
Dissipativity Theory for Nesterov’s Accelerated Method

Bin Hu¹ Laurent Lessard¹

Abstract

In this paper, we adapt the control theoretic concept of dissipativity theory to provide a natural understanding of Nesterov’s accelerated method. Our theory ties rigorous convergence rate analysis to the physically intuitive notion of energy dissipation. Moreover, dissipativity allows one to efficiently construct Lyapunov functions (either numerically or analytically) by solving a small semidefinite program. Using novel supply rate functions, we show how to recover known rate bounds for Nesterov’s method and we generalize the approach to certify both linear and sublinear rates in a variety of settings. Finally, we link the continuous-time version of dissipativity to recent works on algorithm analysis that use discretizations of ordinary differential equations.

1. Introduction

Nesterov’s accelerated method (Nesterov, 2003) has garnered attention and interest in the machine learning community because of its fast global convergence rate guarantees. The original convergence rate proofs of Nesterov’s accelerated method are derived using the method of estimate sequences, which has proven difficult to interpret. This observation motivated a sequence of recent works on new analysis and interpretations of Nesterov’s accelerated method (Bubeck et al., 2015; Lessard et al., 2016; Su et al., 2016; Drusvyatskiy et al., 2016; Flammarion & Bach, 2015; Wibisono et al., 2016; Wilson et al., 2016).

Many of these recent papers rely on Lyapunov-based stability arguments. Lyapunov theory is an analogue to the principle of minimum energy and brings a physical intuition to convergence behaviors. When applying such proof techniques, one must construct a *Lyapunov function*, which is a nonnegative function of the algorithm’s state (an “internal energy”) that decreases along all admissible trajec-

tories. Once a Lyapunov function is found, one can relate the rate of decrease of this internal energy to the rate of convergence of the algorithm. The main challenge in applying Lyapunov’s method is finding a suitable Lyapunov function.

There are two main approaches for Lyapunov function constructions. The first approach adopts the integral quadratic constraint (IQC) framework (Megretski & Rantzer, 1997) from control theory and formulates a linear matrix equality (LMI) whose feasibility implies the linear convergence of the algorithm (Lessard et al., 2016). Despite the generality of the IQC approach and the small size of the associated LMI, one must typically resort to numerical simulations to solve the LMI. The second approach seeks an ordinary differential equation (ODE) that can be appropriately discretized to yield the algorithm of interest. One can then gain intuition about the trajectories of the algorithm by examining trajectories of the continuous-time ODE (Su et al., 2016; Wibisono et al., 2016; Wilson et al., 2016). The work of Wilson et al. (2016) also establishes a general equivalence between Lyapunov functions and estimate sequence proofs.

In this paper, we bridge the IQC approach (Lessard et al., 2016) and the discretization approach (Wilson et al., 2016) by using dissipativity theory (Willems, 1972a;b). The term “dissipativity” is borrowed from the notion of energy dissipation in physics and the theory provides a general approach for the intuitive understanding and construction of Lyapunov functions. Dissipativity for quadratic Lyapunov functions in particular (Willems, 1972b) has seen widespread use in controls. In the sequel, we tailor dissipativity theory to the automated construction of Lyapunov functions, which are not necessarily quadratic, for the analysis of optimization algorithms. Our dissipation inequality leads to an LMI condition that is simpler than the one in Lessard et al. (2016) and hence more amenable to being solved analytically. When specialized to Nesterov’s accelerated method, our LMI recovers the Lyapunov function proposed in Wilson et al. (2016). Finally, we extend our LMI-based approach to the sublinear convergence analysis of Nesterov’s accelerated method in both discrete and continuous time domains. This complements the original LMI-based approach in Lessard et al. (2016), which mainly handles linear convergence rate analyses.

¹University of Wisconsin–Madison, Madison, WI 53706, USA. Correspondence to: Bin Hu <bhu38@wisc.edu>.

An LMI-based approach for sublinear rate analysis similar to ours was independently and simultaneously proposed by Fazlyab et al. (2017). While this work and the present work both draw connections to the continuous-time results mentioned above, different algorithms and function classes are emphasized. For example, Fazlyab et al. (2017) develops LMIs for gradient descent and proximal/projection-based variants with convex/quasi-convex objective functions. In contrast, the present work develops LMIs tailored to the analysis of discrete-time accelerated methods and Nesterov's method in particular.

2. Preliminaries

2.1. Notation

Let $\mathbb{R}$ and $\mathbb{R}_+$ denote the real and nonnegative real numbers, respectively. Let I_p and 0_p denote the $p \times p$ identity and zero matrices, respectively. The Kronecker product of two matrices is denoted $A \otimes B$ and satisfies the properties $(A \otimes B)^\top = A^\top \otimes B^\top$ and $(A \otimes B)(C \otimes D) = (AC) \otimes (BD)$ when the matrices have compatible dimensions. Matrix inequalities hold in the semidefinite sense unless otherwise indicated. A differentiable function $f : \mathbb{R}^p \rightarrow \mathbb{R}$ is m -strongly convex if $f(x) \geq f(y) + \nabla f(y)^\top(x - y) + \frac{m}{2}\|x - y\|^2$ for all $x, y \in \mathbb{R}^p$ and is L -smooth if $\|\nabla f(x) - \nabla f(y)\| \leq L\|x - y\|$ for all $x, y \in \mathbb{R}^p$. Note that f is convex if f is 0-strongly convex. We use x_* to denote a point satisfying $\nabla f(x_*) = 0$. When f is L -smooth and m -strongly convex, x_* is unique.

2.2. Classical Dissipativity Theory

Consider a linear dynamical system governed by the state-space model

$$\xi_{k+1} = A\xi_k + Bw_k. \quad (1)$$

Here, $\xi_k \in \mathbb{R}^{n_\xi}$ is the state, $w_k \in \mathbb{R}^{n_w}$ is the input, $A \in \mathbb{R}^{n_\xi \times n_\xi}$ is the state transition matrix, and $B \in \mathbb{R}^{n_\xi \times n_w}$ is the input matrix. The input w_k can be physically interpreted as a *driving force*. Classical dissipativity theory describes how the internal energy stored in the state ξ_k evolves with time k as one applies the input w_k to drive the system. A key concept in dissipativity theory is the *supply rate*, which characterizes the energy change in ξ_k due to the driving force w_k . The supply rate is a function $S : \mathbb{R}^{n_\xi} \times \mathbb{R}^{n_w} \rightarrow \mathbb{R}$ that maps any state/input pair (ξ, w) to a scalar measuring the amount of energy delivered from w to state ξ . Now we introduce the notion of dissipativity.

Definition 1 *The dynamical system (1) is dissipative with respect to the supply rate S if there exists a function $V : \mathbb{R}^{n_\xi} \rightarrow \mathbb{R}_+$ such that $V(\xi) \geq 0$ for all $\xi \in \mathbb{R}^{n_\xi}$ and*

$$V(\xi_{k+1}) - V(\xi_k) \leq S(\xi_k, w_k) \quad (2)$$

for all k . The function V is called a storage function, which quantifies the energy stored in the state ξ . In addition, (2) is called the dissipation inequality.

The dissipation inequality (2) states that the change of the internal energy stored in ξ_k is equal to the difference between the supplied energy and the dissipated energy. Since there will always be some energy dissipating from the system, the change in the stored energy (which is exactly $V(\xi_{k+1}) - V(\xi_k)$) is always bounded above by the energy supplied to the system (which is exactly $S(\xi_k, w_k)$). A variant of (2) known as the exponential dissipation inequality states that for some $0 \leq \rho < 1$, we have

$$V(\xi_{k+1}) - \rho^2 V(\xi_k) \leq S(\xi_k, w_k), \quad (3)$$

which states that at least a fraction $(1 - \rho^2)$ of the internal energy will dissipate at every step.

The dissipation inequality (3) provides a direct way to construct a Lyapunov function based on the storage function. It is often the case that we have prior knowledge about how the driving force w_k is related to the state ξ_k . Thus, we may know additional information about the supply rate function $S(\xi_k, w_k)$. For example, if $S(\xi_k, w_k) \leq 0$ for all k then (3) directly implies that $V(\xi_{k+1}) \leq \rho^2 V(\xi_k)$, and the storage function V can serve as a Lyapunov function. The condition $S(\xi_k, w_k) \leq 0$ means that the driving force w_k does not inject any energy into the system and may even extract energy out of the system. Then, the internal energy will decrease no slower than the linear rate ρ^2 and approach a minimum value at equilibrium.

An advantage of dissipativity theory is that for any quadratic supply rate, one can automatically construct the dissipation inequality using semidefinite programming. We now state a standard result from the controls literature.

Theorem 2 *Consider the following quadratic supply rate with $X \in \mathbb{R}^{(n_\xi + n_w) \times (n_\xi + n_w)}$ and $X = X^\top$.*

$$S(\xi, w) := \begin{bmatrix} \xi \\ w \end{bmatrix}^\top X \begin{bmatrix} \xi \\ w \end{bmatrix}. \quad (4)$$

If there exists a matrix $P \in \mathbb{R}^{n_\xi \times n_\xi}$ with $P \geq 0$ such that

$$\begin{bmatrix} A^\top P A - \rho^2 P & A^\top P B \\ B^\top P A & B^\top P B \end{bmatrix} - X \leq 0, \quad (5)$$

then the dissipation inequality (3) holds for all trajectories of (1) with $V(\xi) := \xi^\top P \xi$.

Proof. Based on the state-space model (1), we have

$$\begin{aligned} V(\xi_{k+1}) &= \xi_{k+1}^\top P \xi_{k+1} \\ &= (A\xi_k + Bw_k)^\top P (A\xi_k + Bw_k) \\ &= \begin{bmatrix} \xi_k \\ w_k \end{bmatrix}^\top \begin{bmatrix} A^\top P A & A^\top P B \\ B^\top P A & B^\top P B \end{bmatrix} \begin{bmatrix} \xi_k \\ w_k \end{bmatrix}. \end{aligned}$$

Hence we can left and right multiply (5) by $\begin{bmatrix} \xi_k^\top & w_k^\top \end{bmatrix}$ and $\begin{bmatrix} \xi_k^\top & w_k^\top \end{bmatrix}^\top$, and directly obtain the desired conclusion. ■

The left-hand side of (5) is linear in P , so (5) is a *linear matrix inequality* (LMI) for any fixed A, B, X, ρ . The set of P such that (5) holds is therefore a convex set and can be efficiently searched using interior point methods, for example. To apply the dissipativity theory for linear convergence rate analysis, one typically follows two steps.

1. Choose a proper quadratic supply rate function S satisfying certain desired properties, e.g. $S(\xi_k, w_k) \leq 0$.
2. Solve the LMI (5) to obtain a storage function V , which is then used to construct a Lyapunov function.

In step 2, the LMI obtained is typically very small, e.g. 2×2 or 3×3 , so we can often solve the LMI analytically. For illustrative purposes, we rephrase the existing LMI analysis of the gradient descent method (Lessard et al., 2016, §4.4) using the notion of dissipativity.

2.3. Example: Dissipativity for Gradient Descent

There is an intrinsic connection between dissipativity theory and the IQC approach (Megretski et al., 2010; Seiler, 2015). The IQC analysis of the gradient descent method in Lessard et al. (2016) may be reframed using dissipativity theory. Then, the pointwise IQC (Lessard et al., 2016, Lemma 6) amounts to using a quadratic supply rate S with $S \leq 0$. Specifically, assume f is L -smooth and m -strongly convex, and consider the gradient descent method

$$x_{k+1} = x_k - \alpha \nabla f(x_k). \quad (6)$$

We have $x_{k+1} - x_\star = x_k - x_\star - \alpha \nabla f(x_k)$, where $x_\star$ is the unique point satisfying $\nabla f(x_\star) = 0$. Define $\xi_k := x_k - x_\star$ and $w_k := \nabla f(x_k)$. Then the gradient descent method is modeled by (1) with $A := I_p$ and $B := -\alpha I_p$. Since $w_k = \nabla f(\xi_k + x_\star)$, we can define the following quadratic supply rate

$$S(\xi_k, w_k) = \begin{bmatrix} \xi_k \\ w_k \end{bmatrix}^\top \begin{bmatrix} 2mLI_p & -(m+L)I_p \\ -(m+L)I_p & 2I_p \end{bmatrix} \begin{bmatrix} \xi_k \\ w_k \end{bmatrix} \quad (7)$$

By co-coercivity, we have $S(\xi_k, w_k) \leq 0$ for all k . This just restates Lessard et al. (2016, Lemma 6). Then, we can directly apply Theorem 2 to construct the dissipation inequality. We can parameterize $P = p \otimes I_p$ and define the storage function as $V(\xi_k) = p \|\xi_k\|^2 = p \|x_k - x_\star\|^2$. The LMI (5) becomes

$$\left(\begin{bmatrix} (1-\rho^2)p & -\alpha p \\ -\alpha p & \alpha^2 p \end{bmatrix} + \begin{bmatrix} -2mL & m+L \\ m+L & -2 \end{bmatrix} \right) \otimes I_p \leq 0.$$

Hence for any $0 \leq \rho < 1$, we have $p \|x_{k+1} - x_\star\|^2 \leq \rho^2 p \|x_k - x_\star\|^2$ if there exists $p \geq 0$ such that

$$\begin{bmatrix} (1-\rho^2)p & -\alpha p \\ -\alpha p & \alpha^2 p \end{bmatrix} + \begin{bmatrix} -2mL & m+L \\ m+L & -2 \end{bmatrix} \leq 0 \quad (8)$$

The LMI (8) is simple and can be analytically solved to recover the existing rate results for the gradient descent method. For example, we can choose (α, ρ, p) to be $(\frac{1}{L}, 1 - \frac{m}{L}, L^2)$ or $(\frac{2}{L+m}, \frac{L-m}{L+m}, \frac{1}{2}(L+m)^2)$ to immediately recover the standard rate results in Polyak (1987).

Based on the example above, it is evident that choosing a proper supply rate is critical for the construction of a Lyapunov function. The supply rate (7) turns out to be inadequate for the analysis of Nesterov's accelerated method. For Nesterov's accelerated method, the dependence between the internal energy and the driving force is more complicated due to the presence of momentum terms. We will next develop a new supply rate that captures this complicated dependence. We will also make use of this new supply rate to recover the standard linear rate results for Nesterov's accelerated method.

3. Dissipativity for Accelerated Linear Rates

3.1. Dissipativity for Nesterov's Method

Suppose f is L -smooth and m -strongly convex with $m > 0$. Let $x_\star$ be the unique point satisfying $\nabla f(x_\star) = 0$. Now we consider Nesterov's accelerated method, which uses the following iteration rule to find $x_\star$:

$$x_{k+1} = y_k - \alpha \nabla f(y_k), \quad (9a)$$

$$y_k = (1 + \beta)x_k - \beta x_{k-1}. \quad (9b)$$

We can rewrite (9) as

$$\begin{bmatrix} x_{k+1} - x_\star \\ x_k - x_\star \end{bmatrix} = A \begin{bmatrix} x_k - x_\star \\ x_{k-1} - x_\star \end{bmatrix} + B w_k \quad (10)$$

where $w_k := \nabla f(y_k) = \nabla f((1 + \beta)x_k - \beta x_{k-1})$. Also, $A := \tilde{A} \otimes I_p$, $B := \tilde{B} \otimes I_p$, and $\tilde{A}, \tilde{B}$ are defined by

$$\tilde{A} := \begin{bmatrix} 1 + \beta & -\beta \\ 1 & 0 \end{bmatrix}, \quad \tilde{B} := \begin{bmatrix} -\alpha \\ 0 \end{bmatrix}. \quad (11)$$

Hence, Nesterov's accelerated method (9) is in the form of (1) with $\xi_k = \begin{bmatrix} (x_k - x_\star)^\top & (x_{k-1} - x_\star)^\top \end{bmatrix}^\top$.

Nesterov's accelerated method can improve the convergence rate since the input w_k depends on both x_k and x_{k-1} , and drives the state in a specific direction, i.e. along $(1 + \beta)x_k - \beta x_{k-1}$. This leads to a supply rate that extracts energy out of the system significantly faster than with gradient descent. This is formally stated in the next lemma.

Lemma 3 Let f be L -smooth and m -strongly convex with $m > 0$. Let x_* be the unique point satisfying $\nabla f(x_*) = 0$. Consider Nesterov's method (9) or equivalently (10). The following inequalities hold for all trajectories.

$$\begin{aligned} \begin{bmatrix} x_k - x_* \\ x_{k-1} - x_* \\ \nabla f(y_k) \end{bmatrix}^\top X_1 \begin{bmatrix} x_k - x_* \\ x_{k-1} - x_* \\ \nabla f(y_k) \end{bmatrix} &\leq f(x_k) - f(x_{k+1}) \\ \begin{bmatrix} x_k - x_* \\ x_{k-1} - x_* \\ \nabla f(y_k) \end{bmatrix}^\top X_2 \begin{bmatrix} x_k - x_* \\ x_{k-1} - x_* \\ \nabla f(y_k) \end{bmatrix} &\leq f(x_*) - f(x_{k+1}) \end{aligned}$$

where $X_i = \tilde{X}_i \otimes I_p$ for $i = 1, 2$, and $\tilde{X}_i$ are defined by

$$\tilde{X}_1 := \frac{1}{2} \begin{bmatrix} \beta^2 m & -\beta^2 m & -\beta \\ -\beta^2 m & \beta^2 m & \beta \\ -\beta & \beta & \alpha(2 - L\alpha) \end{bmatrix} \quad (12)$$

$$\tilde{X}_2 := \frac{1}{2} \begin{bmatrix} (1 + \beta)^2 m & -\beta(1 + \beta)m & -(1 + \beta) \\ -\beta(1 + \beta)m & \beta^2 m & \beta \\ -(1 + \beta) & \beta & \alpha(2 - L\alpha) \end{bmatrix} \quad (13)$$

Given any $0 \leq \rho \leq 1$, one can define the supply rate as (4) with a particular choice of $X := \rho^2 X_1 + (1 - \rho^2) X_2$. Then this supply rate satisfies the condition

$$S(\xi_k, w_k) \leq \rho^2(f(x_k) - f(x_*)) - (f(x_{k+1}) - f(x_*)). \quad (14)$$

Proof. The proof is similar to the proof of (3.23)–(3.24) in Bubeck (2015), but Bubeck (2015, Lemma 3.6) must be modified to account for the strong convexity of f . See the supplementary material for a detailed proof. ■

The supply rate (14) captures how the driving force w_k is impacting the future state x_{k+1} . The physical interpretation is that there is some amount of hidden energy in the system that takes the form of $f(x_k) - f(x_*)$. The supply rate condition (14) describes how the driving force w_k is coupled with the hidden energy in the future. It says the delivered energy is bounded by a weighted decrease of the hidden energy. Based on this supply rate, one can search Lyapunov function using the following theorem.

Theorem 4 Let f be L -smooth and m -strongly convex with $m > 0$. Let x_* be the unique point satisfying $\nabla f(x_*) = 0$. Consider Nesterov's accelerated method (9). For any rate $0 \leq \rho < 1$, set $\tilde{X} := \rho^2 \tilde{X}_1 + (1 - \rho^2) \tilde{X}_2$ where $\tilde{X}_1$ and $\tilde{X}_2$ are defined in (12)–(13). In addition, let $\tilde{A}, \tilde{B}$ be defined by (11). If there exists a matrix $\tilde{P} \in \mathbb{R}^{2 \times 2}$ with $\tilde{P} \geq 0$ such that

$$\begin{bmatrix} \tilde{A}^\top \tilde{P} \tilde{A} - \rho^2 \tilde{P} & \tilde{A}^\top \tilde{P} \tilde{B} \\ \tilde{B}^\top \tilde{P} \tilde{A} & \tilde{B}^\top \tilde{P} \tilde{B} \end{bmatrix} - \tilde{X} \leq 0 \quad (15)$$

then set $P := \tilde{P} \otimes I_p$ and define the Lyapunov function

$$\mathcal{V}_k := \begin{bmatrix} x_k - x_* \\ x_{k-1} - x_* \end{bmatrix}^\top P \begin{bmatrix} x_k - x_* \\ x_{k-1} - x_* \end{bmatrix} + f(x_k) - f(x_*), \quad (16)$$

which satisfies $\mathcal{V}_{k+1} \leq \rho^2 \mathcal{V}_k$ for all k . Moreover, we have $f(x_k) - f(x_*) \leq \rho^{2k} \mathcal{V}_0$ for Nesterov's method.

Proof. Take the Kronecker product of (15) and I_p , and hence (5) holds with $A := \tilde{A} \otimes I_p$, $B := \tilde{B} \otimes I_p$, and $X := \tilde{X} \otimes I_p$. Let the supply rate S be defined by (4). Then, define the quadratic storage function $V(\xi_k) := \xi_k^\top P \xi_k$ and apply Theorem 2 to show $V(\xi_{k+1}) - \rho^2 V(\xi_k) \leq S(\xi_k, w_k)$. Based on the supply rate condition (14), we can define the Lyapunov function $\mathcal{V}_k := V(\xi_k) + f(x_k) - f(x_*)$ and show $\mathcal{V}_{k+1} \leq \rho^2 \mathcal{V}_k$. Finally, since $P \geq 0$, we have $f(x_k) - f(x_*) \leq \rho^{2k} \mathcal{V}_0$. ■

We can immediately recover the proposed Lyapunov function in Wilson et al. (2016, Theorem 6) by setting $\tilde{P}$ to

$$\tilde{P} = \begin{bmatrix} \sqrt{\frac{L}{2}} \\ \sqrt{\frac{m}{2}} - \sqrt{\frac{L}{2}} \end{bmatrix} \begin{bmatrix} \sqrt{\frac{L}{2}} & \sqrt{\frac{m}{2}} - \sqrt{\frac{L}{2}} \end{bmatrix}. \quad (17)$$

Clearly $\tilde{P} \geq 0$. Now define $\kappa := \frac{L}{m}$. Given $\alpha = \frac{1}{L}$, $\beta = \frac{\sqrt{\kappa}-1}{\sqrt{\kappa}+1}$, and $\rho^2 = 1 - \sqrt{\frac{m}{L}}$, it is straightforward to verify that the left side of the LMI (5) is equal to

$$\frac{m(\sqrt{\kappa}-1)^3}{2(\kappa+\sqrt{\kappa})} \begin{bmatrix} -1 & 1 & 0 \\ 1 & -1 & 0 \\ 0 & 0 & 0 \end{bmatrix},$$

which is clearly negative semidefinite. Hence we can immediately construct a Lyapunov function using (16) to prove the linear rate $\rho^2 = 1 - \sqrt{\frac{m}{L}}$.

Searching for analytic certificates such as (17) can either be carried out by directly analyzing the LMI, or by using numerical solutions to guide the search. For example, numerically solving (15) for any fixed L and m directly yields (17), which makes finding the analytic expression easy.

3.2. Dissipativity Theory for More General Methods

We demonstrate the generality of the dissipativity theory on a more general variant of Nesterov's method. Consider a modified accelerated method

$$x_{k+1} = (1 + \beta)x_k - \beta x_{k-1} - \alpha \nabla f(y_k), \quad (18a)$$

$$y_k = (1 + \eta)x_k - \eta x_{k-1}. \quad (18b)$$

When $\beta = \eta$, we recover Nesterov's accelerated method. When $\eta = 0$, we recover the Heavy-ball method of Polyak

(1987). We can rewrite (18) in state-space form (10) where $w_k := \nabla f(y_k) = \nabla f((1 + \eta)x_k - \eta x_{k-1})$, $A := \tilde{A} \otimes I_p$, $B := \tilde{B} \otimes I_p$, and $\tilde{A}, \tilde{B}$ are defined by

$$\tilde{A} := \begin{bmatrix} 1 + \beta & -\beta \\ 1 & 0 \end{bmatrix}, \quad \tilde{B} := \begin{bmatrix} -\alpha \\ 0 \end{bmatrix}.$$

Lemma 5 *Let f be L -smooth and m -strongly convex with $m > 0$. Let x_* be the unique point satisfying $\nabla f(x_*) = 0$. Consider the general accelerated method (18). Define the state $\xi_k := [(x_k - x_*)^\top \ (x_{k-1} - x_*)^\top]^\top$ and the input $w_k := \nabla f(y_k) = \nabla f((1 + \eta)x_k - \eta x_{k-1})$. Then the following inequalities hold for all trajectories.*

$$\begin{bmatrix} \xi_k \\ w_k \end{bmatrix}^\top (X_1 + X_2) \begin{bmatrix} \xi_k \\ w_k \end{bmatrix} \leq f(x_k) - f(x_{k+1}) \quad (19)$$

$$\begin{bmatrix} \xi_k \\ w_k \end{bmatrix}^\top (X_1 + X_3) \begin{bmatrix} \xi_k \\ w_k \end{bmatrix} \leq f(x_*) - f(x_{k+1}) \quad (20)$$

with $X_i := \tilde{X}_i \otimes I_p$ for $i = 1, 2, 3$, and $\tilde{X}_i$ are defined by

$$\tilde{X}_1 := \frac{1}{2} \begin{bmatrix} -L\delta^2 & L\delta^2 & -(1 - L\alpha)\delta \\ L\delta^2 & -L\delta^2 & (1 - L\alpha)\delta \\ -(1 - L\alpha)\delta & (1 - L\alpha)\delta & \alpha(2 - L\alpha) \end{bmatrix}$$

$$\tilde{X}_2 := \frac{1}{2} \begin{bmatrix} \eta^2 m & -\eta^2 m & -\eta \\ -\eta^2 m & \eta^2 m & \eta \\ -\eta & \eta & 0 \end{bmatrix}$$

$$\tilde{X}_3 := \frac{1}{2} \begin{bmatrix} (1 + \eta)^2 m & -\eta(1 + \eta)m & -(1 + \eta) \\ -\eta(1 + \eta)m & \eta^2 m & \eta \\ -(1 + \eta) & \eta & 0 \end{bmatrix}$$

with $\delta := \beta - \eta$. In addition, one can define the supply rate as (4) with $X := X_1 + \rho^2 X_2 + (1 - \rho^2) X_3$. Then for all trajectories (ξ_k, w_k) of the general accelerated method (18), this supply rate satisfies the inequality

$$S(\xi_k, w_k) \leq \rho^2 (f(x_k) - f(x_*)) - (f(x_{k+1}) - f(x_*)). \quad (21)$$

Proof. A detailed proof is presented in the supplementary material. One mainly needs to modify the proof by taking the difference between β and η into accounts. ■

Based the supply rate (21), we can immediately modify Theorem 4 to handle the more general algorithm (18). Although we do not have general analytical formulas for the convergence rate of (18), preliminary numerical results suggest that there are a family of (α, β, η) leading to the rate $\rho^2 = 1 - \sqrt{\frac{m}{L}}$, and the required value of $\tilde{P}$ is quite different from (17). This indicates that our proposed LMI approach could go beyond the Lyapunov function (17).

Remark 6 *It is noted in Lessard et al. (2016, §3.2) that searching over combinations of multiple IQCs may yield*

improved rate bounds. The same is true of supply rates. For example, we could include $\lambda_1, \lambda_2 \geq 0$ as decision variables and search for a dissipation inequality with supply rate $\lambda_1 S_1 + \lambda_2 S_2$ where e.g. S_1 is (7) and S_2 is (21).

4. Dissipativity for Sublinear Rates

The LMI approach in (Lessard et al., 2016) is tailored for the analysis of linear convergence rates for algorithms that are time-invariant (the A and B matrices in (1) do not change with k). We now show that dissipativity theory can be used to analyze the sublinear rates $O(1/k)$ and $O(1/k^2)$ via slight modifications of the dissipation inequality.

4.1. Dissipativity for $O(1/k)$ rates

The $O(1/k)$ modification, which we present first, is very similar to the linear rate result.

Theorem 7 *Suppose f has a finite minimum f_* . Consider the LTI system (1) with a supply rate satisfying*

$$S(\xi_k, w_k) \leq -(f(z_k) - f_*) \quad (22)$$

for some sequence $\{z_k\}$. If there exists a nonnegative storage function V such that the dissipation inequality (2) holds over all trajectories of (ξ_k, w_k) , then the following inequality holds over all trajectories as well.

$$\sum_{k=0}^T (f(z_k) - f_*) \leq V(\xi_0). \quad (23)$$

In addition, we have the sublinear convergence rate

$$\min_{k:k \leq T} (f(z_k) - f_*) \leq \frac{V(\xi_0)}{T + 1}. \quad (24)$$

If $f(z_{k+1}) \leq f(z_k)$ for all k , then (24) implies that $f(z_k) - f_ \leq \frac{V(\xi_0)}{k+1}$ for all k .*

Proof. By the supply rate condition (22) and the dissipation inequality (2), we immediately get

$$V(\xi_{k+1}) - V(\xi_k) + f(z_k) - f_* \leq 0.$$

Summing the above inequality from $k = 0$ to T and using $V \geq 0$ yields the desired result. ■

To address the sublinear rate analysis, the critical step is to choose an appropriate supply rate. If f is L -smooth and convex, this is easily done. Consider the gradient method (6) and define the quantities $\xi_k := x_k - x_*$, $A := I_p$, and $B := -\alpha I_p$ as in Section 2.3. Since f is L -smooth and convex, define the quadratic supply rate

$$S(\xi_k, w_k) := \begin{bmatrix} \xi_k \\ w_k \end{bmatrix}^\top \begin{bmatrix} 0_p & -\frac{1}{2} I_p \\ -\frac{1}{2} I_p & \frac{1}{2L} I_p \end{bmatrix} \begin{bmatrix} \xi_k \\ w_k \end{bmatrix},$$

which satisfies $S(\xi_k, w_k) \leq f_\star - f(x_k)$ for all k (co-coercivity). Then we can directly apply the LMI (5) with $\rho = 1$ to construct the dissipation inequality. Setting $P = p \otimes I_p$ and defining the storage function as $V(\xi_k) := p\|\xi_k\|^2 = p\|x_k - x_\star\|^2$, the LMI (5) becomes

$$\left(\begin{bmatrix} 0 & -\alpha p \\ -\alpha p & \alpha^2 p \end{bmatrix} + \begin{bmatrix} 0 & \frac{1}{2} \\ \frac{1}{2} & -\frac{1}{2L} \end{bmatrix} \right) \otimes I_p \leq 0,$$

which is equivalent to

$$\begin{bmatrix} 0 & -\alpha p + \frac{1}{2} \\ -\alpha p + \frac{1}{2} & \alpha^2 p - \frac{1}{2L} \end{bmatrix} \leq 0. \quad (25)$$

Due to the $(1, 1)$ entry being zero, (25) holds if and only if

$$\begin{cases} -\alpha p + \frac{1}{2} = 0 \\ \alpha^2 p - \frac{1}{2L} \leq 0 \end{cases} \implies \begin{cases} p = \frac{1}{2\alpha} \\ \alpha \leq \frac{1}{L} \end{cases}$$

We can choose $\alpha = \frac{1}{L}$ and the bound (24) becomes

$$\min_{k \leq T} (f(x_k) - f_\star) \leq \frac{L\|x_0 - x_\star\|^2}{2(T+1)}.$$

Since gradient descent has monotonically nonincreasing iterates, that is $f(x_{k+1}) \leq f(x_k)$ for all k , we immediately recover the standard $O(1/k)$ rate result.

4.2. Dissipativity for $O(1/k^2)$ rates

Certifying a $O(1/k)$ rate for the gradient method required solving a single LMI (25). However, this is not the case for the $O(1/k^2)$ rate analysis of Nesterov's accelerated method. Nesterov's algorithm has parameters that depend on k so the analysis is more involved. We will begin with the general case and then specialize to Nesterov's algorithm. Consider the dynamical system

$$\xi_{k+1} = A_k \xi_k + B_k w_k \quad (26)$$

The state matrix A_k and input matrix B_k change with the time step k , and hence (26) is referred to as a "linear time-varying" (LTV) system. The analysis of LTV systems typically requires a time-dependent supply rate such as

$$S_k(\xi_k, w_k) := \begin{bmatrix} \xi_k \\ w_k \end{bmatrix}^\top X_k \begin{bmatrix} \xi_k \\ w_k \end{bmatrix} \quad (27)$$

If there exists a sequence $\{P_k\}$ with $P_k \geq 0$ such that

$$\begin{bmatrix} A_k^\top P_{k+1} A_k - P_k & A_k^\top P_{k+1} B_k \\ B_k^\top P_{k+1} A_k & B_k^\top P_{k+1} B_k \end{bmatrix} - X_k \leq 0 \quad (28)$$

for all k , then we have $V_{k+1}(\xi_{k+1}) - V_k(\xi_k) \leq S_k(\xi_k, w_k)$ with the time-dependent storage function defined as $V_k(\xi_k) := \xi_k^\top P_k \xi_k$. This is a standard approach for dissipation inequality constructions of LTV systems and can

be proved using the same proof technique in Theorem 2. Note that we need (28) to simultaneously hold for all k . This leads to an infinite number of LMIs in general.

Now we consider Nesterov's accelerated method for a convex L -smooth objective function f (Nesterov, 2003).

$$x_{k+1} = y_k - \alpha_k \nabla f(y_k), \quad (29a)$$

$$y_k = (1 + \beta_k)x_k - \beta_k x_{k-1}. \quad (29b)$$

It is known that (29) achieves a rate of $O(1/k^2)$ when $\alpha_k := 1/L$ and β_k is defined recursively as follows.

$$\zeta_{-1} = 0, \quad \zeta_{k+1} = \frac{1 + \sqrt{1 + 4\zeta_k^2}}{2}, \quad \beta_k = \frac{\zeta_{k-1} - 1}{\zeta_k}.$$

The sequence $\{\zeta_k\}$ satisfies $\zeta_k^2 - \zeta_k = \zeta_{k-1}^2$. We now present a dissipativity theory for the sublinear rate analysis of Nesterov's accelerated method. Rewrite (29) as

$$\begin{bmatrix} x_{k+1} - x_\star \\ x_k - x_\star \end{bmatrix} = A_k \begin{bmatrix} x_k - x_\star \\ x_{k-1} - x_\star \end{bmatrix} + B_k w_k \quad (30)$$

where $w_k := \nabla f(y_k) = \nabla f((1 + \beta_k)x_k - \beta_k x_{k-1})$, $A_k := \tilde{A}_k \otimes I_p$, $B_k := \tilde{B}_k \otimes I_p$, and $\tilde{A}_k, \tilde{B}_k$ are given by

$$\tilde{A}_k := \begin{bmatrix} 1 + \beta_k & -\beta_k \\ 1 & 0 \end{bmatrix}, \quad \tilde{B}_k := \begin{bmatrix} -\alpha_k \\ 0 \end{bmatrix}.$$

Hence, Nesterov's accelerated method (29) is in the form of (26) with $\xi_k := [(x_k - x_\star)^\top (x_{k-1} - x_\star)^\top]^\top$. The $O(1/k^2)$ rate analysis of Nesterov's method (29) requires the following time-dependent supply rate.

Lemma 8 *Let f be L -smooth and convex. Let $x_\star$ be a point satisfying $\nabla f(x_\star) = 0$. In addition, set $f_\star := f(x_\star)$. Consider Nesterov's method (29) or equivalently (30). The following inequalities hold for all trajectories and for all k .*

$$\begin{bmatrix} x_k - x_\star \\ x_{k-1} - x_\star \\ \nabla f(y_k) \end{bmatrix}^\top M_k \begin{bmatrix} x_k - x_\star \\ x_{k-1} - x_\star \\ \nabla f(y_k) \end{bmatrix} \leq f(x_k) - f(x_{k+1})$$

$$\begin{bmatrix} x_k - x_\star \\ x_{k-1} - x_\star \\ \nabla f(y_k) \end{bmatrix}^\top N_k \begin{bmatrix} x_k - x_\star \\ x_{k-1} - x_\star \\ \nabla f(y_k) \end{bmatrix} \leq f(x_\star) - f(x_{k+1})$$

where $M_k := \tilde{M}_k \otimes I_p$, $N_k := \tilde{N}_k \otimes I_p$, and $\tilde{M}_k, \tilde{N}_k$ are defined by

$$\tilde{M}_k := \begin{bmatrix} 0 & 0 & -\frac{1}{2}\beta_k \\ 0 & 0 & \frac{1}{2}\beta_k \\ -\frac{1}{2}\beta_k & \frac{1}{2}\beta_k & \frac{1}{2L} \end{bmatrix}, \quad (31)$$

$$\tilde{N}_k := \begin{bmatrix} 0 & 0 & -\frac{1}{2}(1 + \beta_k) \\ 0 & 0 & \frac{1}{2}\beta_k \\ -\frac{1}{2}(1 + \beta_k) & \frac{1}{2}\beta_k & \frac{1}{2L} \end{bmatrix}. \quad (32)$$

Given any nondecreasing sequence $\{\mu_k\}$, one can define the supply rate as (27) with the particular choice $X_k := \mu_k M_k + (\mu_{k+1} - \mu_k) N_k$ for all k . Then this supply rate satisfies the condition

$$S(\xi_k, w_k) \leq \mu_k (f(x_k) - f_\star) - \mu_{k+1} (f(x_{k+1}) - f_\star). \quad (33)$$

Proof. The proof is very similar to the proof of Lemma 3 with an extra condition $m = 0$. A detailed proof is presented in the supplementary material. ■

Theorem 9 Consider the LTV dynamical system (26). If there exist matrices $\{P_k\}$ with $P_k \geq 0$ and a nondecreasing sequence of nonnegative scalars $\{\mu_k\}$ such that

$$\begin{bmatrix} A_k^\top P_{k+1} A_k - P_k & A_k^\top P_{k+1} B_k \\ B_k^\top P_{k+1} A_k & B_k^\top P_{k+1} B_k \end{bmatrix} - \mu_k M_k - (\mu_{k+1} - \mu_k) N_k \leq 0 \quad (34)$$

then we have $V_{k+1}(\xi_{k+1}) - V_k(\xi_k) \leq S_k(\xi_k, w_k)$ with the storage function $V_k(\xi_k) := \xi_k^\top P_k \xi_k$ and the supply rate (27) using $X_k := \mu_k M_k + (\mu_{k+1} - \mu_k) N_k$ for all k . In addition, if this supply rate satisfies (33), we have

$$f(x_k) - f_\star \leq \frac{\mu_0 (f(x_0) - f_\star) + V_0(\xi_0)}{\mu_k} \quad (35)$$

Proof. Based on the state-space model (26), we can left and right multiply (34) by $[\xi_k^\top \ w_k^\top]$ and $[\xi_k^\top \ w_k^\top]^\top$, and directly obtain the dissipation inequality. Combining this dissipation inequality with (33), we can show

$$V_{k+1}(\xi_{k+1}) + \mu_{k+1} (f_{k+1} - f_\star) \leq V_k(\xi_k) + \mu_k (f_k - f_\star).$$

Summing the above inequality as in the proof of Theorem 7 and using the fact that $P_k \geq 0$ for all k yields the result. ■

We are now ready to show the $O(1/k^2)$ rate result for Nesterov's accelerated method. Set $\mu_k := (\zeta_{k-1})^2$ and $P_k := \frac{L}{2} \begin{bmatrix} \zeta_{k-1} \\ 1 - \zeta_{k-1} \end{bmatrix} \begin{bmatrix} \zeta_{k-1} & 1 - \zeta_{k-1} \end{bmatrix}$. Note that $P_k \geq 0$ and $\mu_{k+1} - \mu_k = \zeta_k$. It is straightforward to verify that this choice of $\{P_k, \mu_k\}$ makes the left side of (34) the zero matrix and hence (35) holds. Using the fact that $\zeta_{k-1} \geq k/2$ (easily proved by induction), we have $\mu_k \geq k^2/4$ and the $O(1/k^2)$ rate for Nesterov's method follows.

Remark 10 Theorem 9 is quite general. The infinite family of LMIs (34) can also be applied for linear rate analysis and collapses down to the single LMI (5) in that case. To apply (34) to linear rate analysis, one needs to slightly modify (31)–(32) such that the strong convexity parameter

m is incorporated into the formulas of M_k, N_k . By setting $\mu_k := \rho^{-2k}$ and $P_k := \rho^{-2k} P$, then the LMI (34) is the same for all k and we recover (5). This illustrates how the infinite number of LMIs (34) can collapse to a single LMI under special circumstances.

5. Continuous-time Dissipation Inequality

Finally, we briefly discuss dissipativity theory for the continuous-time ODEs used in optimization research. Note that dissipativity theory was first introduced in Willems (1972a;b) in the context of continuous-time systems. We denote continuous-time variables in upper case. Consider a continuous-time state-space model

$$\dot{\Lambda}(t) = A(t)\Lambda(t) + B(t)W(t) \quad (36)$$

where $\Lambda(t)$ is the state, $W(t)$ is the input, and $\dot{\Lambda}(t)$ denotes the time derivative of $\Lambda(t)$. In continuous-time, the supply rate is a function $S : \mathbb{R}^{n_\Lambda} \times \mathbb{R}^{n_W} \times \mathbb{R}_+ \rightarrow \mathbb{R}$ that assigns a scalar to each possible state and input pair. Here, we allow S to also depend on time $t \in \mathbb{R}_+$. To simplify our exposition, we will omit the explicit time dependence (t) from our notation.

Definition 11 The dynamical system (36) is dissipative with respect to the supply rate S if there exists a function $V : \mathbb{R}^{n_\Lambda} \times \mathbb{R}_+ \rightarrow \mathbb{R}_+$ such that $V(\Lambda, t) \geq 0$ for all $\Lambda \in \mathbb{R}^{n_\Lambda}$ and $t \geq 0$ and

$$\dot{V}(\Lambda, t) \leq S(\Lambda, W, t) \quad (37)$$

for every trajectory of (36). Here, $\dot{V}$ denotes the Lie derivative (or total derivative); it accounts for Λ 's dependence on t . The function V is called a storage function, and (37) is a (continuous-time) dissipation inequality.

For any given quadratic supply rate, one can automatically construct the continuous-time dissipation inequality using semidefinite programs. The following result is standard in the controls literature.

Theorem 12 Suppose $X(t) \in \mathbb{R}^{(n_\Lambda + n_W) \times (n_\Lambda + n_W)}$ and $X(t)^\top = X(t)$ for all t . Consider the quadratic supply rate

$$S(\Lambda, W, t) := \begin{bmatrix} \Lambda \\ W \end{bmatrix}^\top X \begin{bmatrix} \Lambda \\ W \end{bmatrix} \quad \text{for all } t. \quad (38)$$

If there exists a family of matrices $P(t) \in \mathbb{R}^{n_\Lambda \times n_\Lambda}$ with $P(t) \geq 0$ such that

$$\begin{bmatrix} A^\top P + PA + \dot{P} & PB \\ B^\top P & 0 \end{bmatrix} - X \leq 0 \quad \text{for all } t. \quad (39)$$

Then we have $\dot{V}(\Lambda, t) \leq S(\Lambda, W, t)$ with the storage function defined as $V(\Lambda, t) := \Lambda^\top P \Lambda$.

Proof. Based on the state-space model (36), we can apply the product rule for total derivatives and obtain

$$\begin{aligned}\dot{V}(\Lambda, t) &= \dot{\Lambda}^\top P \Lambda + \Lambda^\top P \dot{\Lambda} + \Lambda^\top \dot{P} \Lambda \\ &= \begin{bmatrix} \Lambda \\ W \end{bmatrix}^\top \begin{bmatrix} A^\top P + P A + \dot{P} & P B \\ B^\top P & 0 \end{bmatrix} \begin{bmatrix} \Lambda \\ W \end{bmatrix}.\end{aligned}$$

Hence we can left and right multiply (39) by $\begin{bmatrix} \Lambda^\top & W^\top \end{bmatrix}$ and $\begin{bmatrix} \Lambda^\top & W^\top \end{bmatrix}^\top$ and obtain the desired conclusion. ■

The algebraic structure of the LMI (39) is simpler than that of its discrete-time counterpart (5) because for given P , the continuous-time LMI is linear in A, B rather than being quadratic. This may explain why continuous-time ODEs are sometimes more amenable to analytic approaches than their discretized counterparts.

We demonstrate the utility of (39) on the continuous-time limit of Nesterov's accelerated method in Su et al. (2016):

$$\ddot{Y} + \frac{3}{t}\dot{Y} + \nabla f(Y) = 0, \quad (40)$$

which we rewrite as (36) with $\Lambda := \begin{bmatrix} \dot{Y}^\top & Y^\top - x_\star^\top \end{bmatrix}^\top$, $W := \nabla f(Y)$, $x_\star$ is a point satisfying $\nabla f(x_\star) = 0$, and A, B are defined by

$$A(t) := \begin{bmatrix} -\frac{3}{t}I_p & 0_p \\ I_p & 0_p \end{bmatrix}, \quad B(t) := \begin{bmatrix} -I_p \\ 0_p \end{bmatrix}.$$

Suppose f is convex and set $f_\star := f(x_\star)$. Su et al. (2016, Theorem 3) constructs the Lyapunov function $\mathcal{V}(Y, t) := t^2(f(Y) - f_\star) + 2\|Y - \frac{t}{2}\dot{Y} - x_\star\|^2$ to show that $\dot{\mathcal{V}} \leq 0$ and then directly demonstrate a $O(1/t^2)$ rate for the ODE (40). To illustrate the power of the dissipation inequality, we use the LMI (39) to recover this Lyapunov function. Denote $G(Y, t) := t^2(f(Y) - f_\star)$. Note that convexity implies $f(Y) - f_\star \leq \nabla f(Y)^\top (Y - x_\star)$, which we rewrite as

$$2t(f(Y) - f_\star) \leq \begin{bmatrix} \dot{Y} \\ Y - x_\star \\ W \end{bmatrix}^\top \begin{bmatrix} 0_p & 0_p & 0_p \\ 0_p & 0_p & tI_p \\ 0_p & tI_p & 0_p \end{bmatrix} \begin{bmatrix} \dot{Y} \\ Y - x_\star \\ W \end{bmatrix}.$$

Since $\dot{G}(Y, t) = 2t(f(Y) - f_\star) + t^2\nabla f(Y)^\top \dot{Y}$, we have

$$\dot{G} \leq \begin{bmatrix} \dot{Y} \\ Y - x_\star \\ W \end{bmatrix}^\top \begin{bmatrix} 0_p & 0_p & \frac{t^2}{2}I_p \\ 0_p & 0_p & tI_p \\ \frac{t^2}{2}I_p & tI_p & 0_p \end{bmatrix} \begin{bmatrix} \dot{Y} \\ Y - x_\star \\ W \end{bmatrix}.$$

Now choose the supply rate S as (38) with $X(t)$ given by

$$X(t) := - \begin{bmatrix} 0_p & 0_p & \frac{t^2}{2}I_p \\ 0_p & 0_p & tI_p \\ \frac{t^2}{2}I_p & tI_p & 0_p \end{bmatrix}.$$

Clearly $S(\Lambda, W, t) \leq -\dot{G}(Y, t)$. Now we can choose $P(t) := 2 \begin{bmatrix} \frac{t}{2}I_p & I_p \end{bmatrix}^\top \begin{bmatrix} \frac{t}{2}I_p & I_p \end{bmatrix}$. Substituting P and X into (39), the left side of (39) becomes identically zero. Therefore, $\dot{V}(\Lambda, t) \leq S(\Lambda, W, t) \leq -\dot{G}(Y, t)$ with the storage function $V(\Lambda, t) := \Lambda^\top P \Lambda$. By defining the Lyapunov function $\mathcal{V}(Y, t) := V(Y, t) + G(Y, t)$, we immediately obtain $\dot{\mathcal{V}} \leq 0$ and also recover the same Lyapunov function used in Su et al. (2016).

Remark 13 As in the discrete-time case, the infinite family of LMIs (39) can also be reduced to a single LMI for the linear rate analysis of continuous-time ODEs. For further discussion on the topic of continuous-time exponential dissipation inequalities, see Hu & Seiler (2016).

6. Conclusion and Future Work

In this paper, we developed new notions of dissipativity theory for understanding of Nesterov's accelerated method. Our approach enjoys advantages of both the IQC framework (Lessard et al., 2016) and the discretization approach (Wilson et al., 2016) in the sense that our proposed LMI condition is simple enough for analytical rate analysis of Nesterov's method and can also be easily generalized to more complicated algorithms. Our approach also gives an intuitive interpretation of the convergence behavior of Nesterov's method using an energy dissipation perspective.

One potential application of our dissipativity theory is for the design of accelerated methods that are robust to gradient noise. This is similar to the algorithm design work in Lessard et al. (2016, §6). However, compared with the IQC approach in Lessard et al. (2016), our dissipativity theory leads to smaller LMIs. This can be beneficial since smaller LMIs are generally easier to solve analytically. In addition, the IQC approach in Lessard et al. (2016) is only applicable to strongly-convex objective functions while our dissipativity theory may facilitate the design of robust algorithm for weakly-convex objective functions. The dissipativity framework may also lead to the design of adaptive or time-varying algorithms.

Acknowledgements

Both authors would like to thank the anonymous reviewers for helpful suggestions that improved the clarity and quality of the final manuscript.

This material is based upon work supported by the National Science Foundation under Grant No. 1656951. Both authors also acknowledge support from the Wisconsin Institute for Discovery, the College of Engineering, and the Department of Electrical and Computer Engineering at the University of Wisconsin–Madison.

References

- Bubeck, S. Convex optimization: Algorithms and complexity. *Foundations and Trends® in Machine Learning*, 8(3-4):231–357, 2015.
- Bubeck, S., Lee, Y., and Singh, M. A geometric alternative to Nesterov's accelerated gradient descent. *arXiv preprint arXiv:1506.08187*, 2015.
- Drusvyatskiy, D., Fazel, M., and Roy, S. An optimal first order method based on optimal quadratic averaging. *arXiv preprint arXiv:1604.06543*, 2016.
- Fazlyab, M., Ribeiro, A., Morari, M., and Preciado, V. M. Analysis of optimization algorithms via integral quadratic constraints: Nonstrongly convex problems. *arXiv preprint arXiv:1705.03615*, 2017.
- Flammarion, N. and Bach, F. From averaging to acceleration, there is only a step-size. In *COLT*, pp. 658–695, 2015.
- Hu, B. and Seiler, P. Exponential decay rate conditions for uncertain linear systems using integral quadratic constraints. *IEEE Transactions on Automatic Control*, 61(11):3561–3567, 2016.
- Lessard, L., Recht, B., and Packard, A. Analysis and design of optimization algorithms via integral quadratic constraints. *SIAM Journal on Optimization*, 26(1):57–95, 2016.
- Megretski, A. and Rantzer, A. System analysis via integral quadratic constraints. *IEEE Transactions on Automatic Control*, 42:819–830, 1997.
- Megretski, A., Jönsson, U., Kao, C. Y., and Rantzer, A. *Control Systems Handbook*, chapter 41: Integral Quadratic Constraints. CRC Press, 2010.
- Nesterov, Y. *Introductory Lectures on Convex Optimization: A Basic Course*. Kluwer Academic Publishers, 2003.
- Polyak, B. T. Introduction to optimization. *Optimization Software*, 1987.
- Seiler, P. Stability analysis with dissipation inequalities and integral quadratic constraints. *IEEE Transactions on Automatic Control*, 60(6):1704–1709, 2015.
- Su, W., Boyd, S., and Candès, E. A differential equation for modeling Nesterov's accelerated gradient method: Theory and insights. *Journal of Machine Learning Research*, 17(153):1–43, 2016.
- Wibisono, A., Wilson, A., and Jordan, M. A variational perspective on accelerated methods in optimization. *Proceedings of the National Academy of Sciences*, pp. 201614734, 2016.
- Willems, J.C. Dissipative dynamical systems Part I: General theory. *Archive for Rational Mech. and Analysis*, 45(5):321–351, 1972a.
- Willems, J.C. Dissipative dynamical systems Part II: Linear systems with quadratic supply rates. *Archive for Rational Mech. and Analysis*, 45(5):352–393, 1972b.
- Wilson, A., Recht, B., and Jordan, M. A Lyapunov analysis of momentum methods in optimization. *arXiv preprint arXiv:1611.02635*, 2016.

Supplementary Material

We will make use of the following result throughout this section.

Proposition S1 *Suppose f is L -smooth and m -strongly convex. Then for all x, y the following inequalities hold.*

$$f(x) - f(y) \geq \nabla f(y)^\top (x - y) + \frac{m}{2} \|x - y\|^2 \quad (\text{S1})$$

$$f(y) - f(x) \geq \nabla f(y)^\top (y - x) - \frac{L}{2} \|y - x\|^2 \quad (\text{S2})$$

Proof. These inequalities follow from the definitions of L -smoothness and m -strong convexity. ■

A. Proof of Lemma 3

Applying (S1) with $(x, y) \mapsto (x_k, y_k)$, we obtain

$$f(x_k) - f(y_k) \geq \nabla f(y_k)^\top (x_k - y_k) + \frac{m}{2} \|x_k - y_k\|^2.$$

Applying (S2) with $(x, y) \mapsto (y_k - \alpha \nabla f(y_k), y_k)$, we obtain

$$f(y_k) - f(y_k - \alpha \nabla f(y_k)) \geq \frac{\alpha}{2} (2 - L\alpha) \|\nabla f(y_k)\|^2.$$

Summing these inequalities, we obtain:

$$f(x_k) - f(y_k - \alpha \nabla f(y_k)) \geq \nabla f(y_k)^\top (x_k - y_k) + \frac{m}{2} \|x_k - y_k\|^2 + \frac{\alpha}{2} (2 - L\alpha) \|\nabla f(y_k)\|^2. \quad (\text{S3})$$

Substituting $x_{k+1} = y_k - \alpha \nabla f(y_k)$ in the left-hand side of (S3), we can rewrite it as

$$\frac{1}{2} \begin{bmatrix} x_k - y_k \\ \nabla f(y_k) \end{bmatrix}^\top \left(\begin{bmatrix} m & 1 \\ 1 & \alpha(2 - L\alpha) \end{bmatrix} \otimes I_p \right) \begin{bmatrix} x_k - y_k \\ \nabla f(y_k) \end{bmatrix} \leq f(x_k) - f(x_{k+1}). \quad (\text{S4})$$

Substituting $y_k = (1 + \beta)x_k - \beta x_{k-1}$ into (S4), we obtain

$$\frac{1}{2} \begin{bmatrix} x_k - x_\star \\ x_{k-1} - x_\star \\ \nabla f(y_k) \end{bmatrix}^\top \left(\begin{bmatrix} \beta^2 m & -\beta^2 m & -\beta \\ -\beta^2 m & \beta^2 m & \beta \\ -\beta & \beta & \alpha(2 - L\alpha) \end{bmatrix} \otimes I_p \right) \begin{bmatrix} x_k - x_\star \\ x_{k-1} - x_\star \\ \nabla f(y_k) \end{bmatrix} \leq f(x_k) - f(x_{k+1}),$$

which directly leads to the formulation of $\tilde{X}_1$ in Lemma 3. Similarly, we apply (S1) with $(x, y) \mapsto (x_\star, y_k)$ and obtain

$$\frac{1}{2} \begin{bmatrix} x_k - x_\star \\ x_{k-1} - x_\star \\ \nabla f(y_k) \end{bmatrix}^\top \left(\begin{bmatrix} (1 + \beta)^2 m & -\beta(1 + \beta)m & -(1 + \beta) \\ -\beta(1 + \beta)m & \beta^2 m & \beta \\ -(1 + \beta) & \beta & \alpha(2 - L\alpha) \end{bmatrix} \otimes I_p \right) \begin{bmatrix} x_k - x_\star \\ x_{k-1} - x_\star \\ \nabla f(y_k) \end{bmatrix} \leq f(x_\star) - f(x_{k+1})$$

which directly leads to the formulation of $\tilde{X}_2$ in Lemma 3. The rest of the proof is straightforward. Actually, we can choose $\tilde{X} := \rho^2 \tilde{X}_1 + (1 - \rho^2) \tilde{X}_2$ and we directly obtain

$$\begin{bmatrix} x_k - x_\star \\ x_{k-1} - x_\star \\ \nabla f(y_k) \end{bmatrix}^\top \left(\tilde{X} \otimes I_p \right) \begin{bmatrix} x_k - x_\star \\ x_{k-1} - x_\star \\ \nabla f(y_k) \end{bmatrix} \leq -(f(x_{k+1}) - f(x_\star)) + \rho^2 (f(x_k) - f(x_\star)).$$

Specifically, $\tilde{X}$ may be computed as

$$\tilde{X} = \frac{1}{2} \begin{bmatrix} (1 + \beta)^2 m - (1 + 2\beta)m\rho^2 & (\rho^2 - 1 - \beta)\beta m & \rho^2 - 1 - \beta \\ (\rho^2 - 1 - \beta)\beta m & \beta^2 m & \beta \\ \rho^2 - 1 - \beta & \beta & \alpha(2 - L\alpha) \end{bmatrix}.$$

■

B. Proof of Lemma 5

Applying (S2) with $(x, y) \mapsto (x_{k+1}, y_k)$, and making the substitutions $x_{k+1} = (1 + \beta)x_k - \beta x_{k-1} - \alpha \nabla f(y_k)$ and $y_k = (1 + \eta)x_k - \eta x_{k-1}$, we obtain:

$$\begin{aligned} f(y_k) - f(x_{k+1}) &\geq \nabla f(y_k)^\top (y_k - x_{k+1}) - \frac{L}{2} \|x_{k+1} - y_k\|^2 \\ &= \nabla f(y_k)^\top ((\beta - \eta)(x_{k-1} - x_k) + \alpha \nabla f(y_k)) - \frac{L}{2} \|(\beta - \eta)(x_{k-1} - x_k) + \alpha \nabla f(y_k)\|^2 \\ &= \frac{1}{2} \begin{bmatrix} x_k - x_\star \\ x_{k-1} - x_\star \\ \nabla f(y_k) \end{bmatrix}^\top \left(\begin{bmatrix} -L(\beta - \eta)^2 & L(\beta - \eta)^2 & -(1 - L\alpha)(\beta - \eta) \\ L(\beta - \eta)^2 & -L(\beta - \eta)^2 & (1 - L\alpha)(\beta - \eta) \\ -(1 - L\alpha)(\beta - \eta) & (1 - L\alpha)(\beta - \eta) & \alpha(2 - L\alpha) \end{bmatrix} \otimes I_p \right) \begin{bmatrix} x_k - x_\star \\ x_{k-1} - x_\star \\ \nabla f(y_k) \end{bmatrix} \end{aligned} \quad (\text{S5})$$

Applying (S1) with $(x, y) \mapsto (x_k, y_k)$ and substituting $y_k = (1 + \eta)x_k - \eta x_{k-1}$, we obtain:

$$\begin{aligned} f(x_k) - f(y_k) &\geq \nabla f(y_k)^\top (x_k - y_k) + \frac{m}{2} \|x_k - y_k\|^2 \\ &= \eta \nabla f(y_k)^\top (x_{k-1} - x_k) + \frac{m\eta^2}{2} \|x_{k-1} - x_k\|^2 \\ &= \frac{1}{2} \begin{bmatrix} x_k - x_\star \\ x_{k-1} - x_\star \\ \nabla f(y_k) \end{bmatrix}^\top \left(\begin{bmatrix} \eta^2 m & -\eta^2 m & -\eta \\ -\eta^2 m & \eta^2 m & \eta \\ -\eta & \eta & 0 \end{bmatrix} \otimes I_p \right) \begin{bmatrix} x_k - x_\star \\ x_{k-1} - x_\star \\ \nabla f(y_k) \end{bmatrix} \end{aligned} \quad (\text{S6})$$

Applying (S1) with $(x, y) \mapsto (x_\star, y_k)$ and again substituting $y_k = (1 + \eta)x_k - \eta x_{k-1}$, we obtain:

$$\begin{aligned} f(x_\star) - f(y_k) &\geq \nabla f(y_k)^\top (x_\star - y_k) + \frac{m}{2} \|x_\star - y_k\|^2 \\ &= -\nabla f(y_k)^\top ((1 + \eta)(x_k - x_\star) - \eta(x_{k-1} - x_\star)) + \frac{m}{2} \|(1 + \eta)(x_k - x_\star) - \eta(x_{k-1} - x_\star)\|^2 \\ &= \frac{1}{2} \begin{bmatrix} x_k - x_\star \\ x_{k-1} - x_\star \\ \nabla f(y_k) \end{bmatrix}^\top \left(\begin{bmatrix} (1 + \eta)^2 m & -\eta(1 + \eta)m & -(1 + \eta) \\ -\eta(1 + \eta)m & \eta^2 m & \eta \\ -(1 + \eta) & \eta & 0 \end{bmatrix} \otimes I_p \right) \begin{bmatrix} x_k - x_\star \\ x_{k-1} - x_\star \\ \nabla f(y_k) \end{bmatrix} \end{aligned} \quad (\text{S7})$$

By adding (S5)–(S7) with the definitions of $\tilde{X}_1$, $\tilde{X}_2$, and $\tilde{X}_3$ in Lemma 5, we obtain:

$$\begin{aligned} \begin{bmatrix} x_k - x_\star \\ x_{k-1} - x_\star \\ \nabla f(y_k) \end{bmatrix}^\top \left((\tilde{X}_1 + \tilde{X}_2) \otimes I_p \right) \begin{bmatrix} x_k - x_\star \\ x_{k-1} - x_\star \\ \nabla f(y_k) \end{bmatrix} &\leq f(x_k) - f(x_{k+1}) \\ \begin{bmatrix} x_k - x_\star \\ x_{k-1} - x_\star \\ \nabla f(y_k) \end{bmatrix}^\top \left((\tilde{X}_1 + \tilde{X}_3) \otimes I_p \right) \begin{bmatrix} x_k - x_\star \\ x_{k-1} - x_\star \\ \nabla f(y_k) \end{bmatrix} &\leq f(x_\star) - f(x_{k+1}) \end{aligned}$$

The rest of the proof follows by substituting above expressions into the weighted sum with ρ^2 . ■

C. Proof of Lemma 8

Since f is L -smooth and convex, we can use the same proof technique as in Lemma 3 while setting $m = 0$ and $\alpha = \frac{1}{L}$. We can thus obtain the following inequalities that parallel (S4).

$$\begin{aligned} \frac{1}{2} \begin{bmatrix} y_k - x_k \\ \nabla f(y_k) \end{bmatrix}^\top \left(\begin{bmatrix} 0 & 1 \\ 1 & -\frac{1}{L} \end{bmatrix} \otimes I_p \right) \begin{bmatrix} y_k - x_k \\ \nabla f(y_k) \end{bmatrix} &\geq f(x_{k+1}) - f(x_k) \\ \frac{1}{2} \begin{bmatrix} y_k - x_\star \\ \nabla f(y_k) \end{bmatrix}^\top \left(\begin{bmatrix} 0 & 1 \\ 1 & -\frac{1}{L} \end{bmatrix} \otimes I_p \right) \begin{bmatrix} y_k - x_\star \\ \nabla f(y_k) \end{bmatrix} &\geq f(x_{k+1}) - f(x_\star) \end{aligned}$$

The conclusion of Lemma 8 follows once we substitute $y_k = (1 - \beta_k)x_k + \beta_k x_{k-1}$. ■