

Remarks on the Polyak-Łojasiewicz inequality and the convergence of gradient systems

Arthur C. B. de Oliveira¹, Leilei Cui², and Eduardo D. Sontag^{1,3}

Abstract—This work explores generalizations of the Polyak-Łojasiewicz inequality (PLI) and their implications for the convergence behavior of gradient flows in optimization problems. Motivated by the continuous-time linear quadratic regulator (CT-LQR) policy optimization problem – where only a weaker version of the PLI is characterized in the literature – this work shows that while weaker conditions are sufficient for global convergence to, and optimality of the set of critical points of the cost function, the “profile” of the gradient flow solution can change significantly depending on which “flavor” of inequality the cost satisfies. After a general theoretical analysis, we focus on fitting the CT-LQR policy optimization problem to the proposed framework, showing that, in fact, it can never satisfy a PLI in its strongest form. We follow up our analysis with a brief discussion on the difference between continuous- and discrete-time LQR policy optimization, and end the paper with some intuition on the extension of this framework to optimization problems with L_1 regularization and solved through proximal gradient flows.

I. INTRODUCTION

Recent advances in Artificial Intelligence (AI) and Machine Learning (ML) have rekindled interest in optimization theory, with many traditional results being revisited in light of the proposed techniques [1]–[8]. In particular, the typical model-free formulation of many successful learning techniques motivates the study of gradient-based optimization methods, which are invaluable in understanding the training of neural networks and similar architectures, typically done through back-propagation algorithms.

Gradient descent or, in continuous time, gradient flow, consists in searching for the argument x that minimizes the value of a given function $f(x)$ by “moving along” the direction of steepest descent of the cost function. Theoretical guarantees are typically desirable, and in search of balancing generality and good properties, often in the optimization literature one deals with specific classes of optimization problems, such as convex optimization [9] or linear programming [10]. In this paper, we will focus on optimization problems that satisfy (to different degrees) a Polyak-Łojasiewicz inequality (PLI), also known as the gradient dominance condition [11], [12].

The PLI is a staple in nonlinear optimization analysis, as, in its strongest form, it guarantees global exponential convergence of the gradient flow to the optimal solution

of the problem [12]. Furthermore, satisfying a PLI globally (gPLI) also guarantees strong robustness properties [2]. However, characterizing such a condition might not be possible for every optimization problem. In [1] the authors noticed that more general conditions than the gPLI can be proposed by using different classes of comparison functions so that different robustness results can be guaranteed.

In particular, the problem of policy optimization for the linear quadratic regulator (LQR) motivates the discussion around weaker versions of the PLI [1], [2], [6], [13]–[18]. For the discrete-time version of the problem, in [14]–[16] the authors show that it satisfies a gPLI, guaranteeing exponential convergence to the optimal feedback law for initialization in the stabilizing set of feedback matrices. However, so far in the literature for the continuous-time LQR policy optimization problem there is no characterization of a gPLI [6], [17], with the analysis in [6] indicating that, at least for the scalar case, the continuous-time LQR does not satisfy a gPLI.

In this work, we are interested in characterizing how generalizations of the PLI affect the rate of convergence of the solution. We begin in Section II by revisiting common assumptions and their consequences regarding the convergence of the gradient flow. We then formally introduce the gPLI and a few other weaker definitions, and discuss their differences and consequences to the convergence of the gradient flow solution. We next deepen the analysis by defining a new family of conditions closely related to, but more general than, the global PLI. We discuss how these weaker conditions relate to each other and how they can characterize weaker forms of convergence than the gPLI. Then, in Section III, we contextualize the theoretical analysis of this paper through the specific problem of the continuous-time LQR policy optimization problem. This problem has no guarantees of satisfying a gPLI, and in fact we show that it can never satisfy such a condition. We characterize which sequences of points of the policy space result in an unbounded value for the gradient of the cost, and which result in a “sub-exponential” convergence profile for the solution. We follow up with a brief discussion on the difference between the continuous- and discrete-time LQR policy optimization, and finalize the paper in Section IV with a comment on possible links between the analysis of this paper and proximal gradient flow for optimization problems with L_1 regularizing terms.

All proofs are omitted due to space constraints but are available in the extended version of this paper on arXiv [19].

The work of EDS and ACO was partially supported by grants AFOSR FA9550-22-1-0316, 21-1-0289, and ONR N00014-21-1-2431

¹Department of Electrical and Computer Engineering, Northeastern University, USA a.castello@northeastern.edu

²Massachusetts Institute of Technology, Massachusetts, USA llcui@mit.edu

³Department of BioEngineering, Northeastern University, USA e.sontag@northeastern.edu

II. THEORETICAL SETUP

Along this paper, let $\mathbb{R}$, $\mathbb{R}_+$ and $\mathbb{R}_{++}$ denote the real, non-negative real, and strictly positive real numbers, respectively. Let $\mathbb{S}_+^n$ and $\mathbb{S}_{++}^n$ be the set of positive semi-definite and positive definite n -by- n matrices.

A given function $\alpha : \mathbb{R}_+ \rightarrow \mathbb{R}$ is said to be *positive-definite* ($\mathcal{PD}$) if $\alpha(0) = 0$ and $\alpha(x) > 0$ for all $x \neq 0$. Similarly, α is said to be of class $\mathcal{K}$ if it is continuous, positive-definite, and strictly increasing. Finally, α is of class $\mathcal{K}_\infty$ if it is of class $\mathcal{K}$ and unbounded.

For a given function $f : \mathcal{X} \rightarrow \mathbb{R}$ bounded below, let $\underline{f} = \inf_{x \in \mathcal{X}} f(x)$, and $x^* = \arg \inf_{x \in \mathcal{X}} f(x)$.

A. Optimization problems and gradient methods

Let $\mathcal{X}$ be an open subset of an Euclidean space (the analysis in this paper can be generalized to manifolds, but we refrain from it for simplicity). Then, an optimization problem consists of searching for the value of an argument/parameter $x \in \mathcal{X}$ that minimizes some cost function $f : \mathcal{X} \rightarrow \mathbb{R}$ (or minimizes the negative of a reward for maximization). Mathematically, we write such a problem as

$$\begin{aligned} & \underset{x}{\text{minimize}} && f(x) \\ & \text{subject to} && x \in \mathcal{X} \end{aligned} \quad (1)$$

which might have one, multiple, or no solution, requiring some assumptions about either the cost f or the search space $\mathcal{X}$ to guarantee the existence and uniqueness of the solution. A common assumption is that of compactness of $\mathcal{X}$, and continuity of $f(x)$ for $x \in \mathcal{X}$, which would guarantee the existence of a minimum (and maximum) in $\mathcal{X}$ since $f(\mathcal{X}) := \{f(x) \mid x \in \mathcal{X}\}$ would be compact. Usually, however, the search space $\mathcal{X}$ is not compact, requiring the adoption of an alternative set of assumptions, outlined next.

Assumption 1: The function f is *real analytic*, *bounded below*, and *proper* (i.e. coercive).

It is easy to prove that Assumption 1 guarantees the existence of a minimum $\underline{f} \in f(\mathcal{X})$ attained at a set of points $\mathcal{T} := \{x \in \mathcal{X} \mid f(x) = \underline{f}\}$. Notice that from the point of view of the optimization problem, any $x \in \mathcal{T}$ is a valid solution of (1), as all result in the same value for the cost function. Then, the optimization problem can be thought of as finding any $x \in \mathcal{T}$.

A natural candidate for solutions to the optimization problem is the set of critical points of f , i.e. the set of points $\mathcal{Z} := \{x \in \mathcal{X} \mid \nabla f(x) = 0\}$. A common strategy for finding an $x \in \mathcal{Z}$ is “moving the parameters along the direction of steepest descent of the function”. Mathematically and in continuous-time, this means imposing the following dynamics for the parameter x

$$\dot{x} = -\nabla f(x), \quad (2)$$

while in discrete time one would impose the following update law for x_k for a small enough $h > 0$

$$x_{k+1} = x_k - h \nabla f(x_k). \quad (3)$$

In this paper we focus on the continuous-time strategy, and one can easily verify that x is an equilibrium of (2) if and only if $x \in \mathcal{Z}$. Nonetheless, there is no a priori guarantee that a solution of (2) initialized in $\mathcal{X}$ will converge to a point in $\mathcal{Z}$, much less in $\mathcal{T}$. So, we next look at what convergence guarantees Assumption 1 allows us to derive, and what other assumptions can be made to improve such guarantees.

B. Convergence guarantees and the Polyak-Łojasiewicz inequality

Consider an optimization problem (1) satisfying Assumption 1, then the following results hold:

Lemma 1: For any *proper* function $f : \mathcal{X} \rightarrow \mathbb{R}$, the solution of the gradient flow (2) initialized at any point $x \in \mathcal{X}$ is *precompact*.

Theorem 1 (Łojasiewicz’s theorem [20]): Let $f : \mathcal{X} \rightarrow \mathbb{R}$ be a real analytic function and let $\phi : \mathbb{R}_+ \times \mathcal{X} \rightarrow \mathcal{X}$, the solution of the initial value problem (2), be precompact, i.e. $\sup_{t>0, x_0 \in \mathcal{X}} \|\phi(t, x_0)\| < \infty$. Then, for all $x_0 \in \mathcal{X}$, $\lim_{t \rightarrow \infty} \phi(t, x_0) \in \mathcal{Z}$, that is, all solutions of the gradient flow (2) initialized in $\mathcal{X}$ converge to a critical point of the cost function f .

Lemma 1 and Theorem 1 guarantee that any solution of (2) initialized in $\mathcal{X}$ will converge to a critical point of the function f . However, $x \in \mathcal{Z}$ is only a necessary condition for the optimality of x . In fact a point $x \in \mathcal{Z}$ can be either a local minimum, a local maximum, or a saddle-point of f . Regarding that, the following result can be stated ([3], [4] appendix A):

Lemma 2 ([3], [4]): Let

$$\mathcal{S} := \{x \in \mathcal{Z} \mid \exists v \in \mathbb{R}^n \text{ s.t. } v^\top \nabla^2 f(x) v < 0\},$$

where $\nabla^2 f$ is the Hessian of f , and let $\phi : \mathbb{R}_+ \times \mathcal{X} \rightarrow \mathcal{X}$ be the solution of the initial value problem (2). Then, the set of $x_0 \in \mathcal{X}$ for which $\lim_{t \rightarrow \infty} \phi(t, x_0) \in \mathcal{S}$ has Lebesgue measure zero. In other words, the center-stable manifold of $\mathcal{S}$ has measure zero.

Notice that Lemma 2 holds even if $\mathcal{S}$ is a continuous set, or the union of continuous sets. In fact, $\mathcal{S}$ need not be even compact, as long as the condition on the value of the Hessian holds for all of its elements. However, despite this result excluding any local maxima and “strict” saddles from the result of a gradient flow solution to problem (1) (with probability one), it is still not enough to guarantee the optimality of the gradient flow solution. Typically, other assumptions are added in the literature to ensure that $\lim_{t \rightarrow \infty} \phi(t, x_0) \in \mathcal{T}$, with, arguably, one of the most popular being the convexity of f . If f is convex in $\mathcal{X}$, then $\mathcal{Z} = \mathcal{T}$ and a gradient flow will eventually find the optimal solution. Furthermore, if f is strongly convex, then a solution of (2) converges to $\mathcal{T}$ exponentially.

In this paper, we first review a condition for f that is weaker than strong convexity but still ensures that a solution of (2) converges to $\mathcal{T}$ exponentially.

Definition 1 (μ -global Polyak-Łojasiewicz inequality):

Given a fixed $\mu > 0$, a function f satisfies a μ -global

Polyak-Łojasiewicz inequality (μ -gPŁI) if

$$\|\nabla f(x)\| \geq \alpha(f(x) - \underline{f}). \quad (4)$$

with $\alpha(r) = \sqrt{\mu r}$, for all $x \in \mathcal{X}$. Furthermore, f is gPŁI if it is μ -gPŁI for some $\mu > 0$.

The property in Definition 1 (often written in the form $\|\nabla f(x)\|^2 \geq \mu(f(x) - \underline{f})$) has been the object of much recent study, and is a natural generalization of convexity (see [12] for a thorough analysis of the relationship between convexity and the gPŁI). An immediate consequence of this property is that all critical points of f must solve (1), *i.e.* if f is gPŁI, then $\mathcal{Z} = \mathcal{T}$, which guarantees the optimality of gradient flow solutions. Further usefulness of this property lies in the fact that it guarantees exponential convergence of the solution of a gradient flow, as we show next.

Definition 2 (μ -global exponential stability): Given a fixed $\mu > 0$, the gradient flow (2) of f is μ -globally exponential stable (μ -GES) if

$$f(\phi(t, x_0)) - \underline{f} \leq (f(x_0) - \underline{f})e^{-\mu t} \text{ for all } t \geq 0. \quad (5)$$

Furthermore, the gradient flow of f is GES if it is μ -GES for some $\mu > 0$.

Lemma 3: The gradient flow (2) of f is μ -GES if and only if f satisfies a μ -global PŁI.

Despite the good convergence properties associated with having an exponential upper-bound, proving that a cost function is μ -global PŁI is not always possible, and in some relevant examples in the literature, the following weaker condition is characterized instead.

Definition 3 (*Semi-global PŁI*): A function f satisfies a *semi-global Polyak-Łojasiewicz inequality* (sgPŁI) if, for every $\epsilon > 0$, there exists a $\mu_\epsilon > 0$ such that

$$\|\nabla f(x)\| \geq \alpha_\epsilon(f(x) - \underline{f}), \quad (6)$$

with $\alpha_\epsilon(r) = \sqrt{\mu_\epsilon r}$, for all $x \in \mathcal{X}_\epsilon := \{x \in \mathcal{X} \mid f(x) - \underline{f} \leq \epsilon\}$.

Similarly to satisfying a global PŁI, satisfying a semi-global PŁI guarantees that all critical points of f must be global minima of (1), *i.e.* $\mathcal{Z} = \mathcal{T}$. In fact, at first glance global and semi-global PŁIs look very similar, with the latter also guaranteeing for any $\epsilon > 0$, an exponential rate of convergence μ_ϵ for all initializations in the sublevel set $\mathcal{X}_\epsilon$. However, their distinction becomes important when analyzing the rate of convergence of gradient flow solutions, as the following lemma illustrates.

Lemma 4: If a function f satisfies a gPŁI, then it also satisfies a sgPŁI. Alternatively, if a function f satisfies a semi-global PŁI but not a global PŁI for any $\mu > 0$, then there must exist a sequence $\{\epsilon_i\}$, $i = 1, 2, \dots$ with $\epsilon_i > 0$ such that

$$\lim_{i \rightarrow \infty} \bar{\mu}_{\epsilon_i} = 0, \quad (7)$$

where $\bar{\mu}_\epsilon$ is the largest μ_ϵ that satisfies (6) for a given $\epsilon > 0$.

Lemma 4 makes the distinction between satisfying a global and a semi-global PŁI clear: if the inequality is only semi-global, then there must be some “unbounded sequence” of points $\{x_i\}$, $x_i \in \mathcal{X}$, with $\lim_{i \rightarrow \infty} f(x_i) - \underline{f} = \infty$, for

which the exponential rate of convergence $\bar{\mu}_{f(x_i) - \underline{f}}$ goes to zero as i goes to infinity. Although intuitively this tells us that the convergence is not purely exponential globally, it is still unclear what the solution of the gradient-flow looks like.

Furthermore, an even weaker version of the PŁI can be characterized as follows:

Definition 4 (ϵ -local PŁI): Given a fixed $\epsilon > 0$, a function f satisfies a ϵ -local Polyak-Łojasiewicz inequality (ϵ -lPŁI) if there exists some $\mu > 0$ such that

$$\|\nabla f(x)\| \geq \alpha(f(x) - \underline{f}), \quad (8)$$

with $\alpha(r) = \sqrt{\mu r}$, for every $x \in \mathcal{X}_\epsilon := \{x \in \mathcal{X} \mid f(x) - \underline{f} \leq \epsilon\}$. Furthermore, f is lPŁI if it is ϵ -lPŁI for some $\epsilon > 0$.

Differently from global and semi-global PŁIs, a local PŁI only gives guarantees for fixed a neighborhood $\mathcal{X}_\epsilon$ around the optimal value. Although inside such neighborhood the same guarantees are obtainable (exponential convergence, optimality of critical points, etc.), no guarantee exists outside of it in general.

Definitions 1, 3 and 4 provide good granularity when analyzing the behavior of the gradient flow for different classes of cost functions, however further precision can be attained with the help of comparison functions, as we discuss next.

C. A generalization of the PŁI

The classic PŁI condition is originally formulated as in Definition 1, using the square root comparison function. However, practical cost functions often fail to satisfy this condition globally — an example being the continuous-time LQR cost, which will be discussed later. This limitation motivated us in [1] to propose a nonlinear version of the PŁI. By simply generalizing $\alpha(r)$ in (4) from $\alpha = \sqrt{\mu r}$ to a positive definite function $\alpha \in \mathcal{PD}$, the convergence of the gradient flow can still be ensured. As an additional benefit, when the gradient flow in (2) is subject to additive noise, the error $f(\phi(t, x_0)) - \underline{f}$ is bounded by an energy-like measure of the noise. This property is formally known as integral input-to-state stability (iISS) [21]. Furthermore, if α is strengthened to a class $\mathcal{K}$ function, the error $f(\phi(t, x_0)) - \underline{f}$ not only converges to zero in the absence of noise but also remains bounded when the noise is below a certain threshold, a property referred to as small-input ISS (siISS) [1]. If α is further strengthened to be a class $\mathcal{K}_\infty$ function, then the error $f(\phi(t, x_0)) - \underline{f}$ converges to zero in the noise-free case and remains bounded under any bounded noise, which corresponds to the classical input-to-state stability (ISS) [22]. Clearly, the global PŁI belongs to the class of $\mathcal{K}_\infty$ functions, while the semi-global PŁI belongs to the class of $\mathcal{PD}$ functions.

For our analysis, we also introduce a new class of functions, called class *satPŁI* functions, which can be represented as

$$\alpha(r) = \sqrt{\frac{ar}{b+r}} \quad \forall r \geq 0, \quad (9)$$

where $a, b > 0$ are constants.

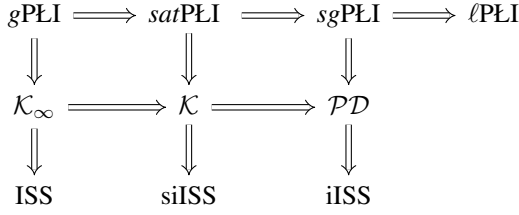


Fig. 1. Diagram of the hierarchy between types of comparison functions and their relationship with the different types of Polyak-Łojasiewicz inequality. Notice in particular that while all functions that satisfy a PLI also have a class $\mathcal{K}_\infty$ lower-bound (as presented in definition 5), the converse is not necessarily true. Similarly, satisfying a sgPLI implies there exists a $\mathcal{PD}$ lower-bound, but the converse is not true. Also notice that the newly introduced class $sat\text{PLI}$ lower bound lies in between $g\text{PLI}$ and $sg\text{PLI}$, and provides a better convergence guarantee. Finally, the ℓPLI stands isolated in the graph, but one should note that it sustains the same convergence and robustness properties than the $g\text{PLI}$, so long as the disturbance is small enough to keep the solution in a neighborhood of the optimum.

With this established, the following definition summarizes the “zoo” of the generalized inequalities based on different classes of comparison functions.

Definition 5: A function f satisfies a *class $\mathcal{K}_\infty$ lower bound* (resp. *class $sat\text{PLI}$, $\mathcal{K}$, or $\mathcal{PD}$*) if

$$\|\nabla f(x)\| \geq \alpha(f(x) - f(x^*)), \quad (10)$$

for all $x \in \mathcal{X}$, with α being a function of class $\mathcal{K}_\infty$ (resp. *class $sat\text{PLI}$, $\mathcal{K}$, or $\mathcal{PD}$*).

Notice that there is a natural order between the comparison function and a natural relation between each of them and the previously defined different types of PLI, all illustrated in Fig. 1. In particular, notice that if the comparison function α is of class $sat\text{PLI}$, then it lies in between a $g\text{PLI}$ and a $sg\text{PLI}$. Further relationships between the different classes can be established, as the following Lemma illustrates.

Lemma 5: For a function f , the following two statements are equivalent

- 1) The function f satisfies a class $\mathcal{K}_{\text{SAT}}$ lower bound.
- 2) The function f satisfies: a local PLI; and a $\mathcal{PD}$ lower bound with comparison function α that is such that

$$\liminf_{r \rightarrow \infty} \alpha(r) > 0.$$

Satisfying different classes of lower bounds can also be shown to have consequences for the convergence rate of solutions. To better illustrate the effects of the different comparison functions, we next define a weaker form of convergence than the one from Definition 2.

Definition 6 (global linear-exponential stability): The gradient flow (2) of f is globally linear-exponential stable (GLES) if there exists a $m > 0$ such that for every $x_0 \in \mathcal{X}$ and every $\underline{t} > 0$ there exists a $\mu_{x_0, \underline{t}} > 0$ for which

$$\begin{aligned}
& (f(\phi(t, x_0)) - \underline{f}) \\
& \leq \begin{cases} (f(\phi(\underline{t}, x_0)) - \underline{f}) - m(t - \underline{t}) & \text{if } t \leq \underline{t} \\ (f(\phi(\underline{t}, x_0)) - \underline{f})e^{-\mu_{x_0, \underline{t}}(t - \underline{t})} & \text{if } t > \underline{t} \end{cases}
\end{aligned} \quad (11)$$

From this definition, we can derive rigorous conditions for the solution to be GLES as follows:

Lemma 6: The solution of the gradient flow (2) of a

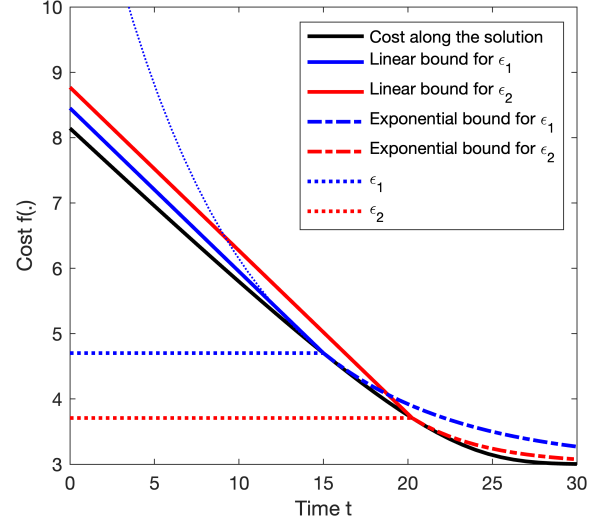


Fig. 2. Gradient Flow trajectory for a cost f that satisfies the conditions in Lemma 6. Notice that for different given values of ϵ , one can find a pair of line and exponential such that the line upper-bounds the solution while it is outside the levelset $\mathcal{X}_\epsilon$, and the exponential upper-bounds it when it is inside the level-set. Furthermore, notice that smaller values of ϵ result in a tighter bound for the exponential section, while also on a looser bound for the linear section, while the opposite occurs when ϵ is larger. Finally, the thin dotted blue line illustrates how the exponential bound for ϵ_1 quickly becomes conservative if $t < \underline{t}$, pointing to linear-exponential as a tighter bound than purely exponential.

function f is GLES if

- The gradient ∇f is globally bounded;
- The function f satisfies a class $\mathcal{K}_{\text{SAT}}$ lower-bound.

We illustrate the results from Lemma 6 in Fig. 2. Observe that the true trajectory of the cost follows a “linear-exponential” profile, which means that for a given “margin of error” $\epsilon := f(\phi(t, x_0)) - \underline{f}$, the solution is upper bounded by a line for values of the cost higher than ϵ and by an exponential for values smaller than ϵ . Also notice from the figure that a smaller margin of error ϵ results in a worse upper bound for the linear part of the solution, but a tighter upper bound for the exponential part.

Next, we informally point out that for a given x_0 and $\underline{t}$, the exponential $(f(\phi(\underline{t}, x_0)) - \underline{f})e^{-\mu_{x_0, \underline{t}}(t - \underline{t})}$ upper-bounds the solution $f(\phi(t, x_0)) - \underline{f}$ for all time $t > 0$, not only $t > \underline{t}$, however, for $t < \underline{t}$, the distance between the exponential upper-bound and the actual solution quickly grows. This can be noticed from Fig. 2 by looking at the thin dotted extension of the blue exponential for $t < \underline{t} = 15$. This means that the solution of this gradient flow is, in a sense, globally upper-bounded by an exponential. However, this bound is not tight for points much larger than $\phi(\underline{t}, x_0)$ (in the sense of $f(\phi(\underline{t}, x_0))$).

Moreover, notice that the conditions in Lemma 6 are only sufficient but not necessary. That is because if f has a globally bounded gradient and satisfies both a local PLI and a $\mathcal{PD}$ lower-bound, but is such that $\liminf_{r \rightarrow \infty} \alpha(r) = 0$ for all positive-definite α for which f satisfies (10), then there is still a linear-exponential upper-bound as described. However,

in that case, the convergence of the solution can also be upper-bounded by a “log-exponential” function constructed similarly to how the linear-exponential bound is built in (11), and as a consequence, any linear function bounding the solution for $t < \underline{t}$ will also not be tight for points much larger than $\phi(\underline{t}, x_0)$ (in the sense of $f(\phi(\underline{t}, x_0))$).

Finally, notice a gap between GES and GLES: if f does not satisfy a global PŁI, but also does not have a globally bounded gradient, then its solution is neither GES nor GLES. This is an important observation for the next section of the paper, as we will show that this is precisely the case for the LQR cost function.

In this section we provided tools for analyzing cost functions f that satisfy Assumption 1 and have one of the properties characterized in Definitions 1, 3, 4, and 5. In the next section, we use these tools to analyze a very important example from the literature: the policy optimization problem for the Linear Quadratic Regulator (LQR).

III. APPLICATIONS TO LQR POLICY OPTIMIZATION

Consider the following continuous-time linear system:

$$\dot{x} = Ax + Bu \quad (12)$$

where $A \in \mathbb{R}^{n \times n}$, and $B \in \mathbb{R}^{n \times m}$ are the system matrices and are such that (A, B) is controllable. Let $\mathcal{G} := \{K \in \mathbb{R}^{m \times n} \mid A - BK \text{ is Hurwitz}\}$. The objective is to determine an output feedback $u = -Ky$ with $K \in \mathcal{G}$ that minimizes

$$\bar{J}(K) = \mathbb{E}_{x_0 \sim \mathcal{X}_0} \left[\int_0^\infty x(t)^\top Q x(t) + x(t)^\top K^\top R K x(t) dt \right], \quad (13)$$

with given matrices $R \in \mathbb{S}_{++}^{m \times m}$ and $Q \in \mathbb{S}_{++}^{n \times n}$, and for x_0 sampled from a probability distribution $\mathcal{X}_0$. It is well known that for linear systems minimizing $\bar{J}(K)$ is equivalent to minimizing the following cost function

$$J(K) = \text{trace}(P_K), \quad (14)$$

where

$$P_K(A - BK) + (A - BK)^\top P_K + K^\top R K + Q = 0, \quad (15)$$

in the sense that a $K^* \in \mathcal{G}$ minimizes $J(\cdot)$ if and only if it minimizes $\bar{J}(\cdot)$. If matrices A and B are known, and the solution is initialized at a point x_0 sampled such that $\mathbb{E}(x_0 x_0^\top) = \Sigma_0 \in \mathbb{S}_{++}^n$, then one can find the optimal feedback matrix $K^* = R^{-1} B^\top P$ where P solves the following Riccati equation

$$A^\top P + PA + PBR^{-1}B^\top P + Q = 0. \quad (16)$$

However, a popular formulation arises when one consider the case where the system matrices are not available, i.e. a model-free or policy optimization approach [1], [3], [4], [6], [13], [14], [16], [17]. In that case, the optimal feedback matrix K^* is obtained by following the negative direction of the gradient $\dot{K} = -\nabla J(K)$, with the gradient ∇J being

given by [23]:

$$\nabla J(K) = -2(B^\top P_K - RK)Y_K, \quad (17)$$

where, for any $K \in \mathcal{G}$, P_K is the solution of (15), and Y_K is the unique positive definite solutions of the following Lyapunov equation

$$Y_K(A - BK)^\top + (A - BK)Y_K + I = 0. \quad (18)$$

In previous works in the literature [5], [6], it was established that the solution for a gradient flow dynamics for solving this problem initialized inside $\mathcal{G}_a := \{K \in \mathbb{R}^{m \times n} \mid J(K) \leq a\}$, satisfies a semi-global PŁI (Lemma 1 of [5], and Theorem 3.16 of [6]), and in [1] it was shown that it actually satisfies a class *sat*PŁI lower-bound. Although, a priori, this does not mean that J defined in (14) does not also satisfy a gPŁI, in the following section we will prove that this is actually the case, i.e. $J(\cdot)$ can never satisfy a gPŁI.

A. The LQR cost lies in the gap

As mentioned, in this section we will show that the LQR cost $J(\cdot)$ has an unbounded gradient (thus not satisfying the conditions for Lemma 6) while also provably not satisfying a gPŁI (thus not admitting a global exponential rate of convergence).

We begin by formally defining “high gain curves” in the space of stable feedback matrices $\mathcal{G}$ and showing that, along any such high gain curve, the gradient is bounded, and thus the LQR cost can never satisfy a global PŁI. Then we show that for any sequence of matrices “approaching the border of instability” (in a sense to be formally defined) the gradient goes to infinity, proving that one cannot upper-bound the gradient globally in $\mathcal{G}$.

Definition 7: A matrix-valued function $\tilde{K} : \mathbb{R}_+ \rightarrow \mathcal{G}$, is called a *high gain curve* of $\mathcal{G}$ if there exists an $\epsilon > 0$ such that the eigenvalues of the closed-loop matrix are strictly in the left half-plane (LHP), i.e. for all $\rho > 0$,

$$\max_i \Re(\lambda_i(A - B(\tilde{K}(\rho) + K^*))) < -\epsilon,$$

and the limit value of the cost is unbounded, i.e. $\lim_{\rho \rightarrow \infty} J(K^* + \tilde{K}(\rho)) = \infty$.

Furthermore, we can guarantee that any (A, B) controllable has a high gain curve, as we show next:

Lemma 7: Let $A \in \mathbb{R}^{n \times n}$ and $B \in \mathbb{R}^{n \times m}$ be the system matrices for (12), with (A, B) being controllable. Then, there always exists a $\tilde{K} : \mathbb{R}_+ \rightarrow \mathcal{G}$ that is a high gain curve of $\mathcal{G} := \{K \in \mathbb{R}^{m \times n} \mid (A - BK) \text{ is Hurwitz}\}$, with the additional property that $B\tilde{K}(\rho)$ is diagonalizable for all $\rho > 0$.

With this, we can state the following lemma regarding the behavior of the gradient along high gain trajectories:

Lemma 8: Let $\tilde{K} : \mathbb{R}_+ \rightarrow \mathcal{G}$ be a high gain curve of $\mathcal{G}$ and let there be a $\bar{\rho} > 0$ such that $B\tilde{K}(\rho)$ is diagonalizable for all $\rho > \bar{\rho}$. Then the limit

$$\lim_{\rho \rightarrow \infty} \nabla J(\tilde{K}(\rho) + K^*) \in \mathbb{R}^{m \times n}, \quad (19)$$

exists and is finite, i.e. the norm of the gradient $\|\nabla J(\tilde{K}(\rho) + K^*)\|$ converges to a constant matrix as $J(\tilde{K}(\rho) + K^*) -$

$J(K^*)$ goes to infinity.

Informally, Lemma 8 proves the boundedness of the norm of the gradient along any high gain curve. As a consequence, it can never admit a class $\mathcal{K}_\infty$ PLI lower-bound, as we formalize in the following corollary.

Corollary 1: There is no $\mu > 0$ such that for all $K \in \mathbb{K}$ it holds that

$$\|\nabla J(K)\| \geq \sqrt{\mu(J(K) - J(K^*))},$$

i.e., the cost function J defined in (14) can never satisfy a gPLI.

From these results, one would think that J has a globally bounded gradient norm, since Lemma 8 shows that the gradient converges to a constant matrix for any high gain curve. However, let $\partial\mathcal{G} := \{K \in \mathbb{R}^{m \times n} \mid \lambda_{\max}(A - BK) = 0\}$ be the border of stability, i.e. the values of K for which $(A - BK)$ has at least one eigenvalue on the imaginary axis, and notice that the value of the gradient (17) explodes to infinity as K approaches $\partial\mathcal{G}$, as we show in the following lemma.

Lemma 9: For any $\bar{K} \in \partial\mathcal{G}$, let $\{K_i\}$ for $i = 1, 2, \dots$ be a sequence of matrices $K_i \in \mathcal{G}$ such that $\lim_{i \rightarrow \infty} K_i = \bar{K}$, then

$$\lim_{i \rightarrow \infty} \|\nabla J(K_i)\|_F = \infty.$$

As a consequence of this fact, gradient $\nabla J(\cdot)$ does not admit a global upper bound, and thus does not satisfy the conditions for Lemma 6.

The fact that the continuous-time LQR cost neither satisfies a gPLI, nor has a globally bounded gradient, makes it hard to provide tight global convergence rate estimates to the policy-optimization algorithm. In fact, the solution can be either exponential or linear-exponential, depending on which region of the state-space it is initialized at. Fortunately, since the gradient is only unbounded for “bounded directions” in the border of the space $\mathcal{G}$, we can enforce boundedness of the norm of the gradient on a subset of $\mathcal{G}$ as follows:

Lemma 10: For a $\delta > 0$, let $\mathcal{G}_\delta := \{K \in \mathcal{G} \mid A - BK + \delta I \text{ is Hurwitz}\}$, then there exists $g > 0$ such that $\|\nabla J(K)\| < g$ for all $K \in \mathcal{G}_\delta$.

The results presented so far illustrate, through the LQR policy optimization problem, the value of understanding exactly what kind of “PLI-like” condition the cost function in question satisfies. To complement our analysis so far, and in hopes of illustrating the different possible behaviors for the solution of the policy optimization for the LQR, we next provide an analysis of the single-input single-state/output LQR case.

B. Convergence analysis of the scalar LQR policy optimization

For the scalar case, the continuous-time system dynamics is given by (12) with $A = a \in \mathbb{R}$, and $B = b = 1$, the later being assumed without loss of generality, since for different values of b , its magnitude could simply be included in the magnitude of the considered input signal u . The weighting matrices of the LQR cost are $Q = q \in \mathbb{R}$ and $R = r \in \mathbb{R}$.

Then, the cost and its gradient can be computed for the scalar case as

$$J(k) = p,$$

and

$$\partial J(k) = -2(p - rk)\ell,$$

where p and ℓ solve (15) and (18), respectively, for scalar parameters. From this, we can recover the results of Lemma 8 for the scalar case. Notice that

$$\begin{aligned} \partial J(k) &= -\frac{2rk(a - k) + (rk^2 + q)}{2(a - k)^2} \\ &= -\frac{2r(\frac{a}{k} - 1) + (r + \frac{q}{k^2})}{2(\frac{a}{k} - 1)^2}, \end{aligned}$$

which implies that $\lim_{k \rightarrow +\infty} \partial J(k) = r/2$, indicating linear convergence as the magnitude of the feedback gain k increases while still keeping $a - k < 0$. This linear convergence, however, becomes less noticeable as the solution approaches the optimum feedback value k^* and instead exponential convergence is observed. To show that we compute

$$m(k) = \frac{\|\partial J(k)\|^2}{J(k) - J(k^*)}$$

for $k = k^* + \epsilon$, and verify that as $\epsilon \rightarrow 0$, $m(k) \rightarrow c$ for some positive constant c . That is done in the following proposition

Proposition 1: For the scalar LQR, let k^* be the value of the feedback gain that minimizes the cost $J(k)$. Then we have that

$$J(k^* + \epsilon) - J(k^*) = \ell r \epsilon^2 =: \delta(\epsilon) \quad (20)$$

$$\partial J(k^* + \epsilon) = \ell(-\delta(\epsilon) + r\epsilon) \quad (21)$$

$$m(k^* + \epsilon) = r\ell(\ell^2 \epsilon^2 - 2\ell\epsilon + 1). \quad (22)$$

These computations allow us to conclude that $m(k^*) = r\ell^* > 0$, characterizing exponential convergence near k^* . Furthermore, $\lim_{\epsilon \rightarrow a - k^*} m(k^* + \epsilon) = \infty$, indicating that the convergence rate explodes for values of k in the boundary of stability.

The result above was already known from the general case, but becomes clearer when derived for the scalar case, independent of matrix equalities and different “high gain trajectories”. Furthermore, the simpler form of the scalar case makes it easier to analyze numerical results.

Take $a = q = r = 1$ and notice from Fig. 3 that the gradient $\partial J(k)$ behaves completely differently if $k > k^*$ or if $k < k^*$. However, notice that if we restrict the domain from $[1, \infty]$ to $[1 + \epsilon, \infty]$, the value of the gradient is now globally bounded. This is the intuition behind Lemma 10.

Furthermore, the linear-exponential convergence behavior described in Lemma 6 becomes very evident for the scalar LQR if k is initialized larger than k^* , as can be seen in Fig. 4a. If, however, the solution is initialized near the border of instability, the convergence is much closer to a decreasing exponential, as evident in Fig. 4b.

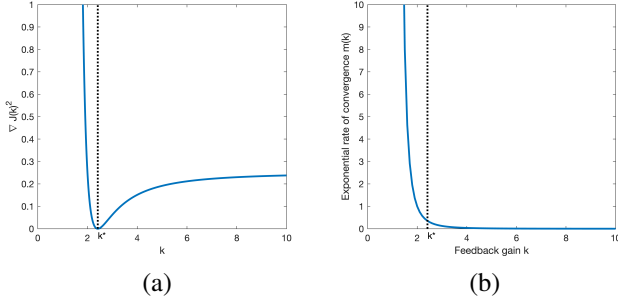


Fig. 3. Illustration of the dual behavior of the LQR cost function (14) for the scalar case. In (a) we plot the squared norm of the gradient $\|\nabla J\|^2$ as a function of the feedback gain k , while in (b) we plot the largest exponential rate of convergence $m(k)$ as defined in (22), and in both plots the dotted vertical line indicates the optimal k^* . Notice that for $k > k^*$ (right side of the dotted line) the gradient is bounded above, and $m(k)$ quickly goes to zero, while for $k < k^*$ (left side of the dotted line) the value of the gradient, and the best exponential rate of convergence both quickly diverge to infinity as k approaches the border of instability.

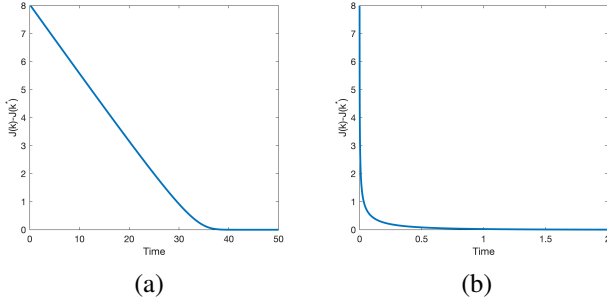


Fig. 4. Simulation results for the gradient flow of the scalar LQR policy optimization with $a = r = q = 1$. Both simulations (a) and (b) were initialized such that $J(k(0)) - J(k^*) \approx 8$, however (a) was initialized for $k(0) > k^*$ and (b) for $k(0) < k^*$. Notice that in (a) the convergence is “linear exponential” as described in Section II-B since, as can be seen in Fig. 3, for $k > k^*$, $\nabla J(k)$ is bounded above. In (b), on the other hand, the convergence is much quicker and exponential, due to the fact that for $k \in [a, k^*]$, the exponential rate of convergence $m(\epsilon)$ defined in (22) is bounded away from zero.

C. Comments on the difference between continuous and discrete-time LQR policy optimization

We conclude the analysis of this paper with a brief overview of the behavior of the discrete-time LQR policy optimization problem. This scenario is studied in different papers in the literature [14]–[17] and for this problem, a global PLI is characterized (see, for example, Lemma 1 in [16]). This is surprising since, by Corollary 1, the continuous-time LQR policy optimization can never admit a global PLI.

Some intuition behind this difference can be obtained by looking at the Euler discretization of the scalar continuous case with step-size $h > 0$, and with $R_d = hr$ and $Q_d = hq$. For this problem, we can define $m_d(k_d^* + \epsilon)$ similarly to how it was done for the continuous-time case in (22).

Furthermore, notice that the feedback gain k_d is bounded between $a < k_d^* + \epsilon < (2 + ha)/h$, which implies that $\mathcal{G}$ is compact in the discrete-time. Moreover, upon explicit computation of $m_d(k_d^* + \epsilon)$, one can check that for any $h > 0$,

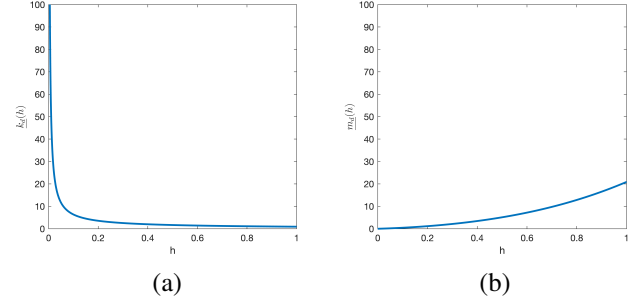


Fig. 5. Visualization of how the discretization step affects the global exponential rate of convergence for the scalar discrete-time LQR policy optimization problem. Notice that for any discretization step $h > 0$, there exists a $\mu = m_d(h)$ that provides global exponential convergence guarantees to the solution, however, that rate of convergence goes to zero as h goes to zero, which is compatible with the observation that CT LQR does not have a global exponential rate of convergence.

$m_d(k_d, h) \rightarrow \infty$ if either $k_d \rightarrow a$ or $k_d \rightarrow (2 + ha)/h$, which implies that $m_d(\cdot, h)$ must admit a minimum value $\underline{m}_d(h)$ attained at some point $\underline{k}_d \in \mathcal{G}$, i.e. $\underline{m}_d(h) = m_d(\underline{k}_d, h)$.

With these established, we pick $a = q = r = 1$ and plot $\underline{k}_d(h)$ and $\underline{m}_d(h)$ in Fig. 5. Notice that as $h \rightarrow 0$ (i.e. as we approach the CT LQR policy optimization problem), the value $\underline{k}_d(h)$ at which $m_d(k_d, h)$ is minimized goes to infinity, while $\underline{m}_d(h)$ goes to zero.

IV. CONCLUSIONS AND FUTURE WORKS

In this paper, we present a brief overview of convergence guarantees for gradient methods in optimization problems. We revisited the Polyak-Łojasiewicz inequality (i.e. gradient dominance condition) and observed how slight changes in its characterization can imply significant changes to the convergence of the gradient flow solution. This motivated the introduction of nonlinear comparison functions as a way of characterizing the behavior of the solution, which we supported with a result that gives conditions for the solution to present a “linear-exponential” behavior in Lemma 6.

The paper follows up with a scenario where the traditional PLI condition does not hold: the continuous-time model-free linear quadratic regulator problem. For this problem we showed that it presents neither globally exponential, nor globally linear-exponential convergence behavior, but a mixture of both depending on how close the solution is initialized to the border of instability. Despite that, we show in Lemma 8 that for any “high gain curve” in the space of stabilizing feedback matrices, the gradient is upper-bounded, which allowed us, through Lemma 10, to characterize global linear-exponential convergence behavior through a judicious restriction of the optimization search space. We then illustrate our results through numerical simulations of the scalar case, where the two regions of the parameter space are clearly defined.

A. A brief comment on lasso and proximal gradient

To finish the paper, we offer a brief comment illustrating the consequences of adding an L_1 regularization term to the cost in (1), and solving it through proximal gradient flow

instead of gradient flow (we refer the reader to [24], [25] for a complete characterization of the expressions used here). Assume $\mathcal{X} \subseteq \mathbb{R}$ for simplicity (scalar parameter), then the optimization problem with an L_1 regularization is written as follows

$$\begin{aligned} & \underset{x}{\text{minimize}} && f(x) + |x| \\ & \text{subject to} && x \in \mathcal{X} \end{aligned}, \quad (23)$$

which cannot be solved through gradient flow due to the non-differentiability of $|\cdot|$. Instead, a solution is found through proximal gradient flow, which can be written as

$$\dot{x} = -(x - \text{prox}_{|\cdot|}(x - \nabla f(x))),$$

where $\text{prox}_{|\cdot|}(x)$ is itself the solution of an optimization problem, but which has the following closed-form solution since we use L_1 -norm:

$$\text{prox}_{|\cdot|}(x) = \text{sign}(x) \max(|x| - 1, 0),$$

which after substituting to the proximal gradient flow results in

$$\dot{x} = -x + \text{sign}(x - \nabla f(x)) \max(|x - \nabla f(x)| - 1, 0).$$

At this point we note informally that if f satisfies the conditions in Lemma 6 (specifically global boundedness of the gradient), then there must exist a $\epsilon > 0$ such that for all x , if $f(x) - \underline{f} \geq \epsilon$, then $x - \nabla f(x) \geq 1$. Therefore, for initializations “large enough” (in the sense that $f(x_0) - \underline{f} \geq \epsilon$) the proximal gradient flow simplifies to

$$\dot{x} = -\nabla f(x) - 1,$$

Alternatively, if $\mathcal{X}$ is such that any sequence $\{x_i\}$ approaching a finite point in its boundary is such that $\lim_{i \rightarrow \infty} \nabla f(x_i) = \infty$, then at a point $x \in \mathcal{X}$ “close enough” to the boundary, $\nabla f(x) \gg x$ which simplifies the proximal gradient flow expression to

$$\dot{x} = -\nabla f(x) + 1.$$

Neither case changes the global behavior predictions given in this paper for the solution while it remains far away from $\mathcal{T}$ (in some sense). In fact, one can argue that if either $|x| \gg 1$ or $|\nabla f(x)| \gg 1$ then $|x - \nabla f(x)| - 1 < 0$ only if $x \approx \nabla f(x)$ (in the sense that $|x - \nabla f(x)|$ is much smaller than either $|x|$ or $|\nabla f(x)|$), which means that $\dot{x} = -x$ is an approximation of $\dot{x} = -\nabla f(x)$. However, notice that it is not straightforward to provide an equivalent local analysis.

Beyond its practical relevance, this observation motivates technically future works on proximal gradient flows, and its equivalent PLIs, which we believe are a natural follow-up to this publication.

REFERENCES

- [1] L. Cui, Z.-P. Jiang, and E. D. Sontag, “Small-disturbance input-to-state stability of perturbed gradient flows: Applications to LQR problem,” *Systems & Control Letters*, vol. 188, p. 105804, 2024.
- [2] E. D. Sontag, “Remarks on input to state stability of perturbed gradient flows, motivated by model-free feedback control learning,” *Systems & Control Letters*, vol. 161, p. 105138, Mar. 2022. [Online]. Available: <https://linkinghub.elsevier.com/retrieve/pii/S0167691122000056>
- [3] A. C. B. De Oliveira, M. Siami, and E. D. Sontag, “Remarks on the gradient training of linear neural network based feedback for the LQR problem,” in *2024 IEEE 63rd Conference on Decision and Control (CDC)*, 2024, pp. 7846–7852.
- [4] A. C. B. de Oliveira, M. Siami, and E. D. Sontag, “Convergence analysis of overparametrized LQR formulations,” *arXiv preprint arXiv:2408.15456*, 2024.
- [5] H. Mohammadi, A. Zare, M. Soltanolkotabi, and M. R. Jovanović, “Convergence and sample complexity of gradient methods for the model-free linear-quadratic regulator problem,” *IEEE Transactions on Automatic Control*, vol. 67, no. 5, pp. 2435–2450, 2021.
- [6] I. Fatkhullin and B. Polyak, “Optimizing static linear feedback: Gradient method,” *SIAM Journal on Control and Optimization*, vol. 59, no. 5, pp. 3887–3911, 2021.
- [7] J. Bhandari and D. Russo, “Global optimality guarantees for policy gradient methods,” *Operations Research*, vol. 72, no. 5, pp. 1906–1927, 2024. [Online]. Available: <https://doi.org/10.1287/opre.2021.0014>
- [8] A. Eftekhari, “Training linear neural networks: Non-local convergence and complexity results,” in *International Conference on Machine Learning*. PMLR, 2020, pp. 2836–2847.
- [9] S. P. Boyd and L. Vandenberghe, *Convex Optimization*. Cambridge, England: Cambridge University Press, 2004.
- [10] V. Chvátal, *Linear programming*. London, UK: Macmillan, 1983.
- [11] B. T. Polyak, “Gradient methods for minimizing functionals,” *USSR Computational Mathematics and Mathematical Physics*, vol. 3, no. 4, pp. 643–653, 1963.
- [12] H. Karimi, J. Nutini, and M. Schmidt, “Linear convergence of gradient and proximal-gradient methods under the Polyak-Łojasiewicz condition,” in *Machine Learning and Knowledge Discovery in Databases*, P. Frasconi, N. Landwehr, G. Manco, and J. Vreeken, Eds. Cham: Springer International Publishing, 2016, pp. 795–811.
- [13] Y. Watanabe and Y. Zheng, “Revisiting strong duality, hidden convexity, and gradient dominance in the linear quadratic regulator,” *arXiv preprint arXiv:2503.10964*, 2025.
- [14] M. Fazel, R. Ge, S. Kakade, and M. Mesbahi, “Global convergence of policy gradient methods for the linear quadratic regulator,” in *International Conference on Machine Learning*. PMLR, 2018, pp. 1467–1476.
- [15] Y. Sun and M. Fazel, “Learning optimal controllers by policy gradient: Global optimality via convex parameterization,” in *2021 60th IEEE Conference on Decision and Control (CDC)*, 2021, pp. 4576–4581.
- [16] B. Hu, K. Zhang, N. Li, M. Mesbahi, M. Fazel, and T. Başar, “Toward a theoretical foundation of policy optimization for learning control policies,” *Annual Review of Control, Robotics, and Autonomous Systems*, vol. 6, no. 1, pp. 123–158, 2023.
- [17] H. Mohammadi, A. Zare, M. Soltanolkotabi, and M. R. Jovanović, “Convergence and sample complexity of gradient methods for the model-free linear-quadratic regulator problem,” *IEEE Transactions on Automatic Control*, vol. 67, no. 5, pp. 2435–2450, May 2022. [Online]. Available: <https://ieeexplore.ieee.org/document/9448427/>
- [18] H. Mohammadi, M. Soltanolkotabi, and M. R. Jovanović, “On the lack of gradient domination for linear quadratic Gaussian problems with incomplete state information,” in *2021 60th IEEE Conference on Decision and Control (CDC)*, 2021, pp. 1120–1124.
- [19] A. C. B. de Oliveira, L. Cui, and E. Sontag, “Remarks on the Polyak-Łojasiewicz inequality and the convergence of gradient systems,” *arXiv preprint arXiv:2503.23641*, 2025.
- [20] S. Łojasiewicz, “Sur les trajectoires du gradient d’une fonction analytique. (Trajectories of the gradient of an analytic function),” *Semin. Geom., Univ. Studi Bologna*, vol. 1982/1983, pp. 115–117, 1984.
- [21] D. Angeli, E. D. Sontag, and Y. Wang, “A characterization of integral input-to-state stability,” *IEEE Transactions on Automatic Control*, vol. 45, no. 6, pp. 1082–1097, 2000.
- [22] E. D. Sontag *et al.*, “Smooth stabilization implies coprime factorization,” *IEEE Transactions on Automatic Control*, vol. 34, no. 4, pp. 435–443, 1989.
- [23] T. Rautert and E. W. Sachs, “Computational design of optimal output feedback controllers,” *SIAM Journal on Optimization*, vol. 7, no. 3, pp. 837–852, Aug. 1997. [Online]. Available: <http://epubs.siam.org/doi/10.1137/S1052623495290441>
- [24] S. Hassan-Moghaddam and M. R. Jovanović, “Proximal gradient flow and Douglas–Rachford splitting dynamics: Global exponential stability via integral quadratic constraints,” *Automatica*, vol. 123, p. 109311, 2021.
- [25] A. Gokhale, A. Davydov, and F. Bullo, “Proximal gradient dynamics: Monotonicity, exponential convergence, and applications,” *IEEE Control Systems Letters*, vol. 8, pp. 2853–2858, 2024.