = \int_{-sv}^{sv} f_v(h) dh = 1.$$

Assume that instead of the function κ , there exists a noisy function $\kappa^\varepsilon \in C^0([0, 1])$, such that $\|\kappa - \kappa^\varepsilon\|_{\infty, [0, 1]} \leq \varepsilon$. The generalized cross validation (GCV) criterion automatically determines the smoothing parameter v [12]. To compute $\mathcal{J}_v \mathcal{W}^\varepsilon$ over the range $I = [0, 1]$, it is necessary to manage the data near the boundaries. The technique presented in the chapter 4 of [12] consists of extending κ^ε to a larger interval $I^v = [-sv, 1 + sv]$. An extension κ^* of κ^ε to $[-sv, 0]$ and $[1, 1 + sv]$ is thus performed under the conditions that $\|\mathcal{J}_v \kappa^* - \kappa^\varepsilon\|_{L^2[0, sv]}$ and $\|\mathcal{J}_v \kappa^* - \kappa^\varepsilon\|_{L^2[1-sv, 1]}$ are minimal. Then this optimization problem at $x = 1$ will have a unique solution as follows

$$\kappa^* = \frac{\int_{1-sv}^1 \left(\kappa^\varepsilon(x) - \int_0^1 f_v(x-h) \kappa(h) dh \right) \left(\int_1^{1+sv} f_v(x-h) dh \right) dx}{\int_{1-sv}^1 \left(\int_1^{1+sv} f_v(x-h) dh \right) dx}.$$

In the continuation, the set of points is defined as:

$$\hat{x}_i := ir, \quad 0 \leq i \leq \mathfrak{M}, \quad (2.2)$$

where $\mathfrak{M}$ is a positive constant and $r = \frac{1}{\mathfrak{M}}$. If the described method of discrete mollification is then applied to $\{\kappa^\varepsilon(\hat{x}_i) : 0 \leq i \leq \mathfrak{M}\}$, the result is a discrete mollified function $\{\mathcal{J}_v \mathcal{W}^\varepsilon(\hat{x}_i) : 0 \leq i \leq \mathfrak{M}\}$. Hence, a stable approximation of $\kappa(x)$ is obtained by interpolating $\mathcal{J}_v \mathcal{W}^\varepsilon(\hat{x}_i)$ from the additional noise function κ^ε .

3. COLLOCATION PROCEDURE

In this section, we will review the main properties of SKCP and some definitions. We present the SKCP collocation scheme to numerically solve the introduced inverse source problem.

3.1. Shifted SKCPs.

Definition 3.1. In the interval Ω_L , the explicit form for the shifted SKCP is defined as [1]

$$T_n(x) = \sum_{k=0}^n \bar{\vartheta}_{k,n}(x/\mathcal{L})^k, \quad (3.1)$$

where

$$\bar{\vartheta}_{k,n} = \begin{cases} \frac{2^{2k-n}}{(2k+1)!} \sum_{j=\lfloor \frac{k+1}{2} \rfloor}^{\frac{n}{2}} \frac{(-1)^{\frac{n}{2}+j+k} (2j+k+1)!}{(2j-k)!}, & n \text{ even}, \\ \frac{2^{2k-n+1}}{(n+1)(2k+1)!} \sum_{j=\lfloor \frac{k}{2} \rfloor}^{\frac{n-1}{2}} \frac{(-1)^{\frac{n+1}{2}+j+k} (j+1)(2j+k+2)!}{(2j-k+1)!}, & n \text{ odd}. \end{cases}$$

Theorem 3.2. ([1]) Assume $\mathcal{S}(t) \in L_w^2(\Omega_L)$ and $|\mathcal{S}(t)^{(3)}| \leq \alpha$ for a positive constant α , such that:

$$\mathcal{S}(t) = \sum_{j=0}^{\infty} g_j T_j(t). \quad (3.2)$$



Then, this series is uniformly convergent to $S(t)$ and we have $|g_j| < \frac{\alpha}{2j^3}$ for all $j > 3$. If

$$\mathcal{S}_{\mathfrak{N}}(t) \simeq \sum_{j=0}^{\mathfrak{N}} g_j T_j(t),$$

is an estimation of $\mathcal{S}(t)$, then

$$|\mathcal{S}(t) - \mathcal{S}_{\mathfrak{N}}(t)| < \frac{\alpha}{2^{\mathfrak{N}}}. \quad (3.3)$$

Let us assume that $\Theta := [0, \mathcal{L}] \times [0, \mathcal{T}]$ and $L_w^2(\Theta)$ is the space of square-integrable functions with the weight function

$$w(x, t) = \sqrt{\frac{1}{L}(x - \frac{x^2}{L})(2\frac{x}{L} - 1)^2} \sqrt{\frac{1}{T}(t - \frac{t^2}{T})(2\frac{t}{T} - 1)^2}.$$

Theorem 3.3. ([7]) Assume $\mathcal{S}(x, t) \in L_w^2(\Theta)$ and $\left\| \frac{\partial^6 \mathcal{S}(x, t)}{\partial x^3 \partial t^3} \right\|_2 \leq \hat{g}$, where $\hat{g}$ is a positive constant, such that:

$$\mathcal{S}(x, t) = \sum_{i=0}^{\infty} \sum_{j=0}^{\infty} g_{i,j} T_i(x) T_j(t). \quad (3.4)$$

So, we have $|g_{i,j}| < \frac{\hat{g}}{i^3 j^3}$, for $i, j > 3$. Also, if an estimate of S is

$$\tilde{\mathcal{S}}(x, t) = \sum_{i=0}^{\mathfrak{N}} \sum_{j=0}^{\mathfrak{M}} g_{i,j} T_i(x) T_j(t), \quad (3.5)$$

where $\mathfrak{N}$ and $\mathfrak{M}$ are positive constants, then

$$|S - \tilde{S}| < \frac{\hat{g}}{2^{\mathfrak{N}+\mathfrak{M}}}.$$

Lemma 3.4. ([7]) According to the assumptions of Theorem 3.3 for S and $\tilde{S}$, we have

$$\begin{aligned} (i) \quad & \left| \frac{\partial \mathcal{S}(x, t)}{\partial x} - \frac{\partial \tilde{\mathcal{S}}}{\partial x} \right| < \eta_1 \frac{\mathfrak{N}}{2^{\mathfrak{N}+\mathfrak{M}-2}}, \\ (ii) \quad & \left| \frac{\partial^2 \mathcal{S}(x, t)}{\partial x^2} - \frac{\partial^2 \tilde{\mathcal{S}}}{\partial x^2} \right| < \eta_2 \frac{\mathfrak{N}^3}{2^{\mathfrak{N}+\mathfrak{M}-8}}, \\ (iii) \quad & \left| \frac{\partial^2 \mathcal{S}(x, t)}{\partial t^2} - \frac{\partial^2 \tilde{\mathcal{S}}}{\partial t^2} \right| < \eta_3 \frac{\mathfrak{M}^3}{2^{\mathfrak{N}+\mathfrak{M}-8}}, \end{aligned}$$

where η_i , $i = 1, 2, 3$, are positive constants.

3.2. Numerical Method. In this subsection, we will present the numerical method to approximate the solution of (1.1)–(1.4) and (1.7), based on SKCPs. To obtain this solution, we let

$$u(x, t) \simeq \mathbf{U}_{\mathfrak{N}, \mathfrak{M}}(x, t) = \sum_{i=0}^{\mathfrak{N}} \sum_{j=0}^{\mathfrak{M}} g_{i,j} T_i(x) T_j(t) = \Psi(x)^T \mathbf{G} \bar{\Psi}(t), \quad (3.6)$$

and

$$f(x) \simeq \tilde{f}_{\mathfrak{N}}(x) = \sum_{i=0}^{\mathfrak{N}} d_i T_i(x) = \Psi(x)^T \mathbf{D}, \quad (3.7)$$

where

$$\Psi(x) = [T_0(x), T_1(x), \dots, T_{\mathfrak{N}}(x)]^T, \quad (3.8)$$

$$\bar{\Psi}(t) = [\bar{T}_0(t), \bar{T}_1(t), \dots, \bar{T}_{\mathfrak{M}}(t)]^T, \quad (3.9)$$



such that $T_i(x)$ is the shifted SKCPs in the interval Ω_L , $\bar{T}_j(t)$ is the shifted SKCPs in the interval Ω_T . In addition, $\mathfrak{M}$ and $\mathfrak{N}$ are two positive constants. Also

$$\mathbf{G} = \begin{pmatrix} g_{0,0} & \cdots & g_{0,\mathfrak{M}} \\ \vdots & \ddots & \vdots \\ g_{\mathfrak{N},0} & \cdots & g_{\mathfrak{N},\mathfrak{M}} \end{pmatrix}_{(\mathfrak{N}+1) \times (\mathfrak{M}+1)},$$

represents the unknown coefficient matrix, and $\mathbf{D} = [d_0, d_1, \dots, d_{\mathfrak{N}}]^T$ denotes the vector of unknown coefficients to be determined.

Theorem 3.5. *If $\bar{\Psi}(t)$ is defined in (3.9), $\mathcal{D}_t^\alpha \bar{\Psi}(t) = \Lambda^\alpha(t)$, where $(1 < \alpha < 2)$, then*

$$\Lambda^\alpha(t) = \left[0, 0, \sum_{r=2}^2 \xi_{r,2}^\alpha(t), \dots, \sum_{r=2}^j \xi_{r,j}^\alpha(t), \dots, \sum_{r=2}^{\mathfrak{M}} \xi_{r,\mathfrak{M}}^\alpha(t) \right]^T, \quad (3.10)$$

that

$$\xi_{r,j}^\alpha(t) = \frac{\Gamma(r+1)}{\mathcal{T}^r \Gamma(r+1-\alpha)} \bar{\vartheta}_{r,j} t^{r-\alpha}.$$

Proof. According to (1.6) and (3.1), we have

$$\mathcal{D}_t^\alpha \bar{T}_0(t) = \mathcal{D}_t^\alpha \bar{T}_1(t) = 0.$$

For $r \geq 2$, we have

$$\mathcal{D}_t^\alpha t^r = \frac{\Gamma(r+1)}{\Gamma(r+1-\alpha)} t^{r-\alpha}. \quad (3.11)$$

Thus, for $j = 1, \dots, \mathfrak{M}$,

$$\mathcal{D}_t^\alpha \bar{\Psi}_j(t) = \sum_{r=0}^j \bar{\vartheta}_{r,j} \mathcal{D}_t^\alpha (t/\mathcal{T})^r = \sum_{r=2}^j \frac{\Gamma(r+1)}{\mathcal{T}^r \Gamma(r+1-\alpha)} \bar{\vartheta}_{r,j} t^{r-\alpha}. \quad (3.12)$$

From (1.1), (3.6), and (3.7),

$$\Psi(x)^T \mathbf{C} \Lambda^\alpha(t) = \Psi_{xx}(x)^T \mathbf{G} \bar{\Psi}(t) + \mu_1(x) \Psi_x(x)^T \mathbf{G} \bar{\Psi}(t) + \mu_2(x) \Psi(x)^T \mathbf{G} \bar{\Psi}(t) + q(t) \Psi(x)^T \mathbf{D}. \quad (3.13)$$

In addition, from the conditions (1.2)–(1.5) and Eq. (3.6), we have

$$\Pi_1(x) = \Psi(x)^T \mathbf{G} \bar{\Psi}(0) - u_0(x), \quad (3.14)$$

$$\Pi_2(x) = \Psi(x)^T \mathbf{G} \bar{\Psi}_t(0) - u_{t_0}(x), \quad (3.15)$$

$$\Omega_0(t) = \Psi(0)^T \mathbf{G} \bar{\Psi}(t) - \lambda_0(t), \quad (3.16)$$

$$\Omega_1(t) = \Psi(\mathcal{L})^T \mathbf{G} \bar{\Psi}(t) - \lambda_1(t), \quad (3.17)$$

$$\Xi(x) = \Psi(x)^T \mathbf{G} \Lambda^\alpha(\hat{t}) - \Psi_{xx}(x)^T \mathbf{G} \bar{\Psi}(\hat{t}) - \mu_1(x) \Psi_x(x)^T \mathbf{G} \bar{\Psi}(\hat{t}) - \mu_2(x) \mathcal{J}_v \phi^\varepsilon(x) - q(\hat{t}) \Psi(x)^T \mathbf{D}, \quad (3.18)$$

where

$$\bar{\Psi}_t(t) = \left[\frac{d}{dt} \bar{T}_0(t), \dots, \frac{d}{dt} \bar{T}_{\mathfrak{M}}(t) \right]^T,$$

$$\Psi_x(x) = \left[\frac{d}{dx} T_0(x), \dots, \frac{d}{dx} T_{\mathfrak{N}}(x) \right]^T,$$

and

$$\Psi_{xx}(x) = \left[\frac{d^2}{dx^2} T_0(x), \dots, \frac{d^2}{dx^2} T_{\mathfrak{N}}(x) \right]^T.$$



Assume that $x_0 = 0$, $x_{\mathfrak{N}} = \mathcal{L}$, $x_1, \dots, x_{\mathfrak{N}-1}$ are the roots of $T_{\mathfrak{N}-1}(x)$, and $t_1, \dots, t_{\mathfrak{M}-1}$ are the roots of $T_{\mathfrak{M}-1}(t)$. Now, by appraising (3.13) in $(\mathfrak{M} - 1) \times (\mathfrak{N} - 1)$ collocation points (x_i, t_j) for $1 \leq i \leq \mathfrak{N} - 1$ and $1 \leq j \leq \mathfrak{M} - 1$, we obtain

$$\begin{aligned} \mathcal{R}(x_i, t_j) &= \Psi(x_i)^T \mathbf{G} \Lambda^\alpha(t_j) - \Psi_{xx}(x_i)^T \mathbf{G} \bar{\Psi}(t_j) \\ &\quad - \mu_1(x_i) \Psi_x(x_i)^T \mathbf{G} \bar{\Psi}(t_j) - \mu_2(x_i) \Psi(x_i)^T \mathbf{G} \bar{\Psi}(t_j) - q(t_j) \Psi(x_i)^T \mathbf{D}. \end{aligned} \quad (3.19)$$

Hence, by evaluating (3.15)–(3.19) at the collocation points and using Eq. (3.13), a system of linear equations of order $(\mathfrak{N} + 1) \times (\mathfrak{M} + 1)$ is obtained as:

$$\begin{cases} \mathcal{R}(x_i, t_j) = 0, & 1 \leq i \leq \mathfrak{N} - 1, 1 \leq j \leq \mathfrak{M}, \\ \Pi_r(x_i) = 0, & 0 \leq i \leq \mathfrak{N}, r = 1, 2, \\ \Omega_z(t_j) = 0, & 1 \leq j \leq \mathfrak{M}, z = 0, 1, \\ \Xi(x_i) = 0, & 0 \leq i \leq \mathfrak{N}, \end{cases} \quad (3.20)$$

in which the unknown coefficients $g_{i,j}$ and d_i , $0 \leq i \leq \mathfrak{N}, 0 \leq j \leq \mathfrak{M}$, should be determined. $\square$

4. CONVERGENCE ANALYSIS

In this section, the convergence of the numerical solution obtained in section 3 is examined.

Theorem 4.1. Suppose that for exact solution u of (1.1), we have the numerical approximation $\mathbf{U}_{\mathfrak{N}, \mathfrak{M}}$, which is determined according to the procedure presented in subsection 3.2. Moreover, the functions $\mu_i(x)$, $i = 1, 2$ and $q(t)$ are bounded and $\mathcal{R}_{\mathfrak{N}, \mathfrak{M}}(x, t)$ is the residual error u . Then, $\|\mathcal{R}_{\mathfrak{N}, \mathfrak{M}}\|_\infty \rightarrow 0$, when $\mathfrak{N}, \mathfrak{M} \rightarrow \infty$.

Proof. From (1.1) and $\mathbf{U}_{\mathfrak{N}, \mathfrak{M}}(x, t)$, we have

$$\mathcal{D}_t^\alpha \mathbf{U}_{\mathfrak{N}, \mathfrak{M}}(x, t) = \frac{\partial^2 \mathbf{U}_{\mathfrak{N}, \mathfrak{M}}}{\partial x^2}(x, t) + \mu_1 \frac{\partial \mathbf{U}_{\mathfrak{N}, \mathfrak{M}}}{\partial x}(x, t) + \mu_2 \mathbf{U}_{\mathfrak{N}, \mathfrak{M}}(x, t) + \tilde{f}_{\mathfrak{N}}(x) q(t) + \mathcal{R}_{\mathfrak{N}, \mathfrak{M}}(x, t). \quad (4.1)$$

First, we consider $\|u(x, t)\|_\infty = \|u\|_\infty = \sup_{(x,t) \in \mathcal{L} \times \mathcal{I}} |u(x, t)|$ and from Eqs. (1.1) and (4.1), the following inequality can be seen:

$$\begin{aligned} \|\mathcal{R}_{\mathfrak{N}, \mathfrak{M}}\|_\infty &\leq \left\| \mathcal{D}_{0,t}^\alpha \left(u - \mathbf{U}_{\mathfrak{N}, \mathfrak{M}} \right) \right\|_\infty + \left\| u_{xx} - \frac{\partial^2 \mathbf{U}_{\mathfrak{N}, \mathfrak{M}}}{\partial x^2} \right\|_\infty \\ &\quad + \left\| \mu_1 \right\|_\infty \left\| u_x - \frac{\partial \mathbf{U}_{\mathfrak{N}, \mathfrak{M}}}{\partial x} \right\|_\infty + \left\| \mu_2 \right\|_\infty \left\| u - \mathbf{U}_{\mathfrak{N}, \mathfrak{M}} \right\|_\infty + \left\| q \right\|_\infty \left\| f - \tilde{f}_{\mathfrak{N}} \right\|_\infty. \end{aligned}$$

By using Lemma 3.4, results

$$\left\| u_{tt} - \frac{\partial^2 \mathbf{U}_{\mathfrak{N}, \mathfrak{M}}}{\partial t^2} \right\|_\infty < \frac{\eta_1 \mathfrak{M}^3}{2^{\mathfrak{N} + \mathfrak{M} - 8}}, \quad (4.2)$$

where η_1 is a positive integer, thus

$$\begin{aligned} \left\| \mathcal{D}_{0,t}^\alpha \left(u - \mathbf{U}_{\mathfrak{N}, \mathfrak{M}} \right) \right\|_\infty &\leq \left(\int_0^t \frac{\|(t-v)^{1-\alpha}\|_\infty}{\Gamma(2-\alpha)} \left\| u_{vv} - \frac{\partial^2 \mathbf{U}_{\mathfrak{N}, \mathfrak{M}}}{\partial^2 v} \right\|_\infty dv \right) \\ &< \frac{\eta_1 \mathfrak{M}^3}{\Gamma(2-\alpha) 2^{\mathfrak{N} + \mathfrak{M} - 8}} \left(\int_0^t \|(t-v)^{1-\alpha}\|_\infty dv \right). \end{aligned}$$

Since $0 < v < t \leq \mathcal{T}$, we get

$$\left\| \mathcal{D}_t^\alpha \left(u - \mathbf{U}_{\mathfrak{N}, \mathfrak{M}} \right) \right\|_\infty < \frac{\eta_1 \mathcal{T}^{2-\alpha} \mathfrak{M}^3}{\Gamma(2-\alpha) 2^{\mathfrak{N} + \mathfrak{M} - 8}}. \quad (4.3)$$



Assume $\|\mu_1(x)\|_\infty < \vartheta_1$, $\|\mu_2(x)\|_\infty < \vartheta_2$ and $\|q(t)\|_\infty < \vartheta_3$. Moreover, by Theorem 3.3,

$$\left\| u_{xx} - \frac{\partial^2 \mathbf{U}_{\mathfrak{N}, \mathfrak{M}}}{\partial x^2} \right\|_\infty < \frac{\eta_2 \mathfrak{N}^3}{2^{\mathfrak{N}+\mathfrak{M}-8}}, \quad (4.4)$$

$$\|\mu_1\|_\infty \left\| u_x - \frac{\partial \mathbf{U}_{\mathfrak{N}, \mathfrak{M}}}{\partial x} \right\|_\infty < \frac{\vartheta_1 \eta_3 \mathfrak{N}}{2^{\mathfrak{N}+\mathfrak{M}-2}}, \quad (4.5)$$

$$\|\mu_2\|_\infty \|u - \mathbf{U}_{\mathfrak{N}, \mathfrak{M}}\|_\infty < \frac{\vartheta_2 \eta_4 \mathfrak{N}}{2^{\mathfrak{N}+\mathfrak{M}}}, \quad (4.6)$$

$$\|q(f - \tilde{f}_{\mathfrak{N}})\|_\infty < \frac{\vartheta_3 \eta_5}{2^{\mathfrak{N}}}, \quad (4.7)$$

where $\eta_2, \eta_3, \eta_4, \eta_5, \vartheta_1, \vartheta_2$, and ϑ_3 are positive constants. Then, from Eqs. (4.2)–(4.7), the following results are obtained

$$\begin{aligned} \|\mathcal{R}_{\mathfrak{N}, \mathfrak{M}}\|_\infty &< \frac{\eta_1 \mathcal{T}^{2-\alpha} \mathfrak{M}^3}{\Gamma(2-\alpha) 2^{\mathfrak{N}+\mathfrak{M}-8}} + \frac{\eta_2 \mathfrak{N}^3}{2^{\mathfrak{N}+\mathfrak{M}-8}} + \frac{\vartheta_1 \eta_3 \mathfrak{N}}{2^{\mathfrak{N}+\mathfrak{M}-2}} + \frac{\vartheta_2 \eta_4}{2^{\mathfrak{N}+\mathfrak{M}}} + \frac{\vartheta_3 \eta_5}{2^{\mathfrak{N}}} \\ &< \frac{\eta_1 \mathcal{T}^{2-\alpha} \mathfrak{M}^3}{\Gamma(2-\alpha) 2^{\mathfrak{N}+\mathfrak{M}-8}} + \frac{\eta_2 \mathfrak{N}^3}{2^{\mathfrak{N}+\mathfrak{M}-8}} + \frac{\vartheta_1 \eta_3 \mathfrak{N}^3}{2^{\mathfrak{N}+\mathfrak{M}-8}} + \frac{\vartheta_2 \eta_4}{2^{\mathfrak{N}+\mathfrak{M}-8}} + \frac{\vartheta_3 \eta_5}{2^{\mathfrak{N}}} \\ &< \hat{\eta} \frac{\mathfrak{M} + 2\mathfrak{N}^3 + 1}{2^{\mathfrak{N}+\mathfrak{M}-8}} + \frac{1}{2^{\mathfrak{N}}}, \end{aligned} \quad (4.8)$$

where

$$\hat{\eta} = \max\left\{\frac{\eta_1 \mathcal{T}^{2-\alpha}}{\Gamma(2-\alpha)}, \eta_2, \vartheta_1 \eta_3, \vartheta_2 \eta_4, \vartheta_3 \eta_5\right\}.$$

So, from Eq. (4.8), we can see when $\mathfrak{N}, \mathfrak{M} \rightarrow \infty$, then $\|\mathcal{R}_{\mathfrak{N}, \mathfrak{M}}\|_\infty \rightarrow 0$. $\square$

5. TEST EXAMPLES

Herein, the correctness and applicability of the proposed method is investigated using some examples. Here, the L_2 -norm error is considered to evaluate the accuracy of the method:

$$\|\epsilon_{\mathfrak{M}}\|_2 = \left(\sum_{i=1}^{\mathfrak{N}} \sum_{j=1}^{\mathfrak{M}} (u(x_i, t_j) - \tilde{u}(x_i, t_j))^2 \right)^{\frac{1}{2}}, \quad (5.1)$$

such that $x_i, t_j, i = 1, \dots, \mathfrak{N}, j = 1, \dots, \mathfrak{M}$, are the list of collocation points. According to the space variable, the convergence order is

$$CO = \log_{\frac{\mathfrak{M}_2}{\mathfrak{M}_1}} \frac{\|\epsilon_{\mathfrak{M}_1}\|_2}{\|\epsilon_{\mathfrak{M}_2}\|_2}. \quad (5.2)$$

We assume that ε is the maximum amount of noise in the input data. By adding some random errors, we simulate the exact data function for the inverse problem. For this purpose,

$$\varrho^\varepsilon(\hat{x}_i) = \varrho(\hat{x}_i)(1 + \varepsilon \times \text{rand}(i)),$$

will be considered as a discrete noisy version of $\varrho(t)$, where $\hat{t}_i, i = 0, \dots, \hat{\mathfrak{M}}$ are defined in (2.2) and the uniformly distributed random numbers in $[-1, 1]$ are defined by the command $\text{rand}(i)$ in Matlab.

Example 5.1. Suppose that we have the equation

$$\mathcal{D}_t^\alpha u(x, t) = u_{xx}(x, t) + x u_x(x, t) + u(x, t) + f(x)q(t), \quad (x, t) \in [0, 1] \times [0, 1],$$



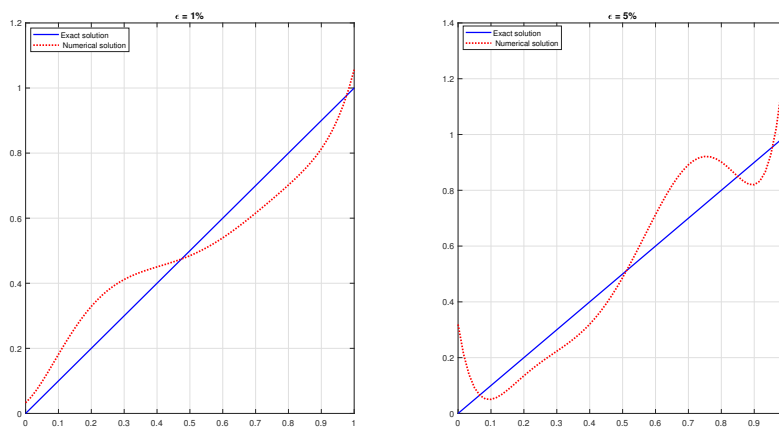


FIGURE 1. The exact and computed approximations of $f(x)$ without regularization, with $\varepsilon = 1\%$ and $\varepsilon = 5\%$ for Example 5.1.

with $\xi_0(x) = 0$, $\xi_1(x) = \pi x$, $\lambda_0(t) = 0$, $\lambda_1(t) = \sin(\pi t)$ and

$$q(x) = -\frac{\pi^3 t^{3-\alpha} \text{HypergeometricPFQ}[\{1\}, \{2 - \frac{\alpha}{2}, \frac{5}{2} - \frac{\alpha}{2}\}, -(\frac{1}{4})\pi^2 t^2]}{(6 - 5\alpha + \alpha^2)\Gamma(2 - \alpha)} - 2\sin(\pi t).$$

The exact solution to this problem is $u(x, t) = x \sin(\pi t)$ and $f(x) = x$.

The exact and computed approximations of $f(x)$ without regularization when $\alpha = 1.5$, $\hat{t} = 1$, $\mathfrak{N} = 14$, $\mathfrak{M} = 10$, and $\mathfrak{M} = 300$ are shown in Figure 1. In Figure 2, the exact solution and the numerical solution using the Tikhonov method and the mollification approach are shown. As can be seen, the mollification method performed better in regularizing the data and provided a more accurate solution than the solution regularized by the Tikhonov method. The mollification method uses a mollifier function to modify the input data. The smoothed data is then used to solve the problem.

The numerical approximation with regularization and the absolute error (AE) when $\alpha = 1.5$, $\hat{t} = 1$, $\mathfrak{N} = 14$, $\mathfrak{M} = 10$, and $\mathfrak{M} = 300$ are shown in Figure 3. The AE for u when $\alpha = 1.5$, $\hat{t} = 1$, $\mathfrak{N} = 14$, $\mathfrak{M} = 10$, and $\mathfrak{M} = 300$ can be seen in Figure 5. The L_2 -norm error, CO , and CPU-times of the numerical solution, for different values of $\mathfrak{M}$ at $\hat{t} = 1$ and $N = 14$, are displayed in Table 1. As the mesh size r is refined in the mollification approach, the numerical error is reduced despite the noise. Here, the theoretical results and the obtained findings are in a good agreement.

Example 5.2. To explain the impact of choosing a suitable regularization approach on the search for a stable solution in more detail. In this example, we examine the problem (1.1)–(1.5) with $\alpha = 1.6$, $\xi_0(x) = \xi_1(x) = \lambda_0(t) = \lambda_1(t) = 0$, $\mu_1(x) = 0$, $\mu_2(x) = 1$. and

$$q(t) = \left(-t^2 + \Pi^2 t^2 - \frac{2t^{2-\alpha}}{(-2 + \alpha)\Gamma(2 - \alpha)} \right),$$

for $(x, t) \in [0, 1] \times [0, 1]$. This problem has the exact solution $u(x, t) = \sin(\pi x) t^2$ and $f(x) = \sin(\pi x)$.

Consider two sample input functions that are defined as follows: $\kappa_1(x) = \sin(\pi x)$ and $\kappa_2(x) = \kappa_1(x) + 0.1$. The algorithm proposed in Section 3.2 is used to calculate the source term $f_1(x)$ and $f_2(x)$, that depend on $\kappa_1(x)$ and $\kappa_2(x)$. Figure 6 illustrates the input functions $\kappa_i(x)$ for $i = 1, 2$ and the approximated source terms $f_i(x)$ for $i = 1, 2$. Despite the small difference between the input functions, the computed source terms show considerable deviations from the expectations. This result highlights the ill-posedness of the inverse problem. Therefore, the application of appropriate regularization techniques is essential to obtain a stable numerical solution.



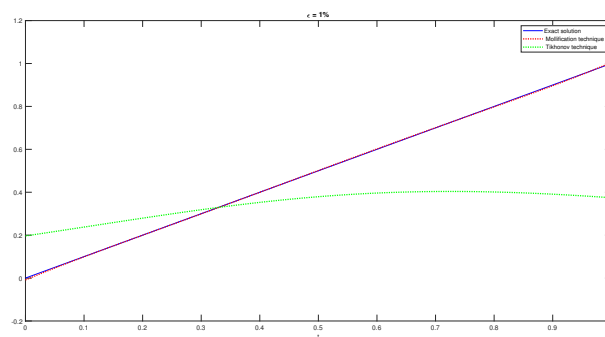


FIGURE 2. The exact and numerical values of $f(x)$, using the Tikhonov method and the mollification method when $\varepsilon = 1\%$ for Example 5.1.

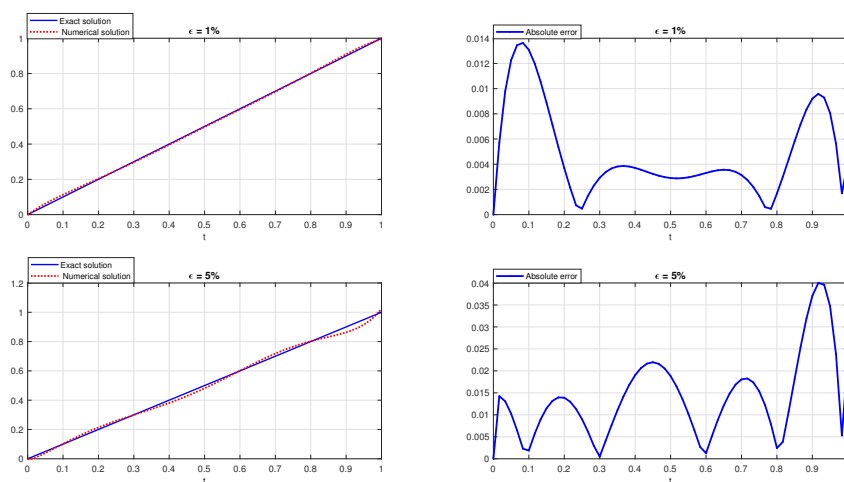


FIGURE 3. The exact and computed approximations with regularization (left) and the AE (right) for $\varepsilon = 1\%$ and $\varepsilon = 5\%$ in Example 5.1.

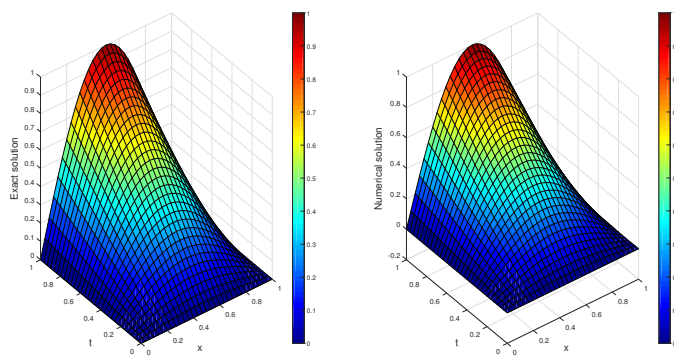


FIGURE 4. The exact and numerical solutions for Example 5.1, with $\varepsilon = 5\%$.

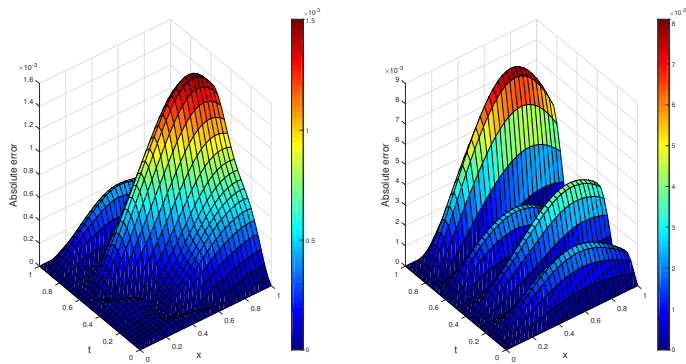


FIGURE 5. The AE of u for Example 5.1, with $\varepsilon = 1\%$ (left) and $\varepsilon = 5\%$ (right) .

TABLE 1. The L_2 -norm error, CO and CPU-time (Sec) for Example 5.1.

$r = 0.003$				$\varepsilon = 5\%$		
$\mathfrak{M}$	$\ \epsilon_{\mathfrak{M}}\ _2$	CO	CPU-time	$\ \epsilon_{\mathfrak{M}}\ _2$	CO	CPU-time
6	3.3444×10^{-4}	—	3.766	4.6758×10^{-4}	—	3.844
8	1.9553×10^{-4}	1.86582	5.406	3.3324×10^{-4}	1.17636	5.47
10	4.6795×10^{-5}	6.40806	7.687	8.02299×10^{-5}	6.38143	7.687
$r = 0.002$				$\varepsilon = 5\%$		
$\mathfrak{M}$	$\ \epsilon_{\mathfrak{M}}\ _2$	CO	CPU-time	$\ \epsilon_{\mathfrak{M}}\ _2$	CO	CPU-time
6	3.2435×10^{-4}	—	3.781	3.28023×10^{-4}	—	3.765
8	1.38681×10^{-4}	2.7063	5.454	1.48369×10^{-4}	2.7578	5.312
10	2.0226×10^{-5}	8.6276	7.794	1.57612×10^{-5}	9.7723	7.469

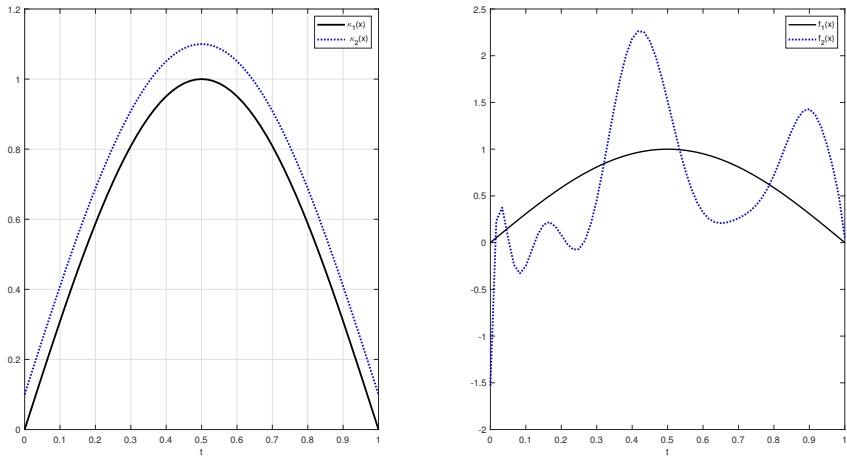


FIGURE 6. The two estimated source terms, $f_1(x)$ and $f_2(x)$ (Right) correspond to the two close input functions $\kappa_1(x)$ and $\kappa_2(x)$ (Left) in Example 5.2.



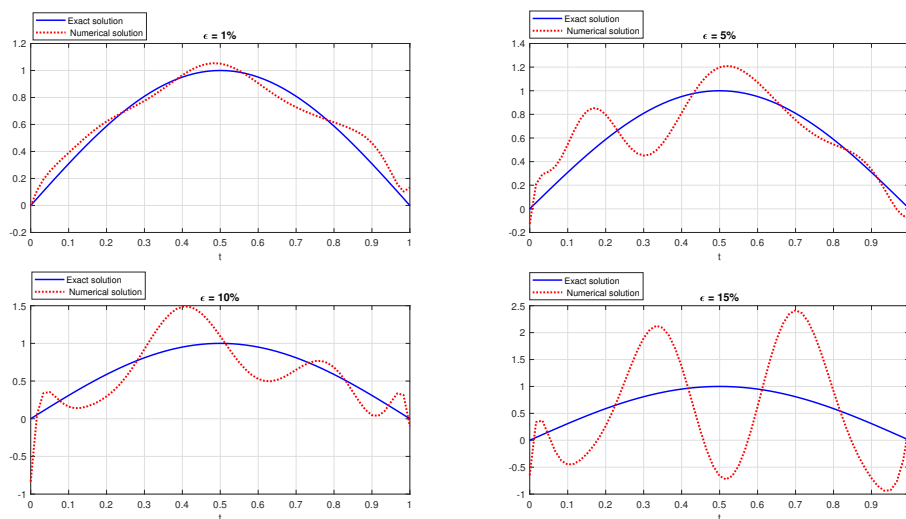


FIGURE 7. The exact function $f(x)$ and its computed approximations without regularization for different noise levels in Example 5.2.

TABLE 2. The L_2 -norm error, CO and CPU-time (Sec) for Example 5.2.

$r = 0.003$		$\varepsilon = 1\%$			$\varepsilon = 5\%$		
M	$\ \epsilon_{\mathfrak{M}}\ _2$	CO	CPU-time		$\ \epsilon_{\mathfrak{M}}\ _2$	CO	CPU-time
2	1.2469×10^{-4}	—	0.767		1.10123×10^{-4}	—	0.782
4	9.3421×10^{-5}	0.4165	1.202		5.0791×10^{-5}	0.9950	1.25
6	3.25217×10^{-5}	2.6025	1.86		2.0868×10^{-5}	2.6921	1.845
$r = 0.002$		$\varepsilon = 1\%$			$\varepsilon = 5\%$		
$\mathfrak{M}$	$\ \epsilon_{\mathfrak{M}}\ _2$	CO	CPU-time		$\ \epsilon_{\mathfrak{M}}\ _2$	CO	CPU-time
2	7.1734×10^{-5}	—	0.781		8.3222×10^{-5}	—	0.752
4	4.3864×10^{-5}	0.7096	1.39		4.75428×10^{-5}	0.9345	1.328
6	1.6819×10^{-5}	2.36412	1.812		2.7453×10^{-5}	2.5627	1.766

Figure 7 shows the exact function $f(x)$ and its calculated approximate values using the proposed collocation scheme without regularization for $\alpha = 1.6$, $\hat{t} = 1$, $\mathfrak{N} = 12$, $\mathfrak{M} = 6$, and $\hat{\mathfrak{M}} = 300$. The exact and the numerical solutions using the mollification and Tikhonov regularizations are shown in Figure 8, where the mollification method clearly performs better than the Tikhonov method. In fact, in the mollification method, the noisy data is smoothed and then used to find the solution. In other words, the mollification method aims to remove the effects of noise or large fluctuations without altering the underlying structure of the data, leading to more accurate and stable computations. Thus, these results clarify that the mollification regularization scheme is a valid method to achieve sustainable solutions compared to the Tikhonov regularization technique.

The exact and computed approximation of $f(x)$ with regularization when $\alpha = 1.6$, $\hat{t} = 1$, $\mathfrak{N} = 12$, $\mathfrak{M} = 6$, and $\hat{\mathfrak{M}} = 300$, are shown in Figure 9. Figure 10 displays the exact and numerical approximation of u and Figure 11 shows the AE of the approximations for u , when $\alpha = 1.6$, $\hat{t} = 1$, $\mathfrak{N} = 12$, $\mathfrak{M} = 6$, and $\hat{\mathfrak{M}} = 300$. Table 2 shows the L_2 -norm error, CO , and CPU-times for the approximated values of u , when $\alpha = 1.6$, $\hat{t} = 1$, and $\mathfrak{N} = 12$. This table confirms that the numerical error decreases with decreasing noise level and decreasing grid size.



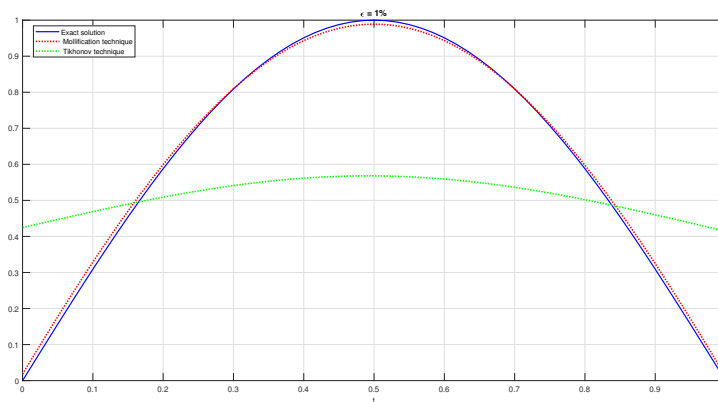


FIGURE 8. The exact and numerical approximations of $f(x)$ with the mollification and the Tikhonov methods when $\varepsilon = 1\%$ in Example 5.2.

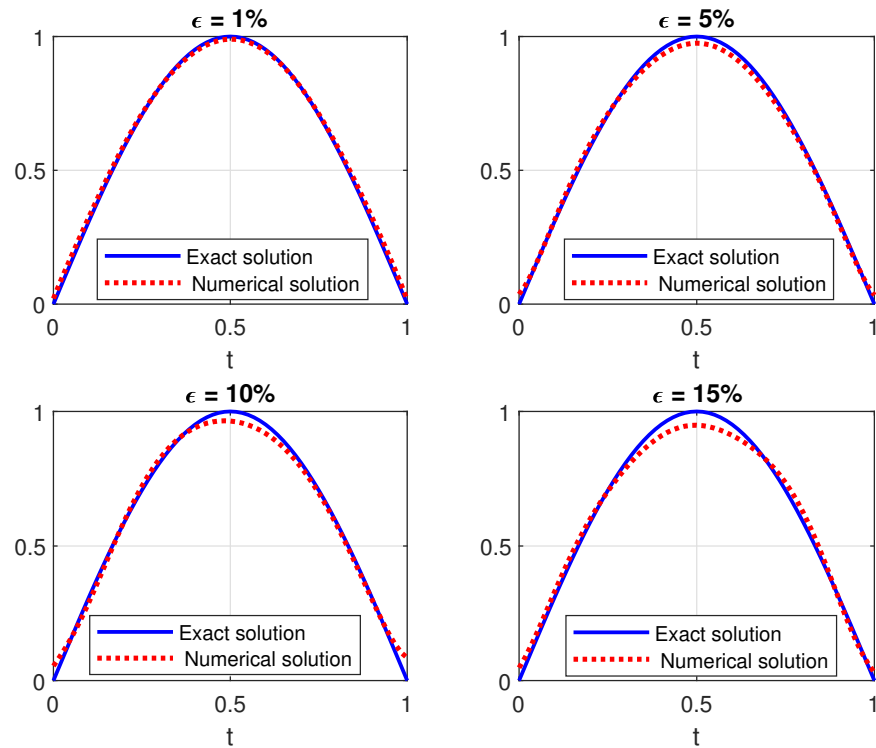


FIGURE 9. The exact function $f(x)$ and its computed approximations with regularization for different noise levels in Example 5.2.

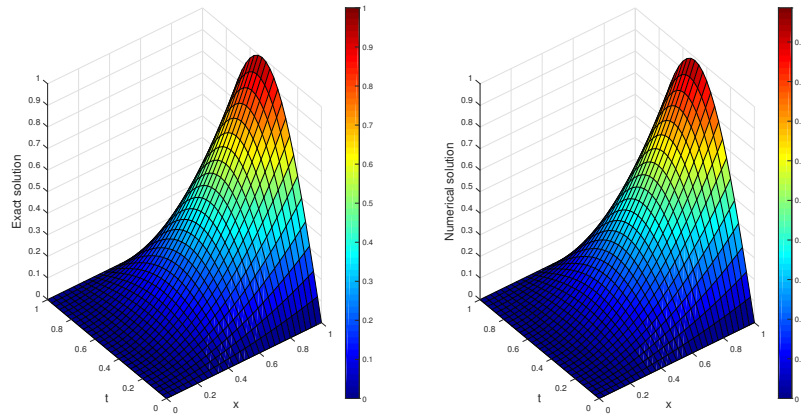


FIGURE 10. The exact and numerical approximation of u with $\varepsilon = 5\%$ in Example 5.2.

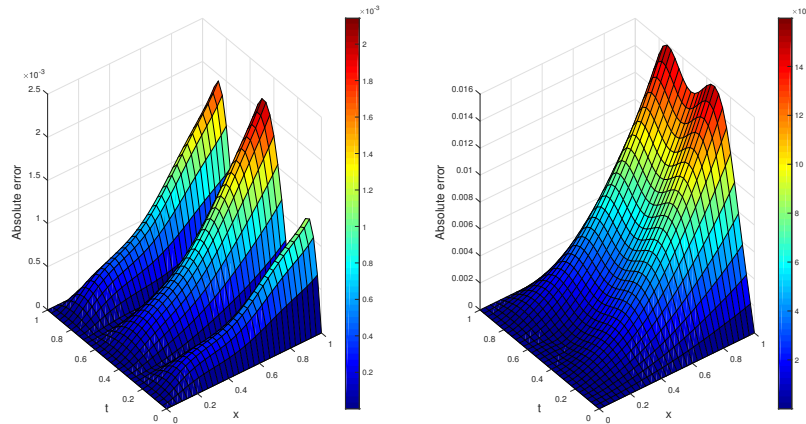


FIGURE 11. The AE for the approximations of u with $\varepsilon = 1\%$ (left) and $\varepsilon = 5\%$ (right) in Example 5.2.

6. CONCLUSION

In this study, a combination of SKCPs collocation method and mollification regularization technique was presented to solve an inverse source problem based on FDWE with an unknown space-dependent source function. The numerical method was expressed and the required convergence conditions were explored. Finally, two numerical implementations were used to investigate the performance of the proposed scheme. The numerical results show the accuracy and high quality of the numerical method.

REFERENCES

- [1] W. M. Abd-Elhameed and Y. H. Youssri, *Sixth-kind chebyshev spectral approach for solving fractional differential equations*, Int. J. Nonlinear Sci. Numer. Simul., 20(2) (2019), 191–203.
- [2] A. Babaei and S. Banihashemi, *Reconstructing unknown nonlinear boundary conditions in a time-fractional inverse reaction-diffusion-convection problem*, Numer. Methods Partial Differ. Equ., 35(3) (2019), 976–992.

- [3] X. Cheng and Z. Li, *Uniqueness and stability for inverse source problem for fractional diffusion-wave equations*, J. Inverse Ill-Posed Probl., 31(6) (2023), 885–904.
- [4] A. Favini and A. Lorenzi, *Differential equations: inverse and direct problems*, Chapman and Hall/CRC, New York, 2006.
- [5] X. H. Gong and T. Wei, *Reconstruction of a time-dependent source term in a time fractional diffusion-wave equation*, Inverse Probl. Sci. Eng., 27 (2019), 1577–1594.
- [6] R. Herrmann, *Fractional calculus: an introduction for physicists*, World Scientific, Singapore, 2011.
- [7] H. Jafari, A. Babaei, and S. Banihashemi, *A novel approach for solving an inverse reaction-diffusion-convection problem*, J. Optim. Theory Appl., 183 (2019), 688–704.
- [8] A. Kirsch, *An introduction to the mathematical theory of inverse problems*, Springer, New York, 2011.
- [9] D. A. Murio, *Stable numerical solution of a fractional-diffusion inverse heat conduction problem*, Comput. Math. Appl., (2007), 1492–1501.
- [10] S. Nemati and A. Babaei, *A numerical method based on the Jacobi polynomials to reconstruct an unknown source term in a time fractional diffusion-wave equation*, Taiwan. J. Math. 23 (2019), 1271–1289.
- [11] Y. Park, L. Reichel, G. Rodriguez, and X. Yu, *Parameter determination for Tikhonov regularization problems in general form*, J. Comput. Appl. Math., 334 (2018), 12–25.
- [12] K. Woodbury, *Inverse Engineering Handbook*, CRC Press, Boca Raton FL, (2003).
- [13] J. Xian and T. Wei, *Determination of the initial data in a time-fractional diffusion-wave problem by a final time data*, Comput. Math. Appl., 78 (2019), 2525–2540.
- [14] X. B. Yan and T. Wei, *Determine a space-dependent source term in a time fractional diffusion-wave equation*, Acta Appl. Math., 165 (2020), 163–181.
- [15] F. Yang, Q. Pu, X. X. Li, and D. G. Li, *The truncation regularization method for identifying the initial value on non-homogeneous time-fractional diffusion wave equations*, J. Math., 7 (2019), 1007.
- [16] J. Zhang and X. Yang, *A class of efficient difference method for time fractional reaction-diffusion equation*, Comput. Appl. Math., 37 (2018), 4376–4396.

