

Nonnegative Low-Rank Matrix Correction under an Orthogonality Constraint in Conservative Vlasov Simulations

Yue Wu*, Stephen Becker†, Jingmei Qiu‡, Xiangxiong Zhang§

July 28, 2026

Abstract

In low-rank numerical methods for Vlasov dynamics, the SVD-type truncation procedure may introduce negative entries into the numerical solution. Such negative values are unphysical because the solution is a probability distribution function. We design optimization-based post-processing algorithms to recover nonnegativity while preserving the macroscopic quantities (density, momentum, and energy) pointwise. The preservation of the macroscopic quantities is written as an orthogonality constraint on the correction term. For a convex formulation based on squared nuclear norm minimization, we show that the proximal operator with the orthogonality constraint is characterized by an implicit singular value thresholding equation, and the threshold can be computed efficiently by bisection. Based on this result, we develop five algorithms for the convex formulation: Douglas–Rachford splitting, restarted dual FISTA, restarted dual accelerated gradient descent, dual PR+ conjugate gradient, and dual L-BFGS. We also consider a non-convex formulation with an explicit rank constraint and develop a tangent-space accelerated alternating projection algorithm that only requires a $2r \times 2r$ SVD per iteration. Numerical results for a Landau damping test case show that the proposed algorithms give comparable correction quality. Among them, the tangent-space accelerated alternating projection is the most cost-efficient, increasingly so as the problem size grows. We further demonstrate the correction as a positivity limiter inside a time-dependent conservative low-rank Vlasov solver, where it removes the negativity introduced by the SVD-type truncation while preserving the conserved mass, momentum, and energy.

Keywords: Nonnegative low-rank matrix approximation · Nuclear norm minimization · Orthogonality constraint · Convex optimization · Vlasov dynamics · Alternating projection on manifolds

Mathematics Subject Classification (2020) 65F55 · 15A83 · 90C25 · 90C26 · 65K10 · 35Q83

*Division of Applied Mathematics, Brown University, 182 George Street, Providence, RI 02912 (yue_wu3@brown.edu).

†Department of Applied Mathematics, University of Colorado Boulder, 526 UCB, Boulder, CO 80309 (stephen.becker@colorado.edu).

‡Department of Mathematical Sciences, University of Delaware, 501 Ewing Hall, Newark, DE 19716 (jingqiu@udel.edu).

§Corresponding author. Department of Mathematics, Purdue University, 150 N. University Street, West Lafayette, IN 47907 (zhan1966@purdue.edu). This work is partially supported by NSF DMS-2208518.

1 Introduction

Numerical simulation of the Vlasov–Poisson (VP) system is fundamental to understanding the complex dynamics of collisionless plasmas and has a wide range of applications in science and engineering, such as fusion energy and space weather prediction. The VP system describes the evolution of the probability distribution function $f(\mathbf{x}, \mathbf{v}, t)$ in phase space, which by definition must be nonnegative. Its velocity moments, including mass density, momentum density, and energy density, satisfy a system of macroscopic conservation laws. Preserving both nonnegativity and conservation at the discrete level is essential for producing physically meaningful solutions and for the nonlinear stability of the numerical scheme, especially in demanding regimes such as low density and pressure.

The high dimensionality of phase space poses a major computational challenge for deterministic VP solvers. Recently, low-rank and adaptive-rank tensor methods [9] have emerged as promising approaches for mitigating the curse of dimensionality by adaptively exploiting low-rank structure in the solution. Broadly, these methods can be divided into two classes. The first is the dynamical low-rank approach, which evolves the low-rank factors by projecting the governing equation onto the tangent space of the low-rank manifold [16]. Robust and rank-adaptive time integrators have been developed within this framework [20, 5, 4, 7]. The second is the step-and-truncate approach, in which the solution is first advanced in an enriched approximation space and subsequently compressed to control the rank. Within a Galerkin framework, the method proposed in [12] dynamically constructs an adaptive low-rank solution basis by adding new basis functions generated from the discretized PDE and removing redundant components through SVD-type truncation. Within a collocation framework [33, 34], selected rows and columns of the solution matrix are adaptively identified, allowing the solution to be updated through CUR-type matrix interpolation without forming the full high-dimensional solution. An important challenge for both classes of low-rank methods is the preservation of the physical conservation laws. Conservative SVD truncation procedures were developed in [8, 13] to preserve local mass and momentum. Recently, the Local Macroscopic Conservative (LoMaC) low-rank tensor methods developed in [14, 11] preserve mass, momentum, and energy locally at the discrete level. Their central idea is to evolve the macroscopic conservation laws alongside the kinetic equation, construct a reference subspace from the resulting macroscopic densities, and project the low-rank kinetic solution onto this subspace through a conservative orthogonal projection.

A remaining issue for the low-rank methods in [13, 14] is the lack of nonnegativity in the numerical solution. Since the distribution function f represents a probability density, it must be nonnegative everywhere. However, SVD-type low-rank approximations of a nonnegative matrix can introduce negative entries, and this problem is further compounded by the time-stepping procedure. Within the low-rank framework, positivity can be enforced by evolving a square-root factorization $f = g^2$, so that the represented solution is nonnegative by construction [31]. This reformulation is tied to the structure of the transport operator and does not extend to second-order operators such as a Fokker–Planck collision term, and a nonlinear transformation of a low-rank function need not remain low rank. We instead leave the low-rank solver unchanged and recover nonnegativity by a minimal post-processing correction, which is agnostic to the underlying operator. Enforcing nonnegativity (or more generally, preserving the invariant domain of the underlying PDE) in high-order numerical methods has been an active area of research. We refer to [30] for a comprehensive survey. The main approaches include polynomial limiters based on convex decomposition of cell averages [24, 32], flux-corrected transport (FCT) methods, and more recently, optimization-based limiters. In

the optimization-based approach, the limiter seeks a minimal modification of the numerical solution while enforcing conservation and positivity (or bounds) as constraints. Such optimization-based methods have been developed for enforcing nonnegativity of polynomial approximations [6] and for invariant-domain-preserving limiters in gas dynamics [18], where first-order operator splitting methods such as the Douglas–Rachford splitting [17] are employed to efficiently solve the resulting constrained minimization problems. At the representation level, [28] proposes a nonnegative tensor train optimization framework for high-dimensional distribution tensors, whose complexity scales linearly with both the dimension d and the number of degrees of freedom per dimension n .

In this work, we design optimization-based post-processing algorithms to recover the nonnegativity of the low-rank solution produced by [13, 14], while preserving the macroscopic quantities (density, momentum, and energy) pointwise. The problem is to find a correction term $\mathbf{X}$ such that $\mathbf{A} + \mathbf{X} \geq \mathbf{0}$ and $\mathbf{X}\mathbf{B} = \mathbf{0}$, where $\mathbf{A}$ is the possibly negative scaled solution matrix and $\mathbf{B}$ encodes the orthogonality constraint associated with the first three velocity moments. We further split $\mathbf{A} = \mathbf{A}_1 + \mathbf{A}_2$, where $\mathbf{A}_1$ is a rank-3 part that carries the macroscopic quantities and $\mathbf{A}_2$ is the remainder. To make subsequent computations efficient, it is desirable for the correction $\mathbf{X}$ to be low-rank as well. A standard approach is to use nuclear norm minimization [3] as a convex relaxation of rank minimization, which leads to the formulation

$$\begin{aligned} \min_{\mathbf{X}} \quad & \|\mathbf{X}\|_* + \frac{a}{2} \|\mathbf{X} - \mathbf{A}_2\|_F^2, \\ \text{subject to:} \quad & \mathbf{X} \geq -\mathbf{A}_1, \quad \mathbf{X}\mathbf{B} = \mathbf{0}, \end{aligned} \tag{1}$$

where $\|\cdot\|_*$ denotes the nuclear norm (sum of singular values) and $a > 0$ is a penalty parameter that controls the trade-off between low rank and approximation error, and the corrected solution is $\mathbf{A}_1 + \mathbf{X}$. However, this formulation has two drawbacks. It is not consistent, since if $\mathbf{A}$ is already nonnegative, the solution is not necessarily the identity correction. It is also not homogeneous, since scaling the input does not scale the output. To address these two issues, we reinterpret $\mathbf{X}$ as a correction term added to $\mathbf{A}$ rather than an approximation to $\mathbf{A}_2$, so the Frobenius term penalizes the magnitude of the correction itself rather than an approximation error. We also replace the nuclear norm by its square, which makes the objective homogeneous of degree two. This leads to the formulation

$$\begin{aligned} \min_{\mathbf{X}} \quad & \|\mathbf{X}\|_*^2 + \frac{a}{2} \|\mathbf{X}\|_F^2, \\ \text{subject to:} \quad & \mathbf{X} \geq -\mathbf{A}, \quad \mathbf{X}\mathbf{B} = \mathbf{0}, \end{aligned} \tag{2}$$

where $\mathbf{X}$ is now a correction term and the corrected solution is $\mathbf{A} + \mathbf{X}$. This formulation is both consistent and homogeneous. One major computational challenge is to evaluate the proximal operator of the squared nuclear norm under the orthogonality constraint $\mathbf{X}\mathbf{B} = \mathbf{0}$. For the standard nuclear norm in (1), the proximal operator is the singular value soft thresholding operator [3]. For the squared nuclear norm in (2), the proximal operator can still be expressed as singular value soft thresholding, but the threshold is now implicit: it depends on the nuclear norm of the unknown minimizer itself. We derive this characterization in Theorem 1 and show that the implicit threshold can be computed by a bisection search that requires only $\mathcal{O}(\log(\min(m, n)))$ singular value comparisons for an $m \times n$ matrix.

We consider two formulations for the nonnegative correction problem. For the convex problem (2), we develop five algorithms: Douglas–Rachford splitting [17] on the primal problem, restarted dual FISTA [2], restarted dual accelerated gradient descent, dual Polak–Ribière plus conjugate gradient, and dual L-BFGS [19]. These algorithms converge to a global minimizer. We also consider

a non-convex formulation with an explicit rank constraint on the correction term. For this formulation, we develop a tangent-space accelerated alternating projection (TAP) algorithm inspired by [26]. The key idea is to work on the tangent space of the low-rank manifold so that each iteration only requires a $2r \times 2r$ SVD for a rank r solution instead of a full SVD. Beyond the fixed correction problem, we demonstrate the correction as a positivity limiter inside the time-dependent LoMaC conservative low-rank scheme of [14]. Applied periodically as the solution is advanced, it removes the negativity introduced by the SVD-type truncation while leaving the conserved moments untouched, so that mass, momentum, and energy remain conserved.

The rest of the paper is organized as follows. In Section 2, we describe the problem background, including the low-rank solution structure for Vlasov dynamics and the conservative orthogonal projection from [13, 14]. In Section 3, we formulate the nonnegative low-rank matrix correction problem with orthogonality constraint and present three optimization formulations. In Section 4, we develop the numerical algorithms for both the convex and non-convex formulations. In Section 5, we present numerical results on a Landau damping test case that demonstrate the effectiveness and compare the performance of all proposed algorithms, study their empirical cost scaling with problem size, and demonstrate the correction as a positivity limiter in a time-dependent conservative low-rank Vlasov solver. We conclude in Section 6.

2 Low-rank Vlasov dynamics and the orthogonality constraint

2.1 Low-rank Vlasov solutions

We consider a 1D1V Vlasov system [12]

$$\begin{aligned} \frac{\partial f}{\partial t} + v \cdot \nabla_x f - E(x, t) \cdot \nabla_v f &= 0, \quad \Omega_x \times \Omega_v \times (0, T], \\ -\Delta_x \phi &= \rho(x, t), \quad \rho(x, t) = \rho_0 - \int_{\mathbb{R}^{d_v}} f(x, v, t) dv, \quad \Omega_x \times (0, T], \\ E(x, t) &= -\nabla_x \phi, \quad \Omega_x \times (0, T], \end{aligned}$$

where f is the probability density function of charged particles, with a corresponding charge density ρ , which generates a self-consistent electrostatic potential ϕ via the Poisson equation. The Vlasov equation is discretized on the phase space $[x_{\min}, x_{\max}] \times [-v_{\max}, v_{\max}]$ with an equidistant $N_x \times N_v$ Cartesian grid

$$\begin{aligned} x_{\text{grid}} : x_{\min} = x_1 &< \cdots < x_{N_x} = x_{\max}, \\ v_{\text{grid}} : -v_{\max} = v_1 &< \cdots < v_{N_v} = v_{\max}. \end{aligned}$$

We denote the mesh size in the v -direction by h_v . The numerical solution $\mathbf{F} \in \mathbb{R}^{N_x \times N_v}$, with $\mathbf{F}_{i,j} \approx f(x_i, v_j)$ approximating the probability distribution function $f(x, v)$, is stored in a low-rank format

$$\mathbf{F} = \mathbf{F}_1 + \mathbf{F}_2 = \mathbf{U}_1 \mathbf{V}_1^T + \mathbf{U}_2 \mathbf{V}_2^T, \quad (3)$$

where $\mathbf{F}_1$ is a rank-3 term that carries the three macroscopic quantities (density, momentum, and internal energy) of $\mathbf{F}$, and $\mathbf{F}_2$ is the remainder. This decomposition is the key structure underlying the conservative low-rank tensor methods in [13, 14].

The three macroscopic quantities are defined as velocity moments of $f(x, v)$:

$$\rho(x) = \int f(x, v) dv, \quad (4a)$$

$$m(x) = \int v f(x, v) dv, \quad (4b)$$

$$e(x) = \int v^2 f(x, v) dv. \quad (4c)$$

Let $\mathbf{1} = (1, \dots, 1) \in \mathbb{R}^{1 \times N_v}$, $\mathbf{v} = (v_1, \dots, v_{N_v}) \in \mathbb{R}^{1 \times N_v}$, and $\mathbf{v}^{\circ 2} = (v_1^2, \dots, v_{N_v}^2) \in \mathbb{R}^{1 \times N_v}$. The discrete macroscopic quantities are computed as

$$\boldsymbol{\rho} = h_v \mathbf{F} \mathbf{1}^T = h_v \mathbf{F}_1 \mathbf{1}^T, \quad (5a)$$

$$\mathbf{m} = h_v \mathbf{F} \mathbf{v}^T = h_v \mathbf{F}_1 \mathbf{v}^T, \quad (5b)$$

$$\mathbf{e} = h_v \mathbf{F} (\mathbf{v}^{\circ 2})^T = h_v \mathbf{F}_1 (\mathbf{v}^{\circ 2})^T, \quad (5c)$$

where the second equalities hold by the requirement that $\mathbf{F}_1$ carries all three macroscopic quantities, i.e., $\mathbf{F}_2$ has zero velocity moments.

2.2 Conservative orthogonal decomposition

Following [13], we construct the decomposition $\mathbf{F} = \mathbf{F}_1 + \mathbf{F}_2$ using a weighted inner product in the velocity space. Let

$$\mathbf{R} = \begin{pmatrix} \mathbf{1} \\ \mathbf{v} \\ \mathbf{v}^{\circ 2} \end{pmatrix} \in \mathbb{R}^{3 \times N_v} \quad (6)$$

be the matrix whose rows span the space of velocity moments, and let

$$\mathbf{S} = \text{diag}(w(v_1), w(v_2), \dots, w(v_{N_v})) \in \mathbb{R}^{N_v \times N_v} \quad (7)$$

be a diagonal weight matrix with pointwise-positive weight function $w(v) > 0$. The requirement that $\mathbf{F}_2$ has zero velocity moments is expressed as

$$\mathbf{F}_2 \mathbf{R}^T = \mathbf{0}, \quad \text{or equivalently,} \quad \mathbf{F}_1 \mathbf{R}^T = \mathbf{F} \mathbf{R}^T. \quad (8)$$

Additionally, we require $\text{rank}(\mathbf{F}_1) = 3$ and that the row space of $\mathbf{F}_1$ lies in the column space of $\mathbf{R} \mathbf{S}$. This gives the system

$$\begin{cases} \mathbf{F}_1 = \mathbf{L} \mathbf{R} \mathbf{S}, \\ \mathbf{F}_1 \mathbf{R}^T = \mathbf{F} \mathbf{R}^T, \end{cases} \quad (9)$$

where $\mathbf{L} \in \mathbb{R}^{N_v \times 3}$. To solve for $\mathbf{F}_1$ explicitly, we compute a QR decomposition

$$\sqrt{\mathbf{S}} \mathbf{R}^T = \mathbf{B} \tilde{\mathbf{R}}, \quad (10)$$

where $\mathbf{B} \in \mathbb{R}^{N_v \times 3}$ has orthonormal columns and $\tilde{\mathbf{R}} \in \text{GL}_3(\mathbb{R})$. Then $\mathbf{F}_1$ is given by

$$\mathbf{F}_1 = \mathbf{F} \sqrt{\mathbf{S}}^{-1} \mathbf{B} \mathbf{B}^T \sqrt{\mathbf{S}}. \quad (11)$$

Thus, $\mathbf{F}_1$ is the orthogonal projection of $\mathbf{F}$ (with respect to the weighted inner product defined by $\mathbf{S}$) onto the three-dimensional subspace spanned by the velocity moment basis functions $\mathbf{1}, \mathbf{v}, \mathbf{v}^{\circ 2}$.

2.3 Rescaled problem and the orthogonality constraint

To formulate the correction problem, we introduce the rescaled matrices

$$\mathbf{A}_1 = \mathbf{F}_1 \sqrt{\mathbf{S}}^{-1}, \quad (12a)$$

$$\mathbf{A}_2 = \mathbf{F}_2 \sqrt{\mathbf{S}}^{-1}, \quad (12b)$$

$$\mathbf{A} = \mathbf{A}_1 + \mathbf{A}_2 = \mathbf{F} \sqrt{\mathbf{S}}^{-1}. \quad (12c)$$

The rescaling converts the weighted inner product to the standard Frobenius inner product, so that the orthogonal projection (11) takes the simple form $\mathbf{A}_1 = \mathbf{A} \mathbf{B} \mathbf{B}^T$. Note that $\mathbf{A}_1$ and $\mathbf{A}_2$ inherit the low-rank structure of $\mathbf{F}_1$ and $\mathbf{F}_2$, and the nonnegativity of $\mathbf{F}$ is equivalent to the nonnegativity of $\mathbf{A}$ (since $\mathbf{S}$ has positive diagonal entries).

A remaining issue for adaptive rank approaches is that the matrix $\mathbf{A}$ (or equivalently $\mathbf{F}$) may contain negative entries. We seek a correction term $\mathbf{X}$ such that $\mathbf{A} + \mathbf{X} \geq 0$. We also require the corrected solution $\mathbf{A} + \mathbf{X}$ to have the same macroscopic quantities as $\mathbf{A}$. Since $\mathbf{A}_1 = \mathbf{A} \mathbf{B} \mathbf{B}^T$ carries the macroscopic information, the requirement that the correction $\mathbf{X}$ does not alter the macroscopic quantities is equivalent to

$$\mathbf{X} \mathbf{B} = \mathbf{0}. \quad (13)$$

This is the orthogonality constraint: the correction term $\mathbf{X}$ must lie in the orthogonal complement of the column space of $\mathbf{B}$. In the following, we use $\mathbf{B} \in \mathbb{R}^{N_v \times 3}$ to denote the matrix with orthonormal columns defined in this section.

3 Nonnegative low-rank matrix correction: optimization formulations

With the setup from the previous section, the goal is to find a correction term $\mathbf{X}$ such that the corrected solution $\mathbf{A} + \mathbf{X}$ is nonnegative and preserves the macroscopic quantities. The constraints are

$$\mathbf{A} + \mathbf{X} \geq 0, \quad \mathbf{X} \mathbf{B} = \mathbf{0}.$$

Ideally, the correction $\mathbf{X}$ should have low rank so that the corrected solution retains a compact low-rank representation. We present three optimization formulations for this problem.

3.1 The approximation formulation

We first seek a nonnegative low-rank approximation $\mathbf{X}$ to the remainder $\mathbf{A}_2$, so that the corrected solution $\mathbf{A}_1 + \mathbf{X}$ is nonnegative and preserves the macroscopic quantities. The ideal problem is

$$\min_{\mathbf{X}} \text{rank}(\mathbf{X}) \quad \text{and} \quad \|\mathbf{X} - \mathbf{A}_2\|_F, \quad \text{subject to} \quad \mathbf{X} + \mathbf{A}_1 \geq 0, \quad \mathbf{X} \mathbf{B} = \mathbf{0}. \quad (14)$$

However, rank minimization is NP-hard in general. A standard convex relaxation technique is to replace the rank function by the nuclear norm [3].

Definition 1 (nuclear norm). *For any matrix $\mathbf{A} \in \mathbb{R}^{m \times n}$, its nuclear norm is defined as*

$$\|\mathbf{A}\|_* = \sum_{k=1}^{\min\{m,n\}} \sigma_k, \quad (15)$$

where σ_k is the k -th singular value of $\mathbf{A}$.

The nuclear norm is the tightest convex lower bound of the rank function on the unit spectral norm ball. Minimizing the nuclear norm therefore promotes low-rank solutions in the same way that minimizing the ℓ^1 norm of a vector promotes sparsity. We record a standard property that will be used in the algorithm derivation, using $\|\cdot\|_2$ to denote the spectral norm.

Proposition 1. *Suppose the compact SVD of $\mathbf{X} \in \mathbb{R}^{m \times n}$ is $\mathbf{X} = \mathbf{U}\mathbf{\Sigma}\mathbf{V}^T$. Then the subdifferential of $\|\cdot\|_*$ at $\mathbf{X}$ is*

$$\partial \|\mathbf{X}\|_* = \{\mathbf{U}\mathbf{V}^T + \mathbf{W} \mid \mathbf{U}^T \mathbf{W} = \mathbf{0}, \mathbf{W}\mathbf{V} = \mathbf{0}, \|\mathbf{W}\|_2 \leq 1\}. \quad (16)$$

Using the nuclear norm as a convex surrogate for the rank, and choosing the Frobenius norm for the approximation error, we obtain the following convex optimization problem:

$$\begin{aligned} \min_{\mathbf{X}} \quad & \|\mathbf{X}\|_* + \frac{a}{2} \|\mathbf{X} - \mathbf{A}_2\|_F^2, \\ \text{subject to:} \quad & \mathbf{X} \geq -\mathbf{A}_1, \mathbf{X}\mathbf{B} = \mathbf{0}, \end{aligned}$$

where $a > 0$ is a parameter that controls the trade-off between low rank and approximation accuracy. This is precisely the formulation in (1).

3.2 The correction formulation (convex)

The formulation (1) has two drawbacks. First, it is not consistent: if $\mathbf{A}$ is already nonnegative, the minimizer $\mathbf{X}$ of (1) is not necessarily $\mathbf{A}_2$, meaning that the post-processing may alter a solution that needs no correction. This can be seen from the following example.

Example 1. *Drop the orthogonality constraint and consider the diagonal case: $\mathbf{A}_1 = \mathbf{0}$ and $\mathbf{A}_2 = \text{diag}(a_1, \dots, a_n)$ with $a_i \geq 0$. Then $\|\mathbf{X}\|_* = \sum_i |x_i|$ reduces to the ℓ^1 norm of the diagonal, and singular value thresholding acts entrywise, so the minimizer of (1) is $\mathbf{X} = \text{diag}(x_1, \dots, x_n)$ with*

$$x_i = \max \left\{ a_i - \frac{1}{a}, 0 \right\},$$

which differs from $\mathbf{A}_2$ unless $a_i \geq 1/a$ for all i .

Second, the formulation is not homogeneous: if the input is scaled by $\alpha > 0$, the minimizer does not scale by the same factor. This can also be verified through Example 1.

To address these two issues, we seek a *correction term* $\mathbf{X}$ to be added to $\mathbf{A}$, rather than an approximation to $\mathbf{A}_2$. We also replace the nuclear norm by its square, which makes the objective homogeneous of degree two. This leads to the formulation

$$\begin{aligned} \min_{\mathbf{X}} \quad & \|\mathbf{X}\|_*^2 + \frac{a}{2} \|\mathbf{X}\|_F^2, \\ \text{subject to:} \quad & \mathbf{X} \geq -\mathbf{A}, \mathbf{X}\mathbf{B} = \mathbf{0}. \end{aligned}$$

which is the formulation (2). Here $\mathbf{X}$ is the correction term and the corrected solution is $\mathbf{A} + \mathbf{X}$. The constraint $\mathbf{X} \geq -\mathbf{A}$ ensures nonnegativity of the corrected solution, and the constraint $\mathbf{X}\mathbf{B} = \mathbf{0}$ ensures preservation of the macroscopic quantities. The squared nuclear norm $\|\mathbf{X}\|_*^2$ promotes low

rank (since $\text{rank}(\mathbf{X} + \mathbf{Y}) \leq \text{rank}(\mathbf{X}) + \text{rank}(\mathbf{Y})$, a small correction should have low rank), while the Frobenius norm penalty $\frac{a}{2} \|\mathbf{X}\|_F^2$ penalizes large corrections. We refer to (2) as the convex formulation in the rest of the paper.

This formulation is consistent: if $\mathbf{A} \geq 0$, then $\mathbf{X} = \mathbf{0}$ is clearly the minimizer. It is also homogeneous: if the input is $\alpha\mathbf{A}$, the minimizer is $\alpha\mathbf{X}$.

Remark 1. *In practice, we have not observed significant differences in the numerical results between the formulations (1) and (2) when the parameter a is well chosen.*

Remark 2 (Non-emptiness of the admissible set). *The admissible set $\{\mathbf{X} \in \mathbb{R}^{m \times n} \mid \mathbf{X} \geq -\mathbf{A}, \mathbf{X}\mathbf{B} = \mathbf{0}\}$ may be empty. It is nonempty if and only if there exists a nonnegative matrix $\tilde{\mathbf{A}} \in \mathbb{R}^{m \times n}$ with $\tilde{\mathbf{A}}\mathbf{B}\mathbf{B}^T = \mathbf{A}\mathbf{B}\mathbf{B}^T$, i.e., the macroscopic part of $\mathbf{A}$ coincides with the macroscopic part of some nonnegative distribution. In the Vlasov dynamics setting, this condition is satisfied whenever the macroscopic moments (ρ, m, e) are realizable, i.e., there exists a nonnegative distribution with these moments. For the first three moments, a necessary and sufficient condition is $\rho > 0$ and $ep \geq m^2$.*

3.3 The rank-constrained formulation (non-convex)

In some applications, it is desirable to explicitly control the rank of the correction term so that the corrected solution can be stored efficiently in a low-rank format. For a given target rank r , we consider the non-convex formulation

$$\begin{aligned} & \min_{\mathbf{X}} \|\mathbf{X}\|_F^2, \\ & \text{subject to: } \text{rank}(\mathbf{X}) \leq r, \mathbf{X} \geq -\mathbf{A}, \mathbf{X}\mathbf{B} = \mathbf{0}. \end{aligned} \tag{17}$$

Unlike the convex formulations (1) and (2), the rank constraint makes the feasible set non-convex, so global optimality cannot be guaranteed. However, this formulation has the computational advantage that the iterates can be maintained in a rank- r factored form throughout the algorithm, as we will see in Section 4. We refer to (17) as the non-convex formulation in the rest of the paper.

4 Algorithms

In this section, we develop algorithms for both the convex formulation (2) and the non-convex formulation (17). For the convex problem, we present five algorithms in Sections 4.2–4.6: Douglas–Rachford splitting on the primal problem, restarted dual FISTA, restarted dual accelerated gradient descent, dual PR+ conjugate gradient, and dual L-BFGS. For the non-convex problem, we develop a tangent-space accelerated alternating projection algorithm in Section 4.7. We begin with a brief review of the optimization tools used throughout this section. Readers familiar with convex optimization may skip to Section 4.2.

4.1 Convex optimization preliminaries

The convex formulation (2) involves minimizing a non-smooth objective (the squared nuclear norm plus a Frobenius penalty) subject to non-smooth constraints (nonnegativity and orthogonality). We briefly review the key concepts needed for the algorithms.

Indicator functions and subdifferentials. To incorporate constraints into an unconstrained framework, we use the indicator function of a convex set $\mathcal{C}$:

$$I_{\mathcal{C}}(\mathbf{X}) = \begin{cases} 0 & \text{if } \mathbf{X} \in \mathcal{C}, \\ +\infty & \text{otherwise,} \end{cases}$$

so that $\min_{\mathbf{X} \in \mathcal{C}} \varphi(\mathbf{X}) = \min_{\mathbf{X}} \varphi(\mathbf{X}) + I_{\mathcal{C}}(\mathbf{X})$. For a convex but non-differentiable function φ , the gradient is replaced by the subdifferential $\partial\varphi(\mathbf{X}_0)$, defined as the set of all $\mathbf{G}$ satisfying $\varphi(\mathbf{Y}) \geq \varphi(\mathbf{X}_0) + \langle \mathbf{G}, \mathbf{Y} - \mathbf{X}_0 \rangle$ for all $\mathbf{Y}$, where the inner product is defined $\langle \mathbf{Y}, \mathbf{X} \rangle = \text{tr}(\mathbf{Y}^T \mathbf{X})$. When φ is differentiable, $\partial\varphi = \{\nabla\varphi\}$. As a scalar example, $\partial|x|(0) = [-1, 1]$. The subdifferential of the nuclear norm is given in Proposition 1.

Proximal operators. Many convex optimization problems can be interpreted as finding the steady state of the (sub)gradient flow

$$\dot{\mathbf{X}} \in -\partial\varphi(\mathbf{X}),$$

where $\partial\varphi$ denotes the subdifferential of a proper, closed, convex function φ . When φ is differentiable, this reduces to the familiar gradient flow

$$\dot{\mathbf{X}} = -\nabla\varphi(\mathbf{X}).$$

Applying the backward Euler method with time step Δt gives

$$\frac{\mathbf{X}^{k+1} - \mathbf{X}^k}{\Delta t} \in -\partial\varphi(\mathbf{X}^{k+1}),$$

or equivalently,

$$\mathbf{0} \in \partial\varphi(\mathbf{X}^{k+1}) + \frac{1}{\Delta t}(\mathbf{X}^{k+1} - \mathbf{X}^k).$$

Remarkably, this is precisely the first-order optimality condition of the optimization problem

$$\mathbf{X}^{k+1} = \arg \min_{\mathbf{X}} \left\{ \varphi(\mathbf{X}) + \frac{1}{2\Delta t} \|\mathbf{X} - \mathbf{X}^k\|_F^2 \right\}. \quad (18)$$

Identifying $\eta = 1/\Delta t$ yields the proximal operator

$$\text{prox}_{\eta^{-1}\varphi}(\mathbf{Z}) = \arg \min_{\mathbf{X}} \left\{ \varphi(\mathbf{X}) + \frac{\eta}{2} \|\mathbf{X} - \mathbf{Z}\|_F^2 \right\},$$

which may therefore be viewed as a single implicit (backward Euler) step of the subgradient flow. Consequently, the proximal point iteration

$$\mathbf{X}^{k+1} = \text{prox}_{\eta^{-1}\varphi}(\mathbf{X}^k)$$

inherits the unconditional stability characteristic of backward Euler and converges to a minimizer of φ under standard assumptions. An important special case is the indicator function of a closed convex set,

$$I_{\mathcal{C}}(\mathbf{X}) = \begin{cases} 0, & \mathbf{X} \in \mathcal{C}, \\ +\infty, & \text{otherwise,} \end{cases}$$

whose proximal operator reduces to the Euclidean projection,

$$\text{prox}_{\eta^{-1}I_{\mathcal{C}}} = \Pi_{\mathcal{C}},$$

independent of η .

Splitting method. Many optimization problems take the composite form

$$\min_{\mathbf{X}} f(\mathbf{X}) + g(\mathbf{X}),$$

where different terms represent different objectives or constraints. The choice of numerical method depends on the structure of these two components. When both f and g are differentiable, their gradients can be combined,

$$\nabla(f + g) = \nabla f + \nabla g,$$

and standard gradient-based methods apply. If only one term (say f) is differentiable while g is nonsmooth, then one may take an explicit gradient step on f followed by an implicit (proximal) step on g , leading to the proximal-gradient method. The more challenging situation arises when both f and g are nonsmooth. In principle, one could apply the proximal point method directly to the sum,

$$\text{prox}_{\eta^{-1}(f+g)},$$

but this generally requires solving a difficult optimization problem involving both functions simultaneously. In contrast, the individual proximal operators,

$$\text{prox}_{\eta^{-1}f} \quad \text{and} \quad \text{prox}_{\eta^{-1}g},$$

are often available in closed form or can be computed efficiently because each function possesses a much simpler structure. Proximal splitting methods exploit this observation by replacing one difficult proximal evaluation of $f + g$ with a sequence of simpler proximal evaluations involving only f and g . This philosophy is analogous to operator splitting methods for differential equations, where a complicated evolution operator is decomposed into simpler subproblems that can be solved individually. Douglas–Rachford splitting is one of the most successful examples of such proximal splitting methods, alternating between the proximal operators of f and g while still converging to a minimizer of the composite problem. Our convex problem (2) falls precisely into this setting. The nonnegativity constraint and the orthogonality-constrained nuclear norm are both nonsmooth, making the proximal operator of their sum difficult to compute directly. However, each individual proximal operator admits an efficient closed-form solution, making Douglas–Rachford splitting a natural and effective approach.

4.2 Douglas–Rachford splitting for the convex problem

4.2.1 Splitting

To apply Douglas–Rachford splitting, the objective should be decomposed into two convex functions whose proximal operators are individually inexpensive to evaluate. The goal is therefore not to split the objective arbitrarily, but rather to isolate the different structures so that each proximal subproblem admits either a closed-form solution or an efficient numerical algorithm.

For our problem, the two nonsmooth components arise from the nonnegativity constraint and the orthogonality-constrained nuclear norm. These two structures are naturally handled by different proximal operators, leading to the decomposition

$$h(\mathbf{X}) = f(\mathbf{X}) + g(\mathbf{X}), \tag{19a}$$

$$f(\mathbf{X}) = \frac{\alpha_1}{2} \|\mathbf{X}\|_F^2 + I_{\{\mathbf{X} | \mathbf{X} \geq -\mathbf{A}\}}(\mathbf{X}), \tag{19b}$$

$$g(\mathbf{X}) = \|\mathbf{X}\|_*^2 + \frac{a_2}{2} \|\mathbf{X}\|_F^2 + I_{\{\mathbf{X}|\mathbf{X}\mathbf{B}=\mathbf{0}\}}(\mathbf{X}), \quad (19c)$$

where $a_1 + a_2 = a$.

The quadratic regularization is split between f and g through the parameters a_1 and a_2 . This decomposition preserves the original objective while allowing the proximal operators of both sub-problems to remain well conditioned. In particular,

- f consists of a quadratic term together with the box constraint $\mathbf{X} \geq -\mathbf{A}$. Its proximal operator is completely separable over the matrix entries and reduces to an elementwise projection.
- g combines the squared nuclear norm and the linear orthogonality constraint $\mathbf{X}\mathbf{B} = \mathbf{0}$. Although this proximal operator is more involved, the orthogonality constraint has a simple linear structure, allowing a closed-form solution to be derived through the optimality conditions.

Since both f and g are proper, closed, and convex, Douglas–Rachford splitting is applicable. Instead of solving one difficult proximal problem involving

$$\text{prox}_{\eta^{-1}}(f+g),$$

the algorithm alternates between the much simpler proximal operators

$\text{prox}_{\eta^{-1}f} \quad \text{and} \quad \text{prox}_{\eta^{-1}g},$

while still converging to a minimizer of the original convex problem.

Algorithm 1: DR for convex problem

```

Input:  $\mathbf{Z}_1, \mathbf{A}, \mathbf{B}, a_1, a_2 \geq 0, \eta > 0, \alpha \in (0, 1]$ , maxit
Output:  $\mathbf{X}_{\text{maxit}}$ 
1 for  $k = 1 : \text{maxit}$  do
2    $\mathbf{X}_k \leftarrow \arg \min_{\mathbf{X}} f(\mathbf{X}) + \frac{\eta}{2} \|\mathbf{X} - \mathbf{Z}_k\|_F^2 ;$            // first proximal operator
3    $\mathbf{Y}_k \leftarrow \arg \min_{\mathbf{Y}} g(\mathbf{Y}) + \frac{\eta}{2} \|\mathbf{Y} - 2\mathbf{X}_k + \mathbf{Z}_k\|_F^2 ;$  // second proximal operator
4    $\mathbf{Z}_{k+1} \leftarrow \mathbf{Z}_k + 2\alpha (\mathbf{Y}_k - \mathbf{X}_k);$ 
5 end

```

The proximal operator of f has a simple closed form. Since f is a quadratic plus an indicator function for the set $\{\mathbf{X} \mid \mathbf{X} \geqslant -\mathbf{A}\}$, the problem is separable and the proximal operator is an elementwise clipping:

$$\arg \min_{\mathbf{X}} f(\mathbf{X}) + \frac{\eta}{2} \|\mathbf{X} - \mathbf{Z}_k\|_F^2 = \pi_+ \left(\frac{\eta}{a_1 + \eta} \mathbf{Z}_k \right), \quad (20)$$

where

$$\pi_+(\mathbf{Z})_{ij} = \max \{-A_{ij}, Z_{ij}\}. \quad (21)$$

4.2.2 The proximal operator of the orthogonality-constrained squared nuclear norm

Unlike the proximal operator of f , which is completely separable over the matrix entries, the proximal operator of g is substantially more challenging. The difficulty arises from two sources. First, the squared nuclear norm

$$\|\mathbf{Y}\|_*^2 = \left(\sum_i \sigma_i(\mathbf{Y}) \right)^2$$

couples all singular values, so the optimization problem is no longer separable. Second, the orthogonality constraint

$$\mathbf{Y}\mathbf{B} = \mathbf{0}$$

introduces a linear constraint on the admissible right singular vectors. Consequently, unlike the proximal operator of f , no elementwise solution is available. Our strategy is therefore different. Rather than solving the minimization problem directly, we first derive its first-order optimality condition, identify its structure, and then show that this condition is equivalent to a modified singular-value thresholding problem. This ultimately leads to a closed-form characterization of the proximal operator.

The optimality condition for

$$\mathbf{Y} = \arg \min_{\mathbf{Y}} g(\mathbf{Y}) + \frac{\eta}{2} \|\mathbf{Y} - \mathbf{M}\|_F^2 \quad (22)$$

is

$$\begin{cases} \mathbf{0} \in \partial \|\mathbf{Y}\|_*^2 + (a_2 + \eta)\mathbf{Y} - \eta\mathbf{M} + \{\mathbf{\Lambda}\mathbf{B}^T \mid \mathbf{\Lambda} \in \mathbb{R}^{N_x \times 3}\}, \\ \mathbf{Y}\mathbf{B} = \mathbf{0}. \end{cases} \quad (23)$$

We first state a lemma for the subgradient of the squared nuclear norm.

Lemma 1. $2\|\mathbf{X}\|_* \partial \|\mathbf{X}\|_* \subset \partial \|\mathbf{X}\|_*^2$.

Proof. For each $\mathbf{Z} \in \partial \|\mathbf{X}\|_*$, we have $\|\mathbf{Y}\|_* - \|\mathbf{X}\|_* \geq \langle \mathbf{Z}, \mathbf{Y} - \mathbf{X} \rangle$ for all $\mathbf{Y}$. Then

$$\begin{aligned} \|\mathbf{Y}\|_*^2 - \|\mathbf{X}\|_*^2 - 2\|\mathbf{X}\|_* \langle \mathbf{Z}, \mathbf{Y} - \mathbf{X} \rangle &\geq \|\mathbf{Y}\|_*^2 - \|\mathbf{X}\|_*^2 - 2\|\mathbf{X}\|_* (\|\mathbf{Y}\|_* - \|\mathbf{X}\|_*) \\ &= (\|\mathbf{Y}\|_* - \|\mathbf{X}\|_*)^2 \geq 0. \end{aligned}$$

Therefore, $2\|\mathbf{X}\|_* \mathbf{Z} \in \partial \|\mathbf{X}\|_*^2$. □

Definition 2. For any matrix $\mathbf{O} \in \mathbb{R}^{n \times k}$ ($k \leq n$) with orthonormal columns, we define the orthogonal projection and its complement as

$$\Pi_{\mathbf{O}} = \mathbf{O}\mathbf{O}^T, \quad \Pi_{\mathbf{O}}^\perp = \mathbf{I} - \mathbf{O}\mathbf{O}^T. \quad (24)$$

The next theorem characterizes the proximal operator of g . It shows that the proximal operator can be expressed in terms of the singular value soft thresholding operator $D_\tau(\cdot)$, which subtracts τ from each singular value and discards the non-positive ones.

Theorem 1. Suppose the compact SVD of the minimizer $\mathbf{Y}$ of (22) is $\mathbf{Y} = \mathbf{U}\mathbf{\Sigma}\mathbf{V}^T$. Then $\mathbf{Y}$ satisfies

$$\mathbf{Y} = D_{\frac{2\|\mathbf{Y}\|_*}{a_2 + \eta}} \left(\frac{\eta}{a_2 + \eta} \mathbf{M} \Pi_{\mathbf{B}}^\perp \right), \quad (25)$$

where $D_\tau(\cdot)$ is the singular value soft thresholding operator with threshold τ .

Proof. Since g is strongly convex, problem (22) has a unique minimizer. We construct a candidate $\mathbf{Y}$ via (25) and verify that it satisfies the first-order optimality condition (23). By Lemma 1 and Proposition 1, $\mathbf{Y}$ satisfies (23) if the following system admits a solution:

$$\begin{aligned} \exists \mathbf{Z} \in \mathbb{R}^{N_x \times 3}, \mathbf{W} \in \{ \mathbf{W} \in \mathbb{R}^{N_x \times N_v} \mid \mathbf{U}^T \mathbf{W} = \mathbf{0}, \mathbf{W} \mathbf{V} = \mathbf{0}, \|\mathbf{W}\|_2 \leq 1 \} : \\ \begin{cases} \mathbf{0} = 2 \|\boldsymbol{\Sigma}\|_* \mathbf{U} \mathbf{V}^T + 2 \|\boldsymbol{\Sigma}\|_* \mathbf{W} + \mathbf{Z} \mathbf{B}^T + (a_2 + \eta) \mathbf{U} \boldsymbol{\Sigma} \mathbf{V}^T - \eta \mathbf{M}, \\ \mathbf{V}^T \mathbf{B} = \mathbf{0}, \end{cases} \end{aligned} \quad (26)$$

which can be rewritten as

$$\begin{aligned} \exists \mathbf{Z} \in \mathbb{R}^{N_x \times 3}, \mathbf{W} \in \{ \mathbf{W} \in \mathbb{R}^{N_x \times N_v} \mid \mathbf{U}^T \mathbf{W} = \mathbf{0}, \mathbf{W} \mathbf{V} = \mathbf{0}, \|\mathbf{W}\|_2 \leq 1 \} : \\ \begin{cases} \frac{\eta}{a_2 + \eta} \mathbf{M} = \mathbf{U} \left(\boldsymbol{\Sigma} + \frac{2 \|\boldsymbol{\Sigma}\|_*}{a_2 + \eta} \right) \mathbf{V}^T + \frac{2 \|\boldsymbol{\Sigma}\|_*}{a_2 + \eta} \mathbf{W} + \mathbf{Z} \mathbf{B}^T, \\ \mathbf{V}^T \mathbf{B} = \mathbf{0}. \end{cases} \end{aligned} \quad (27)$$

We now show that this system admits a solution. This proves that the candidate $\mathbf{Y}$ constructed via (25) satisfies the optimality condition and is therefore the unique minimizer.

Let $\tau = \frac{2 \|\boldsymbol{\Sigma}\|_*}{a_2 + \eta}$, then

$$\begin{aligned} \frac{\eta}{a_2 + \eta} \mathbf{M} &= \mathbf{U} (\boldsymbol{\Sigma} + \tau) \mathbf{V}^T + \tau \mathbf{W} + \mathbf{Z} \mathbf{B}^T \\ &= \mathbf{U} (\boldsymbol{\Sigma} + \tau) \mathbf{V}^T + \tau \mathbf{W} (\mathbf{I} - \mathbf{B} \mathbf{B}^T) + (\mathbf{Z} + \tau \mathbf{W} \mathbf{B}) \mathbf{B}^T \\ &= \begin{pmatrix} \mathbf{U} & \tilde{\mathbf{U}} \end{pmatrix} \begin{pmatrix} \boldsymbol{\Sigma} + \tau & \mathbf{0} \\ \mathbf{0} & \tau \tilde{\boldsymbol{\Sigma}} \end{pmatrix} \begin{pmatrix} \mathbf{V}^T \\ \tilde{\mathbf{V}}^T \end{pmatrix} + \tilde{\mathbf{Z}} \mathbf{B}^T. \end{aligned} \quad (28)$$

Here, $\tilde{\mathbf{Z}} = \mathbf{Z} + \tau \mathbf{W} \mathbf{B} \in \mathbb{R}^{N_x \times 3}$ is arbitrary. The matrix $\tilde{\mathbf{W}} = \mathbf{W} \boldsymbol{\Pi}_B^\perp$, with compact SVD $\tilde{\mathbf{W}} = \tilde{\mathbf{U}} \tilde{\boldsymbol{\Sigma}} \tilde{\mathbf{V}}^T$, satisfies

$$\mathbf{U}^T \mathbf{W} = \mathbf{0} \implies \mathbf{U}^T \tilde{\mathbf{U}} = \mathbf{0}, \quad (29a)$$

$$\mathbf{W} \mathbf{V} = \mathbf{0} \implies \tilde{\mathbf{V}}^T \mathbf{V} = \mathbf{0}, \quad (29b)$$

$$\boldsymbol{\Pi}_B^\perp \mathbf{B} = \mathbf{0} \implies \tilde{\mathbf{V}}^T \mathbf{B} = \mathbf{0}, \quad (29c)$$

$$\|\mathbf{W}^T \mathbf{x}\|_2^2 = \|(\mathbf{I} - \mathbf{B} \mathbf{B}^T) \mathbf{W}^T \mathbf{x}\|_2^2 + \|\mathbf{B} \mathbf{B}^T \mathbf{W}^T \mathbf{x}\|_2^2 \geq \|\tilde{\mathbf{W}}^T \mathbf{x}\|_2^2 \implies \|\tilde{\mathbf{W}}\|_2 \leq 1. \quad (29d)$$

Therefore, such a decomposition exists. The equation

$$\frac{\eta}{a_2 + \eta} \mathbf{M} \boldsymbol{\Pi}_B^\perp = \begin{pmatrix} \mathbf{U} & \tilde{\mathbf{U}} \end{pmatrix} \begin{pmatrix} \boldsymbol{\Sigma} + \tau & \mathbf{0} \\ \mathbf{0} & \tau \tilde{\boldsymbol{\Sigma}} \end{pmatrix} \begin{pmatrix} \mathbf{V}^T \\ \tilde{\mathbf{V}}^T \end{pmatrix} \quad (30)$$

gives a compact SVD. Using the fact that $\|\tilde{\boldsymbol{\Sigma}}\|_2 = \|\tilde{\mathbf{W}}\|_2 \leq 1$, we arrive at equation (25). $\square$

Remark 3 (Efficient computation of the threshold). *Although equation (25) defines $\mathbf{Y}$ implicitly, the threshold τ can be computed efficiently. Suppose we have computed the SVD $\frac{\eta}{a_2 + \eta} \mathbf{M} (\mathbf{I} - \mathbf{B} \mathbf{B}^T) = \hat{\mathbf{U}} \hat{\boldsymbol{\Sigma}} \hat{\mathbf{V}}^T$ with $\hat{\boldsymbol{\Sigma}} = \text{diag}(\sigma_1, \dots, \sigma_n)$. Then τ satisfies*

$$0 = (a_2 + \eta) \tau - 2 \sum_{k=1}^n \max \{ \sigma_k - \tau, 0 \}. \quad (31)$$

The right-hand side is a convex, monotonically increasing, piecewise linear function of τ with break-points at $\sigma_1, \dots, \sigma_n$. The root can be located by binary search over the breakpoints in $\mathcal{O}(\log n)$ comparisons, followed by an explicit solve within the linear piece, for a total cost of $\mathcal{O}(n \log n)$ floating point operations (dominated by sorting the singular values). The SVD itself is only computed once and is the dominant computational cost. A partial SVD, computing only the singular values larger than τ , could be used, but may need to be later expanded since the value of τ is not known in advance, so it is not particularly efficient.

4.3 Restarted dual FISTA for the convex problem

The Douglas–Rachford splitting works directly on the primal problem, but requires computing the proximal operator of g at each iteration, which involves a full SVD (cf. Theorem 1 and Remark 3). An alternative is to work on the dual problem, where the structure of the conjugate functions may lead to cheaper iterations. Since f and g are both strongly convex (with parameters a_1 and a_2 respectively), their convex conjugates f^* and g^* have Lipschitz continuous gradients, making the dual problem amenable to gradient-based methods.

The FISTA algorithm [2] uses a Nesterov-type [21] momentum method to accelerate the convergence of proximal gradient descent from $\mathcal{O}(1/k)$ to $\mathcal{O}(1/k^2)$, where k is the iteration number. FISTA can be further accelerated by adaptive restart strategies [23, 1]. The approach of applying restarted FISTA to a dual problem for enforcing nonnegativity constraints was studied in [6]. We apply the same idea to the dual problem of the splitting formulation (19),

$$\begin{aligned} \min_{\mathbf{\Lambda}} F(\mathbf{\Lambda}), \\ \text{where } F(\mathbf{\Lambda}) = \{f^*(-\mathbf{\Lambda}) + g^*(\mathbf{\Lambda})\}, \end{aligned} \quad (32)$$

where f^* and g^* are convex conjugates of f and g respectively, and $\mathbf{\Lambda} \in \mathbb{R}^{m \times n}$ is the dual variable. Here, we require $a_1 > 0$ and $a_2 > 0$.

Since f is a_1 -strongly convex and g is a_2 -strongly convex, the conjugate functions f^* and g^* are differentiable with $\frac{1}{a_1}$ -Lipschitz and $\frac{1}{a_2}$ -Lipschitz gradients, respectively. In particular, the dual cost function $F(\mathbf{\Lambda}) = f^*(-\mathbf{\Lambda}) + g^*(\mathbf{\Lambda})$ takes finite values and has Lipschitz continuous gradients.

We first write down f^* :

$$\begin{aligned} f^*(\mathbf{\Lambda}) &= \sup_{\mathbf{X} \geq -\mathbf{A}} \langle \mathbf{\Lambda}, \mathbf{X} \rangle - \frac{a_1}{2} \|\mathbf{X}\|_F^2 \\ &= \left(\langle \mathbf{\Lambda}, \mathbf{X} \rangle - \frac{a_1}{2} \|\mathbf{X}\|_F^2 \right) \Big|_{\mathbf{X} = \max\{-\mathbf{A}, a_1^{-1} \mathbf{\Lambda}\}}, \end{aligned} \quad (33)$$

since the problem is separable, and using a similar procedure as the derivation of $\text{prox}_{\eta^{-1}g}$, we have

$$g^*(\mathbf{\Lambda}) = \left(\langle \tilde{\boldsymbol{\sigma}}, \boldsymbol{\sigma} \rangle - \|\boldsymbol{\sigma}\|_{\ell^1}^2 - \frac{a_2}{2} \|\boldsymbol{\sigma}\|_{\ell^2}^2 \right) \Big|_{\boldsymbol{\sigma} = D_{\frac{2\|\boldsymbol{\sigma}\|_{\ell^1}}{a_2}}(\frac{\tilde{\boldsymbol{\sigma}}}{a_2})}, \quad (34)$$

where $\tilde{\boldsymbol{\sigma}}$ is the singular value vector of $\mathbf{\Lambda} \Pi_{\mathbf{B}}^\perp$ and $\|\boldsymbol{\sigma}\|_{\ell^1}$ is the sum over all the absolute values of the entries of $\boldsymbol{\sigma}$.

Since both f^* and g^* are smooth, the proximal-gradient (FISTA) iteration can be set up in two symmetric ways: take the gradient of one conjugate and the proximal operator of the other. We describe both and call them *variant I* (gradient of f^* , proximal operator of g^*) and *variant II* (gradient of g^* , proximal operator of f^*).

Variante I (gradient of f^* , proximal operator of g^*). This variant requires the proximal operator of g^* . Since g is a_2 -strongly convex, g^* has $\frac{1}{a_2}$ -Lipschitz gradient, and $\text{prox}_{\eta_k g^*}$ can be evaluated using the Moreau decomposition $\text{prox}_{\eta_k g^*}(\mathbf{Z}) = \mathbf{Z} - \eta_k \text{prox}_{\eta_k^{-1} g}(\eta_k^{-1} \mathbf{Z})$, where $\text{prox}_{\eta_k^{-1} g}$ is the proximal operator computed in Theorem 1. The iteration is

$$\mathbf{\Lambda}_{k+1} = \text{prox}_{\eta_k g^*}(\mathbf{\Theta}_k + \eta_k \nabla f^*(-\mathbf{\Theta}_k)), \quad (35a)$$

$$t_{k+1} = \frac{1 + \sqrt{1 + 4t_k^2}}{2}, \quad (35b)$$

$$\mathbf{\Theta}_{k+1} = \mathbf{\Lambda}_{k+1} + \frac{t_k - 1}{t_{k+1}}(\mathbf{\Lambda}_{k+1} - \mathbf{\Lambda}_k), \quad (35c)$$

where $\mathbf{\Lambda}_k$ converges fast to $\mathbf{\Lambda}$, and the initial condition is taken to be $\mathbf{\Theta}_0 = \mathbf{\Lambda}_0$ and $t_0 = 1$. We require $\eta_k \leq \frac{1}{L_{f^*}} = a_1$ for stability.

Since this algorithm operates in the dual space, we need to recover the primal variable $\mathbf{X}$ from $\mathbf{\Lambda}$. We use $\mathbf{X} = \nabla f^*(-\mathbf{\Lambda})$, which always satisfies the nonnegativity condition $\mathbf{X} \geq -\mathbf{A}$ (but not necessarily the orthogonality constraint since the iteration never *exactly* reaches optimality). By the Fenchel–Young equality,

$$\langle \mathbf{\Lambda}, \mathbf{X} \rangle = f(\mathbf{X}) + f^*(\mathbf{\Lambda}) \iff \mathbf{X} \in \partial f^*(\mathbf{\Lambda}),$$

the primal variable is given by

$$\mathbf{X} = \nabla f^*(-\mathbf{\Lambda}) = \max\{-\mathbf{A}, -a_1^{-1} \mathbf{\Lambda}\}. \quad (36)$$

Variante II (gradient of g^* , proximal operator of f^*). Here we instead take the gradient of g^* and the proximal operator of f^* . The iteration is

$$\mathbf{\Lambda}_{k+1} = \text{prox}_{\eta_k [f^*(-\cdot)]}(\mathbf{\Theta}_k - \eta_k \nabla g^*(\mathbf{\Theta}_k)), \quad (37a)$$

$$t_{k+1} = \frac{1}{2} \left(1 + \sqrt{1 + 4t_k^2} \right), \quad \mathbf{\Theta}_{k+1} = \mathbf{\Lambda}_{k+1} + \frac{t_k - 1}{t_{k+1}}(\mathbf{\Lambda}_{k+1} - \mathbf{\Lambda}_k), \quad (37b)$$

with step size $\eta_k \leq \frac{1}{L_{g^*}} = a_2$. The gradient $\nabla g^*(\mathbf{\Lambda}) = D_{\frac{2\|\mathbf{X}\|_*}{a_2}}(a_2^{-1} \mathbf{\Lambda} \mathbf{\Pi}_B^\perp)$ requires an SVD, while the proximal operator of f^* is now the cheap closed-form map. By the Moreau decomposition,

$$\text{prox}_{\eta [f^*(-\cdot)]}(\mathbf{Z}) = \mathbf{Z} + \eta \text{prox}_{\eta^{-1} f}(-\eta^{-1} \mathbf{Z}) = \mathbf{Z} + \eta \max\left\{-\mathbf{A}, -\frac{1}{a_1 + \eta} \mathbf{Z}\right\}. \quad (38)$$

The recovered primal variable is now $\mathbf{X} = \nabla g^*(\mathbf{\Lambda}) = D_{\frac{2\|\mathbf{X}\|_*}{a_2}}(a_2^{-1} \mathbf{\Lambda} \mathbf{\Pi}_B^\perp)$, which *exactly* satisfies the orthogonality constraint $\mathbf{X} \mathbf{B} = \mathbf{0}$ and is low-rank, but is not necessarily nonnegative until convergence (the mirror image of variant I, whose iterate is always nonnegative but only orthogonal at convergence).

Remark 4 (Which variant to use). *Both variants converge to the same global minimizer of (2) and, as Table 1 shows, have comparable per-iteration cost (each needs one SVD per iteration) and nearly the same convergence rate, with variant I slightly faster and cheaper (its primal recovery (36) is an elementwise clip rather than an SVD). They differ in which constraint the iterate satisfies exactly before convergence: variant I keeps the nonnegativity $\mathbf{A} + \mathbf{X} \geq \mathbf{0}$, whereas variant II keeps the*

orthogonality $\mathbf{X}\mathbf{B} = \mathbf{0}$. Since the entire purpose of the correction is to recover nonnegativity, we recommend variant I: it returns a nonnegative iterate at every step and can therefore be stopped early while still yielding a physically admissible (nonnegative) solution, whereas variant II may return a solution with small negative entries until it has fully converged.

To eliminate the oscillation caused by the momentum term, restart strategies are employed. For strongly convex problems, a fixed restart with an optimal period produces linear convergence [1]. Since our dual problem (32) is not strongly convex, we also consider adaptive restart methods [22, 23]. The function-value-based and gradient-based restart criteria are

$$F(\mathbf{\Lambda}_k) > F(\mathbf{\Lambda}_{k-1}), \quad (39)$$

and

$$\langle \mathbf{G}(\mathbf{\Theta}_{k-1}), \mathbf{\Lambda}_k - \mathbf{\Lambda}_{k-1} \rangle > 0, \quad (40a)$$

where

$$\mathbf{G}(\mathbf{\Theta}_{k-1}) = \frac{1}{\eta_k} (\mathbf{\Theta}_{k-1} - \mathbf{\Lambda}_k). \quad (40b)$$

4.4 Restarted dual accelerated gradient descent for the convex problem

If both a_1 and a_2 are positive, the dual cost function F in (32) has Lipschitz continuous gradient with Lipschitz constant $L = \frac{1}{a_1} + \frac{1}{a_2}$. In this case, we use a fully explicit accelerated gradient descent method without proximal operators. We apply Nesterov's accelerated gradient descent [21] to the dual problem and use the same restarting strategies as in the previous subsection.

The gradient of the dual cost function F is

$$\nabla F(\mathbf{\Lambda}) = -\nabla f^*(-\mathbf{\Lambda}) + \nabla g^*(\mathbf{\Lambda}), \quad (41)$$

where ∇f^* is given by (36) and ∇g^* is given by

$$\mathbf{X} = \nabla g^*(\mathbf{\Lambda}) = D_{\frac{2\|\mathbf{X}\|_*}{a_2}} \left(\frac{1}{a_2} \mathbf{\Lambda} \mathbf{\Pi}_B^\perp \right). \quad (42)$$

The accelerated gradient iteration is

$$\mathbf{\Lambda}_{k+1} = \mathbf{\Theta}_k - \eta_k (\nabla g^*(\mathbf{\Theta}_k) - \nabla f^*(-\mathbf{\Theta}_k)), \quad (43)$$

$$t_{k+1} = \frac{1 + \sqrt{1 + 4t_k^2}}{2}, \quad (44)$$

$$\mathbf{\Theta}_{k+1} = \mathbf{\Lambda}_{k+1} + \frac{t_k - 1}{t_{k+1}} (\mathbf{\Lambda}_{k+1} - \mathbf{\Lambda}_k), \quad (45)$$

where $\mathbf{\Lambda}_k$ converges to $\mathbf{\Lambda}$, and the initial condition is $\mathbf{\Theta}_0 = \mathbf{\Lambda}_0$ and $t_0 = 1$. We require $\eta_k \leq \frac{1}{L_F} = \frac{a_1 a_2}{a_1 + a_2}$ for stability. Similar to the FISTA, our restarting strategies include fixed restart, function-value-based restart (39), and gradient-based restart (40).

4.5 Dual PR+ conjugate gradient for the convex problem

We next apply the nonlinear conjugate gradient method to the dual problem (32). We use the Polak–Ribière plus (PR+) variant [25] with an explicit step size formula from [27] to avoid line search. The algorithm is

$$\mathbf{G}_k = \nabla F(\mathbf{\Lambda}_k), \quad (46)$$

$$\beta_k = \begin{cases} 0, & k = 1, \\ \frac{\langle \mathbf{G}_k, \mathbf{G}_k - \mathbf{G}_{k-1} \rangle}{\|\mathbf{G}_{k-1}\|_F^2}, & k > 1, \end{cases} \quad (47)$$

$$\mathbf{D}_k = \beta_k \mathbf{D}_{k-1} - \mathbf{G}_k, \quad (48)$$

$$\alpha_k = -\frac{\delta \langle \mathbf{G}_k, \mathbf{D}_k \rangle}{\|\mathbf{D}_k\|_F^2} \quad (49)$$

$$\mathbf{\Lambda}_{k+1} = \mathbf{\Lambda}_k + \alpha_k \mathbf{D}_k. \quad (50)$$

Here, according to [27], we should require $\delta \in (0, \frac{1}{L_F})$ for stability.

4.6 Dual L-BFGS for the convex problem

The L-BFGS algorithm [19] is a quasi-Newton method that uses a limited number of past iterates to approximate the inverse Hessian, requiring only gradient evaluations. Applied to the dual problem (32), the algorithm is

$$\mathbf{G}_k = \nabla F(\mathbf{\Lambda}_k), \quad (51)$$

$$\mathbf{D}_k = -\mathbf{H}_k \mathbf{G}_k, \quad (52)$$

$$\alpha_k = \arg \min_{\alpha \geq 0} F(\mathbf{\Lambda}_k + \alpha \mathbf{D}_k), \quad (53)$$

$$\mathbf{\Lambda}_{k+1} = \mathbf{\Lambda}_k + \alpha_k \mathbf{D}_k, \quad (54)$$

$$\begin{aligned} \mathbf{H}_{k+1} = & \left(\prod_{j=k-p+1}^k \mathbf{J}_j \right)^T \mathbf{H}_0^k \prod_{j=k-p+1}^k \mathbf{J}_j \\ & + \sum_{i=1}^p \gamma_{k-p+i} \left(\prod_{j=k-p+i+1}^k \mathbf{J}_j \right)^T \mathbf{s}_{k-p+i} \mathbf{s}_{k-p+i}^T \prod_{j=k-p+i+1}^k \mathbf{J}_j. \end{aligned} \quad (55)$$

Here, the iterates are assumed to be vectors, $\mathbf{J}_k = \mathbf{I} - \gamma_k \mathbf{y}_k \mathbf{s}_k^T$, $\mathbf{y}_j = \mathbf{G}_{j+1} - \mathbf{G}_j$, $\mathbf{s}_j = \mathbf{\Lambda}_{j+1} - \mathbf{\Lambda}_j$, $\gamma_j = \frac{1}{\mathbf{y}_j^T \mathbf{s}_j}$, and p is the number of history iterates to be stored. The matrix-vector product of $\mathbf{H}_k$ can be efficiently implemented through two iterations presented in [19]. The exact line search can be replaced by a Wolfe-type inexact search.

4.7 Tangent-space accelerated alternating projection for the non-convex problem

For the rank-constrained formulation (17), we develop an alternating projection algorithm that uses the manifold structure to avoid full SVD computations. We first introduce the relevant manifold and its projection properties.

Definition 3. Let $M_{r,B} = \{\mathbf{X} \in \mathbb{R}^{m \times n} \mid \text{rank}(\mathbf{X}) \leq r, \mathbf{X}\mathbf{B} = \mathbf{0}\}$ denote the set of rank-at-most- r matrices satisfying the orthogonality constraint.

Lemma 2. The Frobenius-norm projection onto $M_{r,B}$ is

$$P_{M_{r,B}}(\mathbf{Y}) = P_{M_r}(\mathbf{Y}\mathbf{\Pi}_B^\perp), \quad (56)$$

where $P_{M_r}(\cdot)$ is the standard rank- r SVD truncation.

Proof. Let $\mathbf{Y}\mathbf{\Pi}_B^\perp = \mathbf{U}\mathbf{\Sigma}\mathbf{V}^T$ be a full SVD. For any $\mathbf{Z} \in M_{r,B}$,

$$\begin{aligned} \|\mathbf{Y} - \mathbf{Z}\|_F^2 &= \|\mathbf{Y}\mathbf{\Pi}_B\|_F^2 + \|\mathbf{Y}\mathbf{\Pi}_B^\perp - \mathbf{Z}\|_F^2 \\ &= \|\mathbf{Y}\mathbf{\Pi}_B\|_F^2 + \|\mathbf{Y}\mathbf{\Pi}_B^\perp\|_F^2 - 2\text{tr}(\mathbf{Z}^T \mathbf{Y}\mathbf{\Pi}_B^\perp) + \|\mathbf{Z}\|_F^2 \\ &= \|\mathbf{Y}\mathbf{\Pi}_B\|_F^2 + \left\| \mathbf{\Sigma} - \widehat{\mathbf{Z}} \right\|_F^2 \\ &\geq \|\mathbf{Y}\mathbf{\Pi}_B\|_F^2 + \sum_{k=r+1}^n \Sigma_{k,k}^2, \end{aligned} \quad (57)$$

where $\widehat{\mathbf{Z}} = \mathbf{U}^T \mathbf{Z} \mathbf{V}$ has rank at most r . The inequality is the Eckart–Young–Mirsky theorem: since $\Sigma_{k,k}$ are the singular values of $\mathbf{\Sigma}$, the best rank- r Frobenius approximation retains the largest r of them, with equality if and only if $\widehat{\mathbf{Z}} = P_{M_r}(\mathbf{\Sigma})$, i.e. $\mathbf{Z} = P_{M_r}(\mathbf{Y}\mathbf{\Pi}_B^\perp)$. It is straightforward to check that $\mathbf{Z} = P_{M_r}(\mathbf{Y}\mathbf{\Pi}_B^\perp) \in M_{r,B}$. $\square$

We first consider the alternating projection algorithm that projects onto the nonnegative set and onto $M_{r,B}$, as shown in Algorithm 2.

Algorithm 2: Alternating projection (AP) for the non-convex problem

Input: $\mathbf{X}_0, \mathbf{A}, \mathbf{B}, r, \text{maxit}$
Output: $\mathbf{X}_{\text{maxit}+1}$
1 **for** $k = 1 : \text{maxit}$ **do**
2 $\mathbf{Y}_k = \pi_+(\mathbf{X}_k)$; // element-wise
3 $\mathbf{X}_{k+1} = P_{M_{r,B}}(\mathbf{Y}_k)$; // requires SVD
4 **end**

The expensive step is the projection $P_{M_{r,B}}$, which requires a full SVD. Following [26], we replace this projection by first projecting onto the tangent space of $M_{r,B}$ at the current iterate $\mathbf{X}_k$ and then truncating to rank r . This reduces the SVD cost from a full $m \times n$ SVD to a $2r \times 2r$ SVD. We now introduce the tangent-space machinery. In the following, we assume that the iterates $\mathbf{X}_k$ have exact rank r , so that the tangent space of the fixed-rank manifold is well-defined. This assumption holds generically in practice, especially since we choose r to be small.

Lemma 3. The tangent space of the fixed-rank manifold M_r at $\mathbf{X} = \mathbf{U}\mathbf{\Sigma}\mathbf{V}^T \in M_r$ (compact SVD

form, with Σ having all positive diagonal entries) is

$$\begin{aligned}
T_{M_r}(\mathbf{X}) &= \{\mathbf{U}\mathbf{W}^T + \mathbf{Z}\mathbf{V}^T \mid \mathbf{W} \in \mathbb{R}^{n \times r}, \mathbf{Z} \in \mathbb{R}^{m \times r}\} \\
&= \{\mathbf{U}\mathbf{M}\mathbf{V}^T + \mathbf{U}_p\mathbf{V}^T + \mathbf{U}\mathbf{V}_p^T \mid \mathbf{U}_p^T\mathbf{U} = \mathbf{0}, \mathbf{V}_p^T\mathbf{V} = \mathbf{0}, \mathbf{M} \in \mathbb{R}^{r \times r}, \mathbf{U}_p \in \mathbb{R}^{m \times r}, \mathbf{V}_p \in \mathbb{R}^{n \times r}\} \quad (58) \\
&= \left\{ \begin{pmatrix} \mathbf{U} & \mathbf{U}^\perp \end{pmatrix} \begin{pmatrix} \mathbf{C}_1 & \mathbf{C}_2 \\ \mathbf{C}_3 & \mathbf{0} \end{pmatrix} \begin{pmatrix} \mathbf{V} & \mathbf{V}^\perp \end{pmatrix}^T \mid \mathbf{C}_1 \in \mathbb{R}^{r \times r}, \mathbf{C}_2 \in \mathbb{R}^{r \times (n-r)}, \mathbf{C}_3 \in \mathbb{R}^{(m-r) \times r} \right\}.
\end{aligned}$$

Proof. See [35] and [29]. $\square$

Lemma 4. The least square projection (in Frobenius norm) onto $T_{M_r}(\mathbf{X})$ is

$$P_{T_{M_r}(\mathbf{X})}(\mathbf{Y}) = \Pi_{\mathbf{U}}\mathbf{Y} + \mathbf{Y}\Pi_{\mathbf{V}} - \Pi_{\mathbf{U}}\mathbf{Y}\Pi_{\mathbf{V}}, \quad (59)$$

where $\mathbf{X} = \mathbf{U}\Sigma\mathbf{V}^T$ is a rank- r SVD.

Proof. See [35] and [29]. $\square$

Note that these computations can be done efficiently if the matrix multiplication is done in the right order; e.g., $\Pi_{\mathbf{U}}\mathbf{Y} = \mathbf{U}\mathbf{U}^T\mathbf{Y}$ is calculated as $\mathbf{U}(\mathbf{U}^T\mathbf{Y})$, and the final outer multiplication is not done explicitly but left in factored form (see Remark 5).

Lemma 5. The tangent space of $M_{r,\mathbf{B}}$ at $\mathbf{X} = \mathbf{U}\Sigma\mathbf{V}^T \in M_{r,\mathbf{B}}$ (compact SVD form) is

$$\begin{aligned}
T_{M_{r,\mathbf{B}}}(\mathbf{X}) &= \{\mathbf{U}\mathbf{W}^T + \mathbf{Z}\mathbf{V}^T \mid \mathbf{W} \in \mathbb{R}^{n \times r}, \mathbf{Z} \in \mathbb{R}^{m \times r}, \mathbf{W}^T\mathbf{B} = \mathbf{0}\} \\
&= \{\mathbf{U}\mathbf{M}\mathbf{V}^T + \mathbf{U}_p\mathbf{V}^T + \mathbf{U}\mathbf{V}_p^T \Pi_{\mathbf{B}}^\perp \mid \mathbf{U}_p^T\mathbf{U} = \mathbf{0}, \mathbf{V}_p^T\mathbf{V} = \mathbf{0}, \mathbf{M} \in \mathbb{R}^{r \times r}, \mathbf{U}_p \in \mathbb{R}^{m \times r}, \mathbf{V}_p \in \mathbb{R}^{n \times r}\}. \quad (60)
\end{aligned}$$

Proof. Since $\mathbf{X} \in M_{r,\mathbf{B}}$, we have $\mathbf{V}^T\mathbf{B} = \mathbf{0}$. The orthogonality constraint $\mathbf{W}^T\mathbf{B} = \mathbf{0}$ follows from differentiating $\mathbf{X}\mathbf{B} = \mathbf{0}$ along the manifold. $\square$

Lemma 6. The least square projection (in Frobenius norm) onto $T_{M_{r,\mathbf{B}}}(\mathbf{X})$ is

$$P_{T_{M_{r,\mathbf{B}}}(\mathbf{X})}(\mathbf{Y}) = P_{T_{M_r}(\mathbf{X})}(\mathbf{Y})\Pi_{\mathbf{B}}^\perp = P_{T_{M_r}(\mathbf{X})}(\mathbf{Y}\Pi_{\mathbf{B}}^\perp). \quad (61)$$

Proof. Since $\mathbf{X} \in M_{r,\mathbf{B}}$, we have $\mathbf{V}^T\mathbf{B} = \mathbf{0}$. As a result, there exists a matrix $\mathbf{K}$ such that the matrix $\begin{pmatrix} \mathbf{V} & \mathbf{B} & \mathbf{K} \end{pmatrix}$ is orthonormal. Using matrix $\mathbf{K}$, we can rewrite the tangent space as

$$\begin{aligned}
T_{M_{r,\mathbf{B}}}(\mathbf{X}) &= \left\{ \begin{pmatrix} \mathbf{U} & \mathbf{U}^\perp \end{pmatrix} \begin{pmatrix} \mathbf{C}_1 & \mathbf{0} & \mathbf{C}_2 \\ \mathbf{C}_3 & \mathbf{0} & \mathbf{0} \end{pmatrix} \begin{pmatrix} \mathbf{V} & \mathbf{B} & \mathbf{K} \end{pmatrix}^T \mid \mathbf{C}_1 \in \mathbb{R}^{r \times r}, \mathbf{C}_2 \in \mathbb{R}^{r \times (n-r-3)}, \mathbf{C}_3 \in \mathbb{R}^{(m-r) \times r} \right\}. \quad (62)
\end{aligned}$$

Then, the least square projection is given by

$$\begin{aligned}
P_{T_{M_{r,\mathbf{B}}}(\mathbf{X})}(\mathbf{Y}) &= \Pi_{\mathbf{U}}\mathbf{Y}\Pi_{\mathbf{V}} + \Pi_{\mathbf{U}}^\perp\mathbf{Y}\Pi_{\mathbf{V}} + \Pi_{\mathbf{U}}\mathbf{Y}\Pi_{\mathbf{K}} \\
&= \Pi_{\mathbf{U}}\mathbf{Y}\Pi_{\mathbf{B}}^\perp + \Pi_{\mathbf{U}}^\perp\mathbf{Y}\Pi_{\mathbf{V}} \\
&= \Pi_{\mathbf{U}}\mathbf{Y}\Pi_{\mathbf{B}}^\perp + \Pi_{\mathbf{U}}^\perp\mathbf{Y}\Pi_{\mathbf{V}}\Pi_{\mathbf{B}}^\perp = P_{T_{M_r}(\mathbf{X})}(\mathbf{Y})\Pi_{\mathbf{B}}^\perp, \\
&= \Pi_{\mathbf{U}}\mathbf{Y}\Pi_{\mathbf{B}}^\perp + \Pi_{\mathbf{U}}^\perp\mathbf{Y}\Pi_{\mathbf{B}}^\perp\Pi_{\mathbf{V}} = P_{T_{M_r}(\mathbf{X})}(\mathbf{Y}\Pi_{\mathbf{B}}^\perp). \quad (63)
\end{aligned}$$

□

Remark 5. Since $\text{rank}(P_{T_{M_r, B}}(\mathbf{X})(\mathbf{Y})) \leq 2r$, the tangent-space projection can be stored in a low-rank factored form.

Using these results, we obtain the tangent-space accelerated alternating projection (TAP) algorithm in Algorithm 3. At each iteration, we project $\mathbf{Y}_k$ onto the tangent space of $M_{r, B}$ at $\mathbf{X}_k$, and then truncate to rank r . The implementation of this combined step is given in Algorithm 4.

Algorithm 3: TAP for non-convex problem

Input: $\mathbf{X}_0, \mathbf{A}, \mathbf{B}, r, \text{maxit}$
Output: $\mathbf{X}_{\text{maxit}+1}$
1 **for** $k = 1 : \text{maxit}$ **do**
2 $\mathbf{Y}_k = \pi_+(\mathbf{X}_k)$;
3 $\mathbf{X}_{k+1} = P_{M_r} \left(P_{T_{M_r, B}}(\mathbf{X}_k)(\mathbf{Y}_k) \right)$; // rank r , and SVD also computed & stored
4 **end**

Algorithm 4: TP operator $P_{M_r} \left(P_{T_{M_r, B}}(\mathbf{X})(\mathbf{Y}) \right)$ for non-convex problem

Input: $\mathbf{X}_k = \mathbf{U}\mathbf{\Sigma}\mathbf{V}^T$ (rank- r compact SVD), $\mathbf{B}, \mathbf{Y}_k, r$
Output: $\mathbf{X}_{k+1}$ (in compact SVD form)
1 $\mathbf{Q}_1 \mathbf{R}_1 = \Pi_{\mathbf{U}}^\perp \mathbf{Y}_k \mathbf{V}$; // compact QR
2 $\mathbf{Q}_2 \mathbf{R}_2 = \Pi_{\mathbf{V}}^\perp \mathbf{Y}_k^T \mathbf{U}$; // compact QR
 // by orthogonality: $\mathbf{U}^T \mathbf{Q}_1 = \mathbf{0}$ and $\mathbf{V}^T \mathbf{Q}_2 = \mathbf{0}$
3 $\tilde{\mathbf{Q}} \tilde{\mathbf{R}} = \Pi_{\mathbf{B}}^\perp \begin{pmatrix} \mathbf{V} & \mathbf{Q}_2 \end{pmatrix}$; // compact QR
4 $\mathbf{E} = \begin{pmatrix} \mathbf{U}^T \mathbf{Y}_k \mathbf{V} & \mathbf{R}_2 \\ \mathbf{R}_1 & \mathbf{0} \end{pmatrix} \tilde{\mathbf{R}}^T$; // $\mathbf{E}$ is $2r \times 2r$
5 $\hat{\mathbf{U}} \hat{\mathbf{\Sigma}} \hat{\mathbf{V}}^T = \mathbf{E}$; // $2r \times 2r$ SVD
 // rank- $2r$ factorization: $P_{T_{M_r, B}}(\mathbf{X}_k)(\mathbf{Y}_k) = (\mathbf{U} \quad \mathbf{Q}_1) \hat{\mathbf{U}} \hat{\mathbf{\Sigma}} \hat{\mathbf{V}}^T \tilde{\mathbf{Q}}^T$
6 Truncate to rank r : keep only the top r singular triplets of $\hat{\mathbf{U}} \hat{\mathbf{\Sigma}} \hat{\mathbf{V}}^T$ to form $\mathbf{X}_{k+1}$

Remark 6 (Computational cost). All iterates $\mathbf{X}_k$ are stored in rank- r factored form $\mathbf{U}\mathbf{\Sigma}\mathbf{V}^T$. The computational cost per iteration of Algorithm 4 is $\mathcal{O}(mnr + (m+n)r^2 + r^3)$, which is significantly less than the $\mathcal{O}(mn \min(m, n))$ cost of the full SVD required by the algorithms for the convex formulation.

5 Numerical experiments

5.1 Test problem and setup

We test all algorithms on a 1D1V Landau damping problem discretized on an $N_x = 64, N_v = 128$ grid. The low-rank solution has $\text{rank}(\mathbf{A}_1) = 3$ and $\text{rank}(\mathbf{A}_2) = 26$. The input data are shown

in Figure 1. The scaled matrix $\mathbf{A}$ has a negativity measure of $\|\min\{\mathbf{A}, \mathbf{0}\}\|_F / \|\mathbf{A}\|_F = 3.08\%$, which provides a lower bound on the relative correction error: any nonnegative correction must have $\|\mathbf{X}\|_F / \|\mathbf{A}\|_F \geq 3.08\%$.

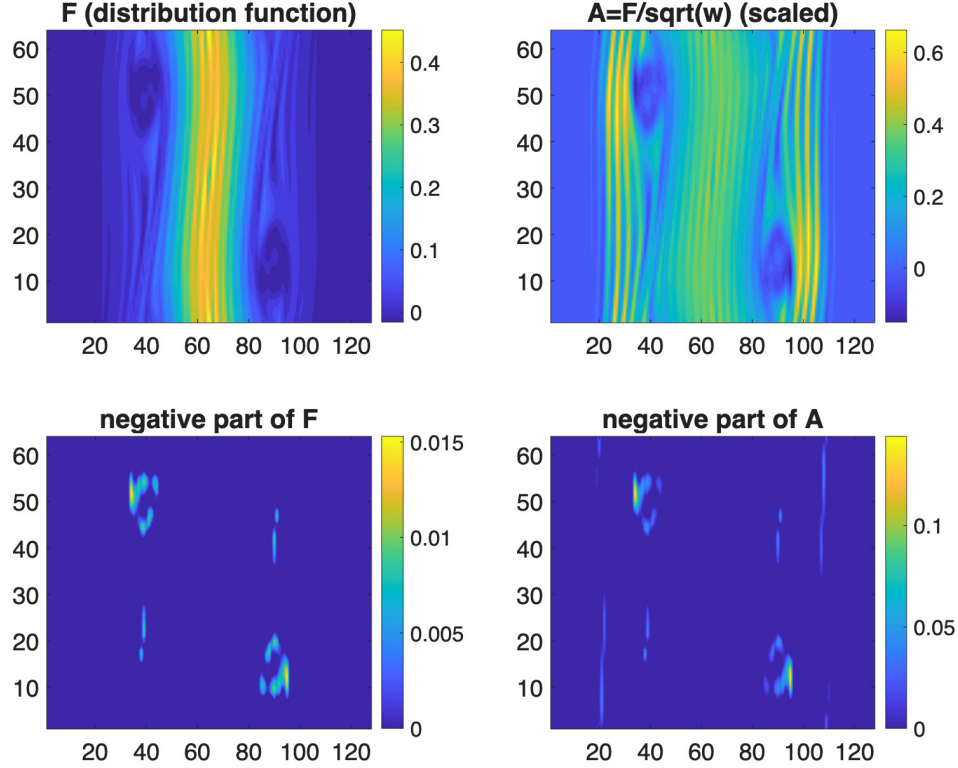


Figure 1: Input matrix $\mathbf{F}$, scaled matrix $\mathbf{A}$, and their negative parts.

For all algorithms applied to the convex formulation (2), after experimentation (see Fig. 2) we find that $a = 1000$ is a good compromise between the two objectives, and henceforth set $a = 1000$. We also set $a_1 = a_2 = a/2$, where a_1 and a_2 are the splitting parameters of the Douglas–Rachford decomposition (19), and note that the choice of splitting $a = a_1 + a_2$ does not produce visible differences in the results. Each convergence plot shows the relative correction $\|\mathbf{X}\|_F / \|\mathbf{A}\|_F$, the orthogonality violation $\|\mathbf{X}\mathbf{B}\|_F$, the rank of $\mathbf{X}$, and the relative change between successive iterates, for several values of the parameter a or the rank constraint r .

The converged relative correction $\|\mathbf{X}\|_F / \|\mathbf{A}\|_F$ and the rank of $\mathbf{X}$ depend on the penalty parameter a . Computed using the Douglas–Rachford splitting with 1000 iterations, as a increases through 1, 100, 1000, 10^4 , 10^5 the relative correction is 4.03%, 3.31%, 3.29%, 3.29%, 3.29%, decreasing monotonically toward the theoretical lower bound of 3.08%, while the rank of $\mathbf{X}$ is 38, 23, 33, 39, 45; this reflects the trade-off between correction quality and low-rank structure controlled by a . In the

comparison below we fix $a = 1000$.

Table 1 compares all algorithms on the 64×128 data, reporting the relative correction $\|\mathbf{X}\|_F / \|\mathbf{A}\|_F$, the rank of $\mathbf{X}$, the orthogonality violation $\|\mathbf{X}\mathbf{B}\|_F$, the minimum of the corrected solution $\min(\mathbf{A} + \mathbf{X})$ (nonnegativity), the number of iterations to convergence, and the total CPU time. Each iterative method is run until the relative change of the iterate falls below 10^{-8} or a cap of 1000 iterations is reached. TAP uses rank $r = 20$. For reference we also include a naive low-rank baseline: take $\mathbf{X} = \max\{-\mathbf{A}, \mathbf{0}\}$ (which makes $\mathbf{A} + \mathbf{X}$ nonnegative), truncate it to rank r by an SVD, and project the result onto $\Pi_{\mathbf{B}}^\perp$ to enforce $\mathbf{X}\mathbf{B} = \mathbf{0}$. This is a direct, non-iterative construction. The baseline is essentially free and satisfies the rank and orthogonality constraints exactly, but it does not enforce nonnegativity, so $\min(\mathbf{A} + \mathbf{X}) = -1.7 \times 10^{-2} < 0$. All convex algorithms converge to the same global minimizer (relative correction 3.29%, rank 33). The methods that recover the primal iterate as $\mathbf{X} = \nabla f^*(-\mathbf{A})$ (DR, dual FISTA I, dual AGD, dual PR+ CG, and dual L-BFGS) keep $\min(\mathbf{A} + \mathbf{X}) = 0$ exactly at every iterate and satisfy the orthogonality constraint only in the limit. Conversely, dual FISTA II recovers $\mathbf{X} = \nabla g^*(\mathbf{A})$, which is exactly orthogonal and low rank ($\|\mathbf{X}\mathbf{B}\|_F$ at machine precision) but only nonnegative in the limit (see Section 4.3). The two FISTA variants converge at nearly the same rate, with variant I slightly faster and cheaper per iteration. Douglas–Rachford splitting reaches the minimizer in the fewest iterations, whereas the dual PR+ conjugate gradient is the slowest and does not meet the tolerance within 1000 iterations (its iterate has not collapsed to rank 33 and its orthogonality violation remains 6×10^{-9}). Although L-BFGS converges in relatively few iterations, its line search evaluates the objective several times per iteration, so its total time is larger. TAP attains a comparable correction with exact orthogonality and machine-precision nonnegativity at a lower rank and the lowest cost among the iterative methods.

method	$\ \mathbf{X}\ _F / \ \mathbf{A}\ _F$	rank($\mathbf{X}$)	$\ \mathbf{X}\mathbf{B}\ _F$	$\min(\mathbf{A} + \mathbf{X})$	iters	time (s)
DR splitting	3.29%	33	7.4×10^{-10}	0	65	0.10
dual FISTA I	3.29%	33	3.2×10^{-12}	0	201	0.28
dual FISTA II	3.29%	33	1.4×10^{-16}	-1.9×10^{-16}	214	0.39
dual AGD	3.29%	33	3.8×10^{-12}	0	308	0.33
dual PR+ CG	3.29%	47	6.4×10^{-9}	0	1000	1.08
dual L-BFGS	3.29%	33	5.6×10^{-10}	0	150	0.55
TAP	3.31%	20	5.2×10^{-17}	-1.3×10^{-16}	45	0.04
baseline	2.93%	20	7.4×10^{-17}	-1.7×10^{-2}	1	0.001

Table 1: Comparison of all algorithms on the 64×128 Landau data: relative correction, rank, orthogonality violation $\|\mathbf{X}\mathbf{B}\|_F$, nonnegativity $\min(\mathbf{A} + \mathbf{X})$, iterations to convergence, and total CPU time. Convex methods use $a = 1000$, the DR step size $\eta = 7a$ and relaxation $\alpha = 0.88$, and TAP uses $r = 20$. Convergence is declared when the relative change of the iterate drops below 10^{-8} , capped at 1000 iterations, and the baseline is a direct one-shot construction (its “iterations” entry is 1). The convex methods reach the same minimizer. All of them except dual FISTA II enforce nonnegativity exactly ($\min(\mathbf{A} + \mathbf{X}) = 0$ at every iterate), whereas dual FISTA II instead enforces orthogonality exactly ($\mathbf{X}\mathbf{B} = \mathbf{0}$), as discussed in Section 4.3. PR+ conjugate gradient does not meet the tolerance within 1000 iterations. TAP and the baseline satisfy $\mathbf{X}\mathbf{B} = \mathbf{0}$ to machine precision by construction. The baseline is the fastest but violates nonnegativity, $\min(\mathbf{A} + \mathbf{X}) < 0$. CPU times are indicative and machine-dependent.

5.2 Results for the convex formulation

Douglas–Rachford splitting. We use step size $\eta = 7a$ and relaxation parameter $\alpha = 0.88$ (jointly tuned for the fastest convergence at $a = 1000$, Figure 3) in Algorithm 1, starting from $\mathbf{Z}_1 = \mathbf{0}$ and running 1000 iterations. The convergence curves for different values of a are shown in Figure 2(a).

Dual methods. The restarted dual FISTA, restarted dual accelerated gradient descent, PR+ conjugate gradient, and L-BFGS all have comparable per-iteration cost to the Douglas–Rachford splitting. For the FISTA and accelerated gradient methods, we use step size $\eta_k = a_1$ and test three restarting strategies: fixed restart (period 100), function-value-based restart, and gradient-based restart. Figure 2(b) shows the fixed-restart FISTA and compares its two dual variants (Section 4.3), and Figure 3 compares the convergence of all nine convex algorithms at the representative value $a = 1000$.

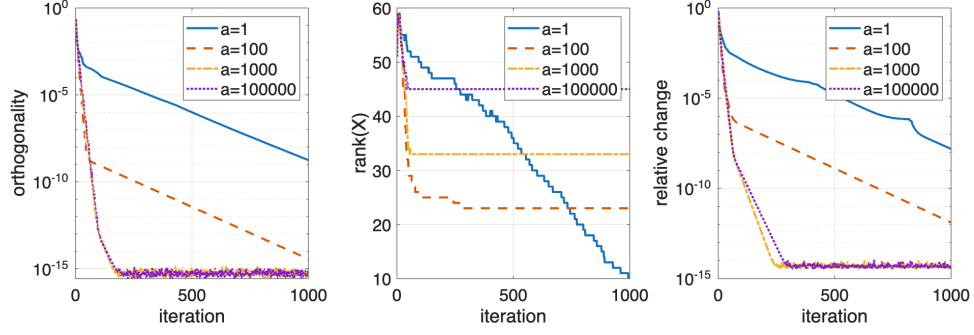
5.3 Results for the non-convex formulation

For the tangent-space accelerated alternating projection (TAP, Algorithm 3), we test different rank constraints r and run 200 iterations starting from $\mathbf{X}_0 = -\min\{\mathbf{A}, \mathbf{0}\}$. The convergence curves are shown in Figure 4.

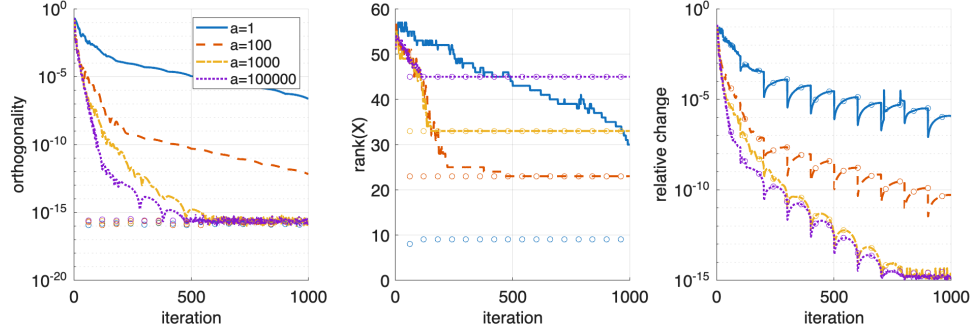
5.4 Empirical cost scaling with problem size

The timing experiments above use a single 64×128 matrix. To verify the cost scaling empirically, we run the conservative low-rank scheme of [13] for a strong Landau damping problem on a sequence of refined phase-space grids and measure the cost as the matrix size grows. We use the strong Landau damping initial condition $f_0(x, v) = (1 + \frac{1}{2} \cos \frac{x}{2}) \frac{1}{\sqrt{2\pi}} e^{-v^2/2}$ on $\Omega_x \times \Omega_v = [0, 4\pi] \times [-6, 6]$ and take the low-rank solution snapshot at $t = 10$, by which time the SVD-type truncation has introduced negative entries. For each grid we form the rescaled matrix $\mathbf{A}$ and its conservative orthogonal split $\mathbf{A} = \mathbf{A}_1 + \mathbf{A}_2$ with $\mathbf{A}_1 = \mathbf{A}\mathbf{B}\mathbf{B}^T$, exactly as in Sections 2.3–3.2. We use four grids with $N_v = 2N_x$, listed in Table 2. The orthogonality $\mathbf{A}_2\mathbf{B} = \mathbf{0}$ holds to machine precision ($\|\mathbf{A}_2\mathbf{B}\|_F \lesssim 10^{-13}$) on every grid. We emphasize that this strong Landau damping data (perturbation amplitude $\frac{1}{2}$) is different from the data used in the rest of the paper (Figure 1 and Table 1). We use it here because the same low-rank solver then generates a self-consistent family of nonnegative-correction problems across all grid resolutions, which is what we need to measure the cost scaling. Figure 5 shows the low-rank solution and its negative part for the 64×128 and 256×512 grids: the solution exhibits the same filamentary structure on both grids, while the truncation-induced negative entries become finer and smaller in magnitude as the grid is refined.

For convenience we write $m = N_x$ and $n = N_v$, so the low-rank field is an $m \times n$ matrix. The per-iteration cost of every convex algorithm is dominated by a single full $m \times n$ SVD, which costs $\mathcal{O}(mn \min(m, n))$, whereas a TAP iteration is dominated by one tangent-space projection (a $2r \times 2r$ SVD plus $\mathcal{O}(mnr)$ work, Algorithm 4) with a fixed rank r . Since $n = 2m$ here, we have $\min(m, n) = m$ and $mn = 2m^2$, so the convex per-iteration cost grows as $\mathcal{O}((mn)^{3/2})$ while the TAP per-iteration cost grows only as $\mathcal{O}(mn)$. For each grid size we time both the dominant kernel and the full algorithm (Douglas–Rachford splitting for the convex problem, with $a = 10^3$, and TAP for the non-convex problem, with $r = 20$). To make the comparison fair we strip the convergence



(a) Douglas–Rachford splitting



(b) Dual FISTA with fixed restart, the two variants overlaid

Figure 2: Convergence of two convex algorithms for different values of a . (a) Douglas–Rachford splitting. (b) Dual FISTA with fixed restart, overlaying Variant I (lines, gradient of f^* and proximal operator of g^*) and Variant II (circles, gradient of g^* and proximal operator of f^*), with curves of the same color corresponding to the same value of a . Each row shows the orthogonality violation $\|\mathbf{X}\mathbf{B}\|_F$, the rank of $\mathbf{X}$, and the relative change between iterates. Variant II keeps the orthogonality violation at machine precision throughout, since its iterate satisfies $\mathbf{X}\mathbf{B} = \mathbf{0}$ exactly.

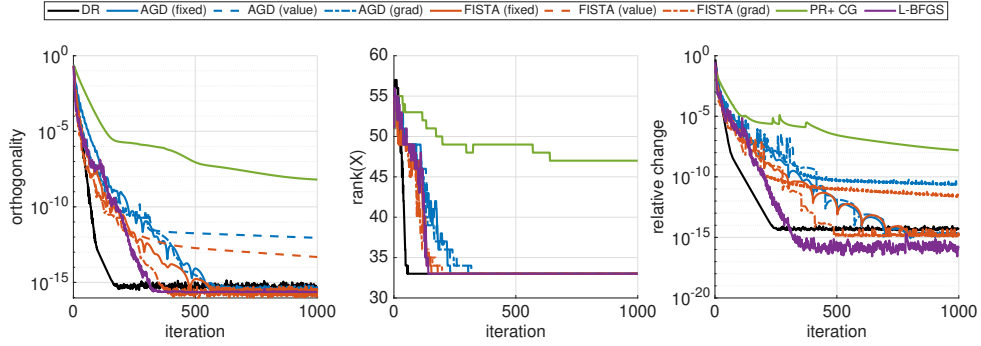


Figure 3: Convergence of all nine convex algorithms at the single penalty $a = 1000$ ($a_1 = a_2 = a/2$): the orthogonality violation $\|\mathbf{XB}\|_F$, the rank of $\mathbf{X}$, and the relative change between iterates versus iteration. Color encodes the method family (Douglas–Rachford, accelerated gradient, FISTA, PR+ conjugate gradient, L-BFGS) and line style encodes the restart strategy (fixed, function-value, gradient), with dual FISTA shown for Variant I. The DR step size $\eta = 7a$ and relaxation $\alpha = 0.88$, jointly tuned at this a , give the fastest convergence (65 iterations, versus 114 at the untuned $\eta = 10^{0.5}a$, $\alpha = 0.9$). All methods reach the same minimizer (rank 33, relative correction 3.29%) except the PR+ conjugate gradient, which has not converged within 1000 iterations.

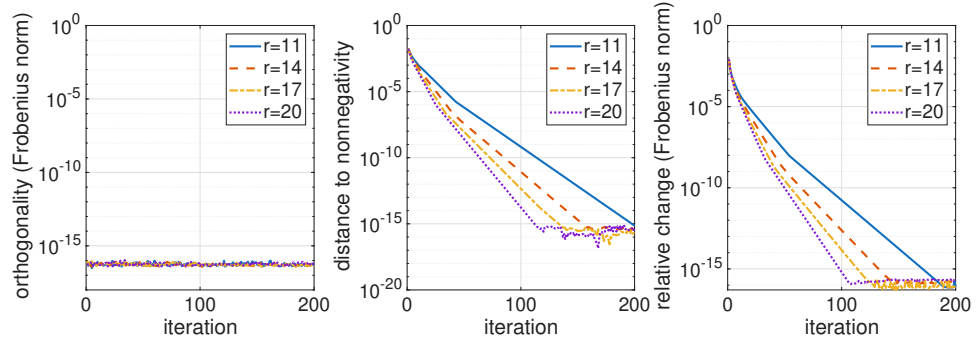


Figure 4: Tangent-space accelerated alternating projection (TAP): convergence for different rank constraints r . Left: orthogonality violation $\|\mathbf{XB}\|_F$, which stays at machine precision throughout because TAP keeps $\mathbf{XB} = \mathbf{0}$ exactly at every iterate. Middle: distance to the nonnegative set, $\max(-\min\{\mathbf{A} + \mathbf{X}, \mathbf{0}\})$. Right: relative change between successive iterates.

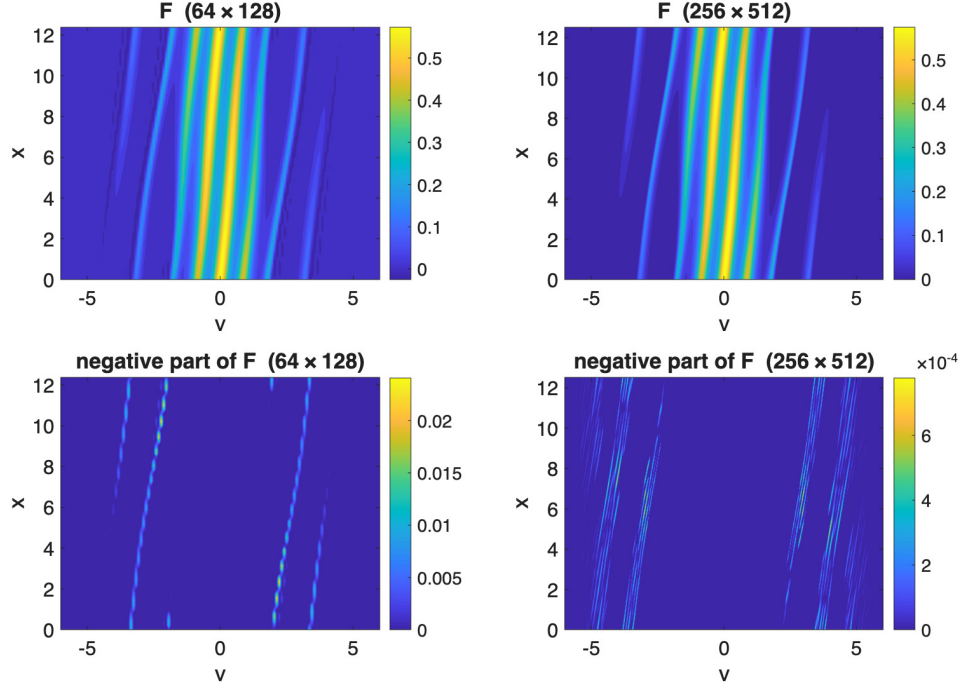


Figure 5: Low-rank Landau damping solution F (top) and its negative part $-\min\{F, 0\}$ (bottom) on the 64×128 (left) and 256×512 (right) grids, in the style of the left column of Figure 1. The negative entries, introduced by the SVD-type truncation, concentrate along the filament edges and become finer and smaller in magnitude under refinement.

grid $N_x \times N_v$	mn	convex (DR splitting)		non-convex (TAP)		ratio
		full SVD (s)	DR / iter (s)	tangent proj. (s)	TAP / iter (s)	
32×64	2,048	1.9×10^{-4}	2.6×10^{-4}	3.1×10^{-4}	3.0×10^{-4}	0.9
64×128	8,192	5.5×10^{-4}	8.9×10^{-4}	4.4×10^{-4}	4.9×10^{-4}	1.8
128×256	32,768	3.1×10^{-3}	4.1×10^{-3}	6.3×10^{-4}	8.7×10^{-4}	4.7
256×512	131,072	1.7×10^{-2}	2.8×10^{-2}	1.1×10^{-3}	1.6×10^{-3}	18.0

Table 2: Per-iteration cost scaling across grid sizes for the Landau damping problem (median of three trials, convergence diagnostics excluded). For the convex problem we report the cost of one full $m \times n$ SVD (the dominant kernel) and the full per-iteration cost of the Douglas–Rachford splitting. For the non-convex problem we report the cost of one tangent-space projection (the $2r \times 2r$ -SVD kernel) and the full per-iteration cost of TAP, with $r = 20$ fixed. The last column is the per-iteration ratio, DR / TAP. The negativity measure $\|\min\{\mathbf{A}, \mathbf{0}\}\|_F / \|\mathbf{A}\|_F$ ranges from 2.8% to 11.3% and $\text{rank}(\mathbf{A}_2)$ from 23 to 41 across the four grids.

diagnostics from both inner loops, so the reported times reflect only the algorithmic work. The results are reported in Table 2 and Figure 6.

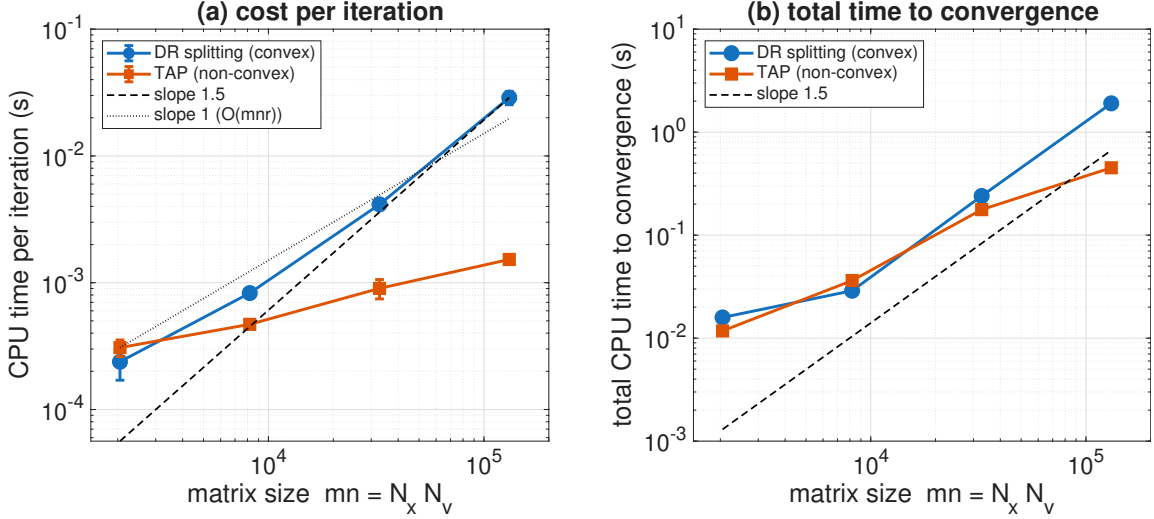


Figure 6: Cost scaling versus matrix size $mn = N_x N_v$ (log-log). (a) Per-iteration cost of the Douglas–Rachford splitting (convex) and the TAP algorithm (non-convex), plotted as the median over 20 runs with error bars of ± 1 standard deviation. The DR cost follows the predicted $(mn)^{3/2}$ slope at the larger sizes, while the TAP cost grows *sub-linearly* (empirical slope ≈ 0.4), below its asymptotic $\mathcal{O}(mnr)$ rate (slope 1). In this size range a TAP iteration is dominated by the size-independent cost of the tangent-space projection (the $2r \times 2r$ SVD and the QR of the tangent factors, $\mathcal{O}(r^3 + (m+n)r^2)$) rather than by the $\mathcal{O}(mnr)$ matrix products, which take only microseconds, so the linear regime would set in only at much larger mn . The per-iteration gap between DR and TAP nonetheless reaches about $18\times$ at 256×512 . (b) Total wall-clock time to convergence (relative change of the iterate below 10^{-8} , with the DR step size tuned per grid): the two methods are comparable on the smaller grids, while TAP is about $4\times$ faster at 256×512 . With its tuned step size DR needs fewer iterations than TAP, but its per-iteration full SVD makes each iteration much more expensive.

The measured costs confirm the predicted per-iteration scaling (Figure 6(a)). The convex per-iteration cost follows the $(mn)^{3/2}$ reference slope at the larger grids (the DR per-iteration time increases by a factor of about 6.9 over the last fourfold increase in mn , approaching the asymptotic factor of 8), while the TAP per-iteration cost grows much more slowly. The per-iteration advantage of TAP therefore grows monotonically with the problem size, from being slightly more expensive on the smallest 32×64 grid (where the highly optimized full SVD of a tiny matrix is cheaper than the TAP projection overhead) to being 18 times cheaper on the 256×512 grid.

The total cost is the per-iteration cost times the number of iterations to convergence, so we also examine how the iteration count scales with the problem size. We run both algorithms until the relative change of the iterate $\|\mathbf{X}_k - \mathbf{X}_{k-1}\|_F / \|\mathbf{X}_{k-1}\|_F$ falls below 10^{-8} , with $a = 10^3$ and $r = 20$ for TAP. For DR we tune the step size at each grid, so that DR is compared at its best at every

size. The optimal step grows with the problem (from $\eta \approx 2a$ on the coarsest grid to $\eta \approx 11a$ on the finest, with relaxation α between 0.90 and 0.94). The two algorithms reach comparable relative corrections $\|\mathbf{X}\|_F / \|\mathbf{A}\|_F$ (last two columns of Table 3), so the comparison is fair. With its step size tuned, DR converges in markedly fewer iterations than TAP, and the gap widens with refinement (DR from 30 to 82 iterations versus TAP from 32 to 264, Table 3). Each DR iteration nevertheless requires a full SVD, whereas a TAP iteration requires only a $2r \times 2r$ SVD, so DR’s per-iteration cost grows much faster (Figure 6(a)). The two effects offset in the total wall-clock time (Figure 6(b)), which is comparable for the two methods on the three smaller grids (DR is even slightly faster at 64×128) while TAP is about $4\times$ faster on the finest 256×512 grid. Since the per-iteration cost gap scales as $\min(m, n)/r = m/r$ while DR’s iteration advantage grows only mildly, we expect TAP’s total-time advantage to keep widening on still finer grids.

grid $N_x \times N_v$	mn	iterations to conv.		total time (s)		$\ \mathbf{X}\ _F / \ \mathbf{A}\ _F$	
		DR	TAP	DR	TAP	DR	TAP
32×64	2,048	30	32	0.016	0.012	11.6%	11.6%
64×128	8,192	31	74	0.029	0.036	5.8%	6.2%
128×256	32,768	55	196	0.24	0.18	8.1%	9.5%
256×512	131,072	82	264	1.90	0.45	2.8%	3.7%

Table 3: Iterations to convergence (relative change of the iterate below 10^{-8}) and total wall-clock time for DR splitting and TAP ($r = 20$) across grid sizes, with $a = 10^3$. The DR step size is tuned per grid (the optimum grows with the grid: $\eta/a = 2.0, 2.5, 7.1, 11.2$, relaxation $\alpha = 0.90$ on the two coarse grids and 0.94 on the two fine grids). With its step size tuned, DR converges in fewer iterations than TAP, increasingly so as the grid is refined. Each DR iteration nevertheless requires a full SVD, so its per-iteration cost grows faster, and the total time favors TAP only on the finest grid. The last two columns show that the two algorithms reach comparable relative corrections. CPU times are indicative and machine-dependent.

Remark 7 (Accelerating the per-iteration SVD). *The per-iteration cost of the convex algorithms is dominated by a single full $m \times n$ SVD (Table 2), so it is natural to ask whether a cheaper low-rank factorization can replace it. A maxvol-based cross (CUR) approximation [10], which samples a few rows and columns without ever forming the full matrix, is slower than the built-in SVD at every size we tested (for instance about 8 times slower at 128×256). The reason is that it is designed to minimize the number of entry evaluations, an advantage only when the entries are expensive to compute, as for implicit functions or high-dimensional tensors, whereas our matrices are explicit and stored in memory, so a single level-3 BLAS SVD is faster.*

A randomized SVD [15] is more promising, but its benefit is strongly rank-dependent, and this is the essential point. Its cost scales as $\mathcal{O}(mnr)$, so it beats the full SVD, which costs $\mathcal{O}(mn \min(m, n))$, only when the numerical rank r is a small fraction of $\min(m, n)$. This holds for the data matrix $\mathbf{A}$, which is genuinely low rank: there the randomized SVD reproduces the same singular triplets and is faster on the finer grids, by about three times on the finest (Table 4, upper block). It does not hold for the matrices that the DR iteration actually factors. Each iteration applies the SVD not to $\mathbf{A}$ but to a proximal input $\mathbf{W}$, a dense intermediate matrix formed from the current iterates that is numerically full rank, and the soft-thresholded output retains a large fraction of the spectrum,

roughly half of $\min(m, n)$ (Table 4, lower block). On these matrices the same randomized SVD gives no speedup and is in fact slower at every size. Replacing the SVD by a randomized one therefore does not accelerate DR, even though it reproduces the same iterates to controllable accuracy. The low $\mathcal{O}(mnr)$ per-iteration cost is instead achieved by the tangent-space method, which keeps the iterate in rank- r factored form and never forms or factors a full-rank matrix.

grid $N_x \times N_v$	r	full SVD (s)	randomized SVD (s)	speedup
<i>data matrix</i> $\mathbf{A}$ (genuinely low rank)				
64×128	31	2.9×10^{-4}	6.6×10^{-4}	$0.44\times$
128×256	43	2.2×10^{-3}	1.6×10^{-3}	$1.35\times$
256×512	44	9.5×10^{-3}	3.2×10^{-3}	$3.02\times$
<i>DR proximal matrix</i> $\mathbf{W}$ (numerically full rank)				
64×128	58	3.6×10^{-4}	1.3×10^{-3}	$0.29\times$
128×256	84	2.2×10^{-3}	3.9×10^{-3}	$0.58\times$
256×512	121	1.2×10^{-2}	1.3×10^{-2}	$0.95\times$

Table 4: Rank-dependence of the randomized-SVD speedup. The same randomized SVD (subspace iteration with oversampling) is applied to two matrices at each grid: the data matrix $\mathbf{A}$, which is genuinely low rank, and a representative Douglas–Rachford proximal input $\mathbf{W}$, a dense intermediate matrix formed from the current iterates that is numerically full rank. Here r is the numerical rank of $\mathbf{A}$ and the retained rank of the thresholded $\mathbf{W}$, and the randomized SVD targets this r . The speedup is the full-SVD time divided by the randomized-SVD time (fastest of seven trials), so a value above one means the randomized SVD is faster. On the low-rank $\mathbf{A}$ it is faster on the finer grids, whereas on the full-rank $\mathbf{W}$ it is slower at every size, because its $\mathcal{O}(mnr)$ cost helps only when $r \ll \min(m, n)$. CPU times are indicative and machine-dependent.

5.5 Application as a positivity limiter in the LoMaC scheme

So far we have applied the correction to a single fixed matrix. We now use it as a *positivity limiter* inside the time-dependent LoMaC low-rank scheme of [14], which conserves total mass, momentum *and* energy at the discrete level. As the low-rank solution is advanced, the SVD-type truncation at each step introduces small negative entries. We apply the correction periodically (every k steps and on the final step, each application an SVD-scale computation) to remove the unphysical negative entries while preserving the macroscopic moments, since its constraint $\mathbf{XB} = \mathbf{0}$ leaves the three conserved densities (mass, momentum and energy) untouched. We compare two correction algorithms as the limiter. The first is the convex Douglas–Rachford splitting (DR, Algorithm 1), which imposes *no a-priori rank* on the correction: the corrected field takes whatever rank nonnegativity requires, and its rank is reduced only by the scheme’s own conservative truncation $\mathcal{T}_{\varepsilon_{\text{tr}}}$ at a tolerance ε_{tr} . (For DR we use the step size $\eta = 10a$, found near-optimal by a step-size sweep on a representative correction subproblem: 83 iterations to drive $\|\mathbf{XB}\|_F$ below 10^{-10} , versus 190 at $\eta = 10^{0.5}a$ and 602 at $\eta = a$.) The second is the rank-constrained TAP (Algorithm 3). The essential point is that its rank must *not* be fixed at some foreign value. Instead, we let it track the solution, taking $r = N + 3$ where N is the current numerical rank of the low-rank solution at each

application. This gives TAP the rank the field already carries, plus a small margin for the extra rank that nonnegativity introduces, so that, like DR, it is never artificially capped. (A fixed rank, e.g. $r = 20$, would cap the attainable nonnegativity, as we note below.)

We test the limiter on the bump-on-tail instability used to validate the LoMaC scheme (Example 5.4 of [14]), a well-resolved, genuinely low-rank run of the same family as the data in Figure 1, with initial datum

$$f_0(x, v) = (1 + \alpha \cos(kx)) \left(n_p e^{-v^2/2} + n_b e^{-(v-u)^2/(2v_t^2)} \right),$$

$\alpha = 0.04$, $k = 0.3$, $n_p = \frac{9}{10\sqrt{2\pi}}$, $n_b = \frac{2}{10\sqrt{2\pi}}$, $u = 4.5$, $v_t = 0.5$, on $[0, 2\pi/k] \times [-13, 13]$, integrated to $T = 30$ with $\varepsilon = 10^{-4}$, using the validated LoMaC discretization (linear fifth-order upwind in x , CFL = 0.1 on both grids, so that the two runs share the same time-step criterion). We run two grids, 64×128 and 128×256 , and apply the limiter every $k = 1000$ steps. Table 5 and Figure 7 summarize the results.

In this regime the LoMaC solution is genuinely low rank (numerical rank 28 on 64×128 and 49 on 128×256 at $T = 30$), and the scheme conserves total mass, momentum *and* energy to machine precision ($\sim 10^{-14}$ for all three). The SVD truncation nonetheless leaves a negativity $\|\min\{\mathbf{F}, \mathbf{0}\}\|_F / \|\mathbf{F}\|_F \approx 2\text{--}3 \times 10^{-4}$. The limiter removes it, and because the correction satisfies $\mathbf{X}\mathbf{B} = \mathbf{0}$ it does so without disturbing the three conserved quantities. Here the two correction algorithms differ in an instructive way (Table 5): TAP keeps $\mathbf{X}_k\mathbf{B} = \mathbf{0}$ *exactly at every iterate*, so the limited run conserves mass, momentum and energy to machine precision ($\sim 10^{-14}$), just like the unlimited LoMaC scheme. DR enforces $\mathbf{X}\mathbf{B} = \mathbf{0}$ only in the limit, so each application perturbs the moments by its orthogonality residual $\|\mathbf{X}\mathbf{B}\|_F$, and conservation is preserved only to that level ($\sim 10^{-10}$ at $\varepsilon_{\text{tr}} = 10^{-4}$, improving to $\sim 10^{-12}$ as the correction is converged more tightly at $\varepsilon_{\text{tr}} = 10^{-10}$). For exact conservation, then, the rank-constrained TAP limiter is preferable, which mirrors the constraint-preservation property noted in Section 5.6.

The other point concerns the interplay between nonnegativity and rank, and it explains why the stored solution is not nonnegative to machine zero by default. The correction itself enforces nonnegativity *absolutely*: the corrected field satisfies $\mathbf{A} + \mathbf{X} \geq \mathbf{0}$ to machine precision *before* truncation (for DR, $\mathbf{X}$ is the exact pointwise clip $\mathbf{X} = \max\{-\mathbf{A}, \cdot\}$, so $\mathbf{A} + \mathbf{X} \geq \mathbf{0}$ identically). Negativity reappears *only* when this corrected field is re-expressed in the rank-bounded low-rank format. A nonnegative field generically has higher rank than the smooth, slightly negative one, so the conservative truncation $\mathcal{T}_{\varepsilon_{\text{tr}}}$ must discard a singular-value tail of norm $O(\varepsilon_{\text{tr}})$. Because the corrected field sits exactly at zero where it had been negative, discarding this tail perturbs it by $O(\varepsilon_{\text{tr}})$ and pushes those entries just below zero, leaving a negativity of magnitude $O(\varepsilon_{\text{tr}})$. The negativity floor of the stored low-rank solution is therefore set by the truncation tolerance ε_{tr} , not by the correction: at $\varepsilon_{\text{tr}} = 10^{-4}$ (the scheme's own truncation level) the negativity is $\sim 2 \times 10^{-6}$ at a modestly higher rank ($49 \rightarrow 51\text{--}53$ on 128×256), while tightening to $\varepsilon_{\text{tr}} = 10^{-10}$ drives it toward machine precision ($\sim 10^{-12}$, $\min \mathbf{F} \sim -10^{-11}$ to -10^{-12}) at a higher rank ($49 \rightarrow 84$ for DR, $49 \rightarrow 91$ for TAP). Provided the rank is allowed to grow, both algorithms reach comparable floors (Table 5 and Figure 7). This is also why letting TAP's rank track the solution ($r = N + 3$) matters: a fixed externally-imposed rank caps the attainable nonnegativity (with $r = 20$ the rank-20 correction stalls at $\min(\mathbf{A} + \mathbf{X}) \approx -1.4 \times 10^{-10}$ regardless of ε_{tr}), whereas the adaptive rank, like the unconstrained convex correction, removes that cap. In short, the residual negativity is governed by how much rank one keeps in the low-rank format, not by the correction itself.

Figure 7(a) shows the negativity of the stored low-rank field (that is, after the correction *and* the re-truncation) just after each limiter application. The DR and TAP curves overlap at each ε_{tr} ,

reaching $\sim 2 \times 10^{-6}$ at $\varepsilon_{\text{tr}} = 10^{-4}$ and $\sim 10^{-12}$ at $\varepsilon_{\text{tr}} = 10^{-10}$, confirming that the floor follows ε_{tr} rather than the correction. (Between applications the per-step truncation at ε lets the negativity regrow, bounded by the no-limiter level, so a moderate k controls it while the final output, limited on the last step, is nonnegative to ε_{tr} .) Figure 7(b) shows the rank: both limiters hold the solution at the higher rank that nonnegativity demands, increasingly so as ε_{tr} is tightened. The cost is modest and comparable for the two methods: with $k = 1000$ the limiter performs 15 (64×128) or 30 (128×256) correction calls, raising the wall-clock time on 64×128 from 23 to 28s and on 128×256 from 174s to 208–230s (a 20–30% overhead, dominated by the higher rank sustained between calls rather than by the correction solves). TAP is the cheaper of the two here, and its per-iteration advantage over DR grows on finer grids (Section 5.4).

mesh	limiter	mass dev.	mom. dev.	energy dev.	negativity	rank	CPU (s)
64×128	none	5.6×10^{-15}	2.5×10^{-14}	1.5×10^{-14}	1.8×10^{-4}	28	23
64×128	DR, $\varepsilon_{\text{tr}} = 10^{-4}$	4.4×10^{-14}	8.3×10^{-14}	1.6×10^{-13}	1.1×10^{-6}	27	28
64×128	TAP, $\varepsilon_{\text{tr}} = 10^{-4}$	6.6×10^{-15}	2.5×10^{-14}	1.6×10^{-14}	1.5×10^{-6}	28	28
64×128	DR, $\varepsilon_{\text{tr}} = 10^{-10}$	2.7×10^{-14}	2.2×10^{-13}	5.5×10^{-13}	2.1×10^{-12}	54	28
64×128	TAP, $\varepsilon_{\text{tr}} = 10^{-10}$	3.8×10^{-15}	2.6×10^{-14}	1.7×10^{-14}	2.8×10^{-12}	53	28
128×256	none	9.8×10^{-15}	4.0×10^{-14}	3.5×10^{-14}	3.0×10^{-4}	49	174
128×256	DR, $\varepsilon_{\text{tr}} = 10^{-4}$	1.7×10^{-11}	1.1×10^{-10}	6.3×10^{-11}	2.0×10^{-6}	53	222
128×256	TAP, $\varepsilon_{\text{tr}} = 10^{-4}$	9.9×10^{-15}	3.3×10^{-14}	3.0×10^{-14}	2.0×10^{-6}	51	208
128×256	DR, $\varepsilon_{\text{tr}} = 10^{-10}$	5.6×10^{-13}	2.2×10^{-12}	6.7×10^{-13}	2.9×10^{-13}	84	230
128×256	TAP, $\varepsilon_{\text{tr}} = 10^{-10}$	9.3×10^{-15}	3.4×10^{-14}	2.9×10^{-14}	1.4×10^{-12}	91	208

Table 5: Correction as a positivity limiter in the LoMaC low-rank scheme [14] for the bump-on-tail run, $T = 30$, applied every 1000 steps, comparing the convex DR limiter (no rank cap) with the rank-constrained TAP limiter at rank $r = N + 3$ ($N =$ current solution rank). The negativity is $\|\min\{\mathbf{F}, \mathbf{0}\}\|_F / \|\mathbf{F}\|_F$ for the limited (post-application) solution. The LoMaC scheme conserves mass, momentum and energy to machine precision. TAP keeps $\mathbf{X}_k \mathbf{B} = \mathbf{0}$ exactly at every iterate, so the limited run conserves all three to machine precision. DR enforces $\mathbf{X} \mathbf{B} = \mathbf{0}$ only at convergence, so it preserves conservation to its orthogonality residual ($\sim 10^{-10}$ at $\varepsilon_{\text{tr}} = 10^{-4}$, $\sim 10^{-12}$ at $\varepsilon_{\text{tr}} = 10^{-10}$). Both drive the negativity to the re-truncation floor ε_{tr} , with the kept rank rising as ε_{tr} is tightened. Both grids use CFL = 0.1. The per-step cost of the LoMaC integration scales as $\mathcal{O}((m+n)r^2)$, linear in the grid size and quadratic in the rank r , so it remains a low-rank cost even as the limiter raises r . Enforcing nonnegativity increases the rank only modestly and transiently: at $\varepsilon_{\text{tr}} = 10^{-4}$, the scheme’s own truncation level, the rank is essentially unchanged, whereas at $\varepsilon_{\text{tr}} = 10^{-10}$ it roughly doubles (for instance $28 \rightarrow 53$ on 64×128 and $49 \rightarrow 91$ on 128×256), but this higher rank is sustained only briefly after each application before the per-step truncation relaxes it. The time-averaged overhead is therefore only 20–30% and stays far below the full-grid $\mathcal{O}(mn \min(m, n))$ cost. CPU times are indicative and machine-dependent.

One might worry that the numerical rank of the limited solution, which reaches a sizeable fraction of $\min(N_x, N_v)$ on the 128×256 grid of Figure 7, would approach full rank as the mesh is refined. It does not: the near-full appearance is an artifact of the small $\min(N_x, N_v)$ at this resolution, not of the low-rank structure. To exhibit the trend cheaply we run the same test on the coarser 32×64 and 64×128 grids (Figure 8). The peak numerical rank grows only sub-linearly with

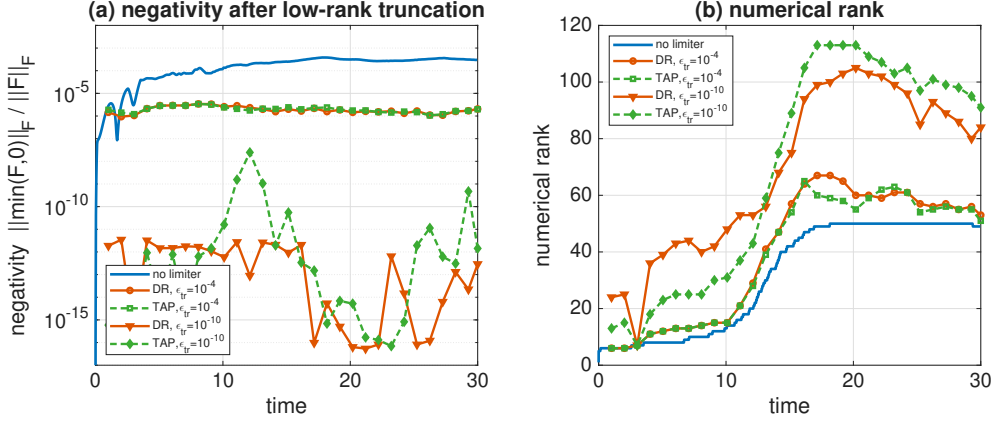


Figure 7: Positivity limiter in the LoMaC low-rank scheme for the bump-on-tail run (128×256 , $T = 30$), applied every 1000 steps, comparing the convex DR limiter (no rank cap) with the rank-constrained TAP limiter at rank $r = N + 3$. (a) Negativity $\|\min\{\mathbf{F}, \mathbf{0}\}\|_F / \|\mathbf{F}\|_F$ of the stored low-rank field versus time. The correction enforces nonnegativity to machine precision *before* truncation ($\mathbf{A} + \mathbf{X} \geq \mathbf{0}$ identically). The residual negativity plotted here is reintroduced *solely* by re-expressing the corrected field in the rank-bounded low-rank format, by discarding singular values below ε_{tr} , so it scales with ε_{tr} . Without the limiter it grows to $\sim 3 \times 10^{-4}$. With either limiter each application resets it to this re-truncation floor, $\sim 2 \times 10^{-6}$ at $\varepsilon_{\text{tr}} = 10^{-4}$ and $\sim 10^{-12}$ at $\varepsilon_{\text{tr}} = 10^{-10}$ (DR and TAP curves overlap). (b) Numerical rank: enforcing nonnegativity raises the kept rank above the no-limiter value (49), more so as ε_{tr} is tightened (≈ 51 – 53 at $\varepsilon_{\text{tr}} = 10^{-4}$, ≈ 84 – 91 at $\varepsilon_{\text{tr}} = 10^{-10}$). The rank is reduced only by the scheme’s conservative truncation, never by a mechanical cap, and DR and adaptive-rank TAP track each other.

the grid (for the no-limiter run it is 18, 28, 50 across 32×64 , 64×128 , 128×256 , each refinement multiplying it by well under the factor 2 that a full-rank field would require), so as a fraction of the maximum $\min(N_x, N_v)$ it *decreases* under refinement. On the coarsest grid the tight-tolerance ($\varepsilon_{\text{tr}} = 10^{-10}$) limiter is pinned at the full-rank ceiling, whereas on 128×256 it already sits at 80–88% of it. Doubling the resolution therefore makes the rank a smaller fraction of full, not larger, confirming that the limited solution remains genuinely low rank.

5.6 Discussion

We summarize the main observations from the experiments of this section, which comprise a fixed-size correction problem (the 64×128 Landau data, Table 1 and Figures 2–4), a cost-scaling study across four grid sizes (Section 5.4), and the use of the correction as a positivity limiter in a time-dependent solver (Section 5.5).

Approximation quality. All the algorithms produce corrections of comparable quality. For $a \geq 10$ in the convex formulation and for every rank constraint r tested in TAP, the relative correction $\|\mathbf{X}\|_F / \|\mathbf{A}\|_F$ is close to the theoretical lower bound (here 3.08%) and decreases monotonically as a or r increases. All convex algorithms reach the same global minimizer (3.29%, rank 33), reflecting the global optimality of the convex formulation, and TAP attains a comparable correction (3.31%) at the prescribed lower rank $r = 20$. The penalty a and the rank bound r play the same role of trading correction quality against the rank of $\mathbf{X}$. The naive low-rank baseline matches the rank and orthogonality constraints but cannot enforce nonnegativity ($\min(\mathbf{A} + \mathbf{X}) < 0$), and its smaller relative correction (2.93%, below the 3.08% bound) is possible only because it violates nonnegativity, which is exactly the property the optimization-based corrections are designed to deliver.

Convergence speed. On the fixed-size problem, Douglas–Rachford splitting converges in the fewest iterations among the convex algorithms (about 65 at $a = 1000$, Table 1), once its step size and relaxation are tuned (Figure 3). The dual methods are slower. Among them the gradient-based restart performs best, the L-BFGS method needs few iterations but an expensive line search, and the PR+ conjugate gradient is the slowest and does not meet the tolerance within 1000 iterations. The tangent-space accelerated alternating projection converges in the fewest iterations overall (about 45).

Per-iteration cost and scaling. The two formulations differ fundamentally in per-iteration cost: every convex iteration requires a full $m \times n$ SVD ($\mathcal{O}(mn \min(m, n))$), whereas a TAP iteration requires only a $2r \times 2r$ SVD ($\mathcal{O}(mnr)$). The scaling study confirms this empirically. The per-iteration cost ratio of DR to TAP grows from order one to about 18 as the grid is refined from 32×64 to 256×512 (Table 2 and Figure 6(a)). With its step size tuned at each grid, DR converges in fewer iterations than TAP, but its much higher per-iteration cost dominates the total wall-clock time: the two methods are comparable on the smaller grids and TAP is about $4\times$ faster on the finest grid (Table 3 and Figure 6(b)), an advantage that widens under refinement since the cost gap scales as m/r . TAP is therefore the most cost-efficient method, increasingly so at scale, while the convex formulation offers global optimality and lets the correction rank adapt freely.

Constraint preservation. The formulations also differ in which constraint the iterate satisfies before convergence. TAP keeps the orthogonality $\mathbf{X}_k \mathbf{B} = \mathbf{0}$ exactly at every iterate, and the iterate

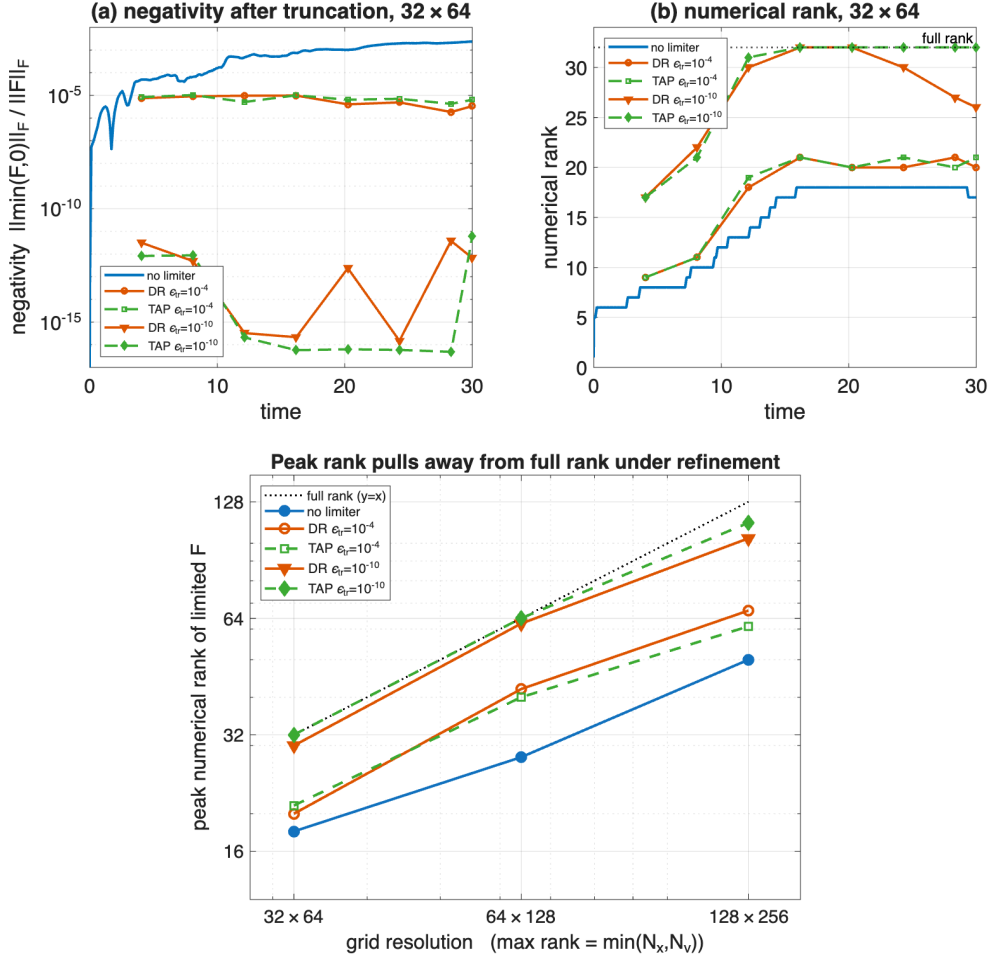


Figure 8: Rank of the limited solution versus mesh resolution, addressing whether it approaches full rank under refinement. Top: the coarse-grid (32×64) counterpart of Figure 7, showing the negativity (left) and the numerical rank (right) versus time; on this small grid the tight-tolerance limiter is pinned at the full-rank ceiling $\min(N_x, N_v) = 32$ (dotted). Bottom: peak numerical rank of the limited solution versus grid resolution for each configuration, against the full-rank line $y = \min(N_x, N_v)$ (dotted). Every curve pulls away from the full-rank line as the mesh is refined, since the absolute rank grows sub-linearly, so the near-full rank seen in Figure 7 is a coarse-grid artifact and doubling the resolution reduces the rank as a fraction of the maximum. All runs use $\text{CFL} = 0.1$.

stays low rank, whereas the convex methods that recover the primal variable by clipping (DR, dual FISTA I, dual AGD, PR+ CG, L-BFGS) keep the nonnegativity $\mathbf{A} + \mathbf{X}_k \geq \mathbf{0}$ exactly at every iterate. The dual FISTA variant II is the mirror case, keeping $\mathbf{X}_k \mathbf{B} = \mathbf{0}$ exactly instead. Both constraints hold at convergence, but the per-iterate property matters when the iteration is stopped early or embedded in a larger computation.

Use as a positivity limiter. The last experiment (Section 5.5) embeds the correction as a positivity limiter in the conservative LoMaC low-rank Vlasov solver. Both the convex (DR) and the rank-constrained (TAP) corrections remove the negativity introduced by the SVD truncation while leaving the conserved moments untouched, since $\mathbf{X} \mathbf{B} = \mathbf{0}$. The per-iterate constraint preservation above is decisive here: TAP, which keeps $\mathbf{X}_k \mathbf{B} = \mathbf{0}$ exactly, preserves mass, momentum and energy to machine precision, whereas DR preserves them only to its orthogonality residual. The correction itself enforces nonnegativity to machine precision, so the only residual negativity comes from re-expressing the corrected field in the rank-bounded low-rank format, and its level is set by the re-truncation tolerance, reaching machine zero as that tolerance is tightened at the cost of a higher but still moderate rank. For this to succeed, the correction rank must be allowed to track the solution rather than be fixed in advance.

6 Conclusion

We developed optimization-based post-processing algorithms to recover nonnegativity in low-rank numerical solutions of Vlasov dynamics while preserving the macroscopic conservation laws pointwise. The conservation requirement is written as an orthogonality constraint on the correction term. We studied two optimization formulations. The first is the convex problem (2) based on squared nuclear norm minimization. For this problem, Theorem 1 shows that the proximal operator can be computed through an implicit singular value thresholding equation with a threshold determined by bisection. We developed five algorithms for this formulation: Douglas–Rachford splitting, restarted dual FISTA, restarted dual accelerated gradient descent, dual PR+ conjugate gradient, and dual L-BFGS. The second is the rank-constrained non-convex formulation (17). For this formulation, we developed a tangent-space accelerated alternating projection algorithm that only requires a $2r \times 2r$ SVD per iteration. The numerical results for the Landau damping test case show that all algorithms produce comparable correction quality. The TAP algorithm is the most cost-efficient because of its low per-iteration cost, and this advantage becomes more pronounced as the problem size increases. Among the convex algorithms, Douglas–Rachford splitting converges the fastest. The gradient-restarted dual methods give the best performance among the dual approaches. The convex formulation has the advantage of global optimality, while the TAP algorithm naturally preserves low-rank storage throughout the iteration. Finally, we demonstrated that the correction can be used as a positivity limiter inside the time-dependent LoMaC low-rank scheme, which conserves mass, momentum and energy, using either the convex DR algorithm or the rank-constrained TAP algorithm with its rank tracking the solution ($r = N + 3$). On the validated, genuinely low-rank bump-on-tail run up to $T = 30$, both limiters remove the truncation negativity. Because the correction rank is not capped a priori, enforcing nonnegativity raises the kept rank, and the residual negativity is set by the re-truncation tolerance, driven toward machine precision as that tolerance is tightened. Since TAP maintains the orthogonality constraint exactly at every iterate, the TAP limiter preserves all three conserved quantities to machine precision, whereas DR

preserves them to its orthogonality residual. Both methods incur a modest, comparable overhead.

Declaration of AI-assisted technologies

All original ideas, as well as codes, are attributed to the authors. During the preparation of this work, the authors used Anthropic’s Claude Code to assist with mathematical discussions, further numerical implementation, and drafting of the manuscript. After using this tool, the authors carefully reviewed and edited the content as needed and take full responsibility for the final publication.

Declarations

Funding. X. Zhang is partially supported by NSF DMS-2208518. J.-M. Qiu acknowledges support provided by Air Force Office of Scientific Research FA9550-24-1-0254 and Department of Energy DE-SC0023164. S. Becker acknowledges support from the Department of Energy DE-SC0023346.

Competing Interests. The authors have no competing interests to declare that are relevant to the content of this article.

Data and Code Availability. The test data and the code used to generate the numerical results in this paper are available from the authors upon request.

Author Contributions. All authors contributed to the conception of the study, the methodology, the numerical implementation, and the writing of the manuscript. All authors read and approved the final manuscript.

References

- [1] A. Beck. *First-Order Methods in Optimization*, volume 25 of *MOS-SIAM Series on Optimization*. SIAM, Philadelphia, PA, 2017.
- [2] A. Beck and M. Teboulle. A fast iterative shrinkage-thresholding algorithm for linear inverse problems. *SIAM Journal on Imaging Sciences*, 2(1):183–202, 2009.
- [3] J.-F. Cai, E. J. Candès, and Z. Shen. A singular value thresholding algorithm for matrix completion. *SIAM Journal on optimization*, 20(4):1956–1982, 2010.
- [4] G. Ceruti, J. Kusch, and C. Lubich. A rank-adaptive robust integrator for dynamical low-rank approximation. *BIT Numerical Mathematics*, 62(4):1149–1174, 2022.
- [5] G. Ceruti and C. Lubich. An unconventional robust integrator for dynamical low-rank approximation. *BIT Numerical Mathematics*, 62(1):23–44, 2022.
- [6] Y. Chen, D. Xiu, and X. Zhang. On enforcing non-negativity in polynomial approximations in high dimensions. *SIAM Journal on Scientific Computing*, 47(2):A866–A888, 2025.
- [7] A. Dektor and L. Einkemmer. Interpolatory dynamical low-rank approximation for the 3+ 3d Boltzmann-BGK equation. *Journal of Computational Physics*, 547:114515, 2026.
- [8] L. Einkemmer and I. Joseph. A mass, momentum, and energy conservative dynamical low-rank scheme for the Vlasov equation. *Journal of Computational Physics*, 443:110495, 2021.

- [9] L. Einkemmer, K. Kormann, J. Kusch, R. G. McClarren, and J.-M. Qiu. A review of low-rank methods for time-dependent kinetic simulations. *Journal of Computational Physics*, 538:114191, 2025.
- [10] S. A. Goreinov, I. V. Oseledets, D. V. Savostyanov, E. E. Tyrtyshnikov, and N. L. Zamarashkin. How to find a good submatrix. In V. Olshevsky and E. Tyrtyshnikov, editors, *Matrix Methods: Theory, Algorithms and Applications*, pages 247–256. World Scientific, 2010.
- [11] W. Guo, J. F. Ema, and J.-M. Qiu. A local macroscopic conservative (lomac) low rank tensor method with the discontinuous Galerkin method for the Vlasov dynamics. *Communications on Applied Mathematics and Computation*, 6(1):550–575, 2024.
- [12] W. Guo and J.-M. Qiu. A low rank tensor representation of linear transport and nonlinear Vlasov solutions and their associated flow maps. *Journal of Computational Physics*, 458:111089, 2022.
- [13] W. Guo and J.-M. Qiu. A conservative low rank tensor method for the Vlasov dynamics. *SIAM Journal on Scientific Computing*, 46(1):A232–A263, 2024.
- [14] W. Guo and J.-M. Qiu. A local macroscopic conservative (LoMaC) low rank tensor method for the Vlasov dynamics. *Journal of Scientific Computing*, 101(3):61, 2024.
- [15] N. Halko, P.-G. Martinsson, and J. A. Tropp. Finding structure with randomness: probabilistic algorithms for constructing approximate matrix decompositions. *SIAM Review*, 53(2):217–288, 2011.
- [16] O. Koch and C. Lubich. Dynamical low-rank approximation. *SIAM Journal on Matrix Analysis and Applications*, 29(2):434–454, 2007.
- [17] P.-L. Lions and B. Mercier. Splitting algorithms for the sum of two nonlinear operators. *SIAM Journal on Numerical Analysis*, 16(6):964–979, 1979.
- [18] C. Liu, D. Misis, C.-W. Shu, and X. Zhang. Efficient optimization-based invariant-domain-preserving limiters in solving gas dynamics equations. *Journal of Computational Physics*, 558:114839, 2026.
- [19] D. C. Liu and J. Nocedal. On the limited memory BFGS method for large scale optimization. *Mathematical Programming*, 45(1–3):503–528, 1989.
- [20] C. Lubich and I. V. Oseledets. A projector-splitting integrator for dynamical low-rank approximation. *BIT Numerical Mathematics*, 54(1):171–188, 2014.
- [21] Y. E. Nesterov. A method of solving a convex programming problem with convergence rate $o(\frac{1}{k^2})$. *Doklady Akademii Nauk SSSR*, 269(3):543–547, 1983.
- [22] Y. E. Nesterov. Gradient methods for minimizing composite functions. *Mathematical Programming*, 140(1):125–161, Aug 2013.
- [23] B. O’Donoghue and E. J. Candès. Adaptive restart for accelerated gradient schemes. *Foundations of Computational Mathematics*, 15(3):715–732, Jun 2015.

- [24] B. Perthame and C.-W. Shu. On positivity preserving finite volume schemes for Euler equations. *Numerische Mathematik*, 73(1):119–130, 1996.
- [25] M. J. D. Powell. Convergence properties of algorithms for nonlinear optimization. *SIAM Review*, 28(4):487–500, 1986.
- [26] G.-J. Song, M. K. Ng, and T.-X. Jiang. Tangent space based alternating projections for nonnegative low rank matrix approximation. *IEEE Transactions on Knowledge and Data Engineering*, 35(11):11917–11934, 2023.
- [27] J. Sun and J. Zhang. Global convergence of conjugate gradient methods without line search. *Annals of Operations Research*, 103(1–4):161–173, Mar 2001.
- [28] X. Tang, R. Dwaraknath, and L. Ying. Variational inference and density estimation with non-negative tensor train. *arXiv preprint arXiv:2507.21519*, 2025.
- [29] B. Vandereycken. Low-rank matrix completion by Riemannian optimization. *SIAM Journal on Optimization*, 23(2):1214–1236, 2013.
- [30] K. Wu, X. Zhang, and C.-W. Shu. High order numerical methods preserving invariant domain for hyperbolic and related systems. *arXiv preprint arXiv:2512.09116*, 2025.
- [31] E. Ye and N. F. G. Loureiro. Quantized tensor networks for solving the Vlasov–Maxwell equations. *Journal of Plasma Physics*, 90(3):805900301, 2024.
- [32] X. Zhang and C.-W. Shu. On positivity-preserving high order discontinuous Galerkin schemes for compressible Euler equations on rectangular meshes. *Journal of Computational Physics*, 229(23):8918–8934, 2010.
- [33] N. Zheng, D. Hayes, A. Christlieb, and J.-M. Qiu. A Semi-Lagrangian adaptive-rank (SLAR) method for linear advection and nonlinear Vlasov-Poisson system. *Journal of Computational Physics*, 532:113970, 2025.
- [34] N. Zheng, W. A. Sands, D. Hayes, A. J. Christlieb, and J.-M. Qiu. A semi-Lagrangian adaptive rank (SLAR) method for high-dimensional Vlasov dynamics. *SIAM Journal on Scientific Computing*, page accepted, 2026.
- [35] S. Zheng, W. Huang, B. Vandereycken, and X. Zhang. Riemannian optimization using three different metrics for Hermitian PSD fixed-rank constraints. *Computational Optimization and Applications*, 91(3):1135–1184, 2025.