

PI^σ - PI^σ Continuous Iterative Learning Control for Nonlinear Systems with Arbitrary Relative Degree

Lorenzo Cenceschi¹, Franco Angelini^{1,2}, Cosimo Della Santina^{3,4,5}, and Antonio Bicchi^{1,2,6}

Abstract—Online-Offline Iterative Learning Control provides an effective and robust solution to learn precise trajectory tracking when dealing with repetitive tasks. Yet, these algorithms were developed under the assumption that the relative degree between input and output is one. This prevents applications in many practically meaningful situations - e.g. mechanical systems control. To overcome this issue, this manuscript proposes a PI^σ - PI^σ algorithm fusing information from the whole visible dynamics. We provide sufficient convergence conditions when the controlled system has a generic constant relative degree, and it is possibly subject to measurement delay. The controller is validated on several simulation scenarios, including learning to swing-up a soft pendulum.

I. INTRODUCTION

Since its introduction [1], [2], Iterative Learning Control (ILC) has imposed itself as an effective and simple way of implementing precise repetitive motions. An in depth overview on recent advancements in ILC literature is provided by the survey papers [3], [4]. In its classic offline form, ILC iteratively learns an open-loop action by combining error evolutions from previous iterations [5], [6].

Most of ILC algorithms are designed under the hypothesis that the relative degree r is one. Yet, this condition is not fulfilled in many practically meaningful application domains. For example, relative degree is two for collocated control of mechanical systems, and may get up to twice the number of the degrees of freedoms for the underactuated case. A solution is provided in [7], where the tracking error is differentiated r times. Note that the first $r-1$ derivatives can be obtained as combination of systems states (see [8] and Sec. II-A). Therefore the algorithm is said to be of D-type, since it requires to apply a further numerical derivation to the system state - with well-known issues in terms of numerical stability. P-type algorithms are those which use only combinations of the system state to produce the learning signal. A P-type ILC algorithm for systems with arbitrary relative degree is proposed in [9], where the derivative operator is

replaced by a temporal shift. An alternative solution based on sample data update law is proposed in [10]. A learning rule for an underactuated robot with relative degree two is proposed and experimentally validated in [11].

However, all these ILC algorithms are of the offline kind. Due to purely open-loop nature, offline ILC is prone to robustness issues, for example when subject to iteration dependent disturbances [12] or to mismatches in the initial condition [13]. Feedback loops have been introduced to solve this issue, and several ways of combining the two components proposed [14]–[16]. Among them, a relevant trend goes under the names of “online-offline” or “current-iteration” or “open-closed-loop” ILC algorithms [17]–[19]. Here, the feedback controller is directly incorporated in the learning scheme, by adding the control action to the learning signal. Several instances of these algorithms have been proposed over the years, as for example including integral actions [20], a time-variable forgetting factor [21], or its extension to networked systems [22].

We are not aware of any offline-online ILC algorithm that can deal with systems with relative degree $r > 1$. In this paper, we propose a simple instance of such algorithm. A number of other effects are also considered, namely: (i) it is of type P, (ii) it can have a non-unitary forgetting factor, (iii) it can incorporate a delay in the feedback loop, (iv) it can include $\sigma \in \mathbb{N}$ integrators on both feedback action and ILC offline update. Due to the latter characteristics, and in accordance with the standard notation found in the state of the art, we refer to this algorithm as of PI^σ-PI^σ-type. This also serves to stress that all the signals used to produce online feedback and learning the feedforward are combination of state variables (and time integrations). No derivative (or other not causal operator) is used in this work.

II. PRELIMINARIES AND NOTATION

A. Considered system

Let us consider a control-affine nonlinear single input single output system with n states and a fixed relative degree r . According to classic results in nonlinear control theory [8], a set of coordinates can always be selected such that the nonlinear system has the structure (called normal form)

$$\dot{\xi} = \begin{bmatrix} 0 & 1 & 0 & \dots & 0 \\ 0 & 0 & 1 & \dots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \dots & 1 \\ 0 & 0 & \dots & \dots & 0 \end{bmatrix} \xi + \begin{bmatrix} 0 \\ 0 \\ \vdots \\ 0 \\ f(\xi, \eta) \end{bmatrix} + \begin{bmatrix} 0 \\ 0 \\ \vdots \\ 0 \\ g(\xi, \eta) \end{bmatrix} u, \quad (1)$$

$$\dot{\eta} = q(\xi, \eta), \quad y = (\xi)_1,$$

This research has received funding from the European Union’s Horizon 2020 Research and Innovation Programme under Grant Agreement No. 780883 (THING), No. 871237 (SOPHIA), No. 101016970 (NI). ¹Centro di Ricerca “Enrico Piaggio”, Università di Pisa, Largo Lucio Lazzarino 1, 56122 Pisa, Italy. ²Dipartimento di Ingegneria dell’Informazione, Università di Pisa, Largo Lucio Lazzarino 1, 56122 Pisa, Italy. ³Cognitive Robotics Department, Delft University of Technology, 2628 CD Delft, The Netherlands. ⁴Robotic Mechatronic Center (RMC), Institute of Robotics and Mechatronics, German Aerospace Center (DLR), D-82234 Oberpfaffenhofen, Germany. ⁵Technical University Munich, Chair of Sensor Based Robots and Intelligent Assistance Systems, Department of Informatics, D-85748 Garching, Germany. ⁶Soft Robotics for Human Cooperation and Rehabilitation, Fondazione Istituto Italiano di Tecnologia, via Morego, 30, 16163 Genova, Italy. Contacts: {l.cenceschi, cosimodellasantina}@gmail.com

where $u \in \mathbb{R}$ and $y \in \mathbb{R}$ are the control input and output, respectively. $\xi \in \mathbb{R}^r$ is the portion of the state vector visible from the output y , and $(\xi)_i$ is the i -th element. Its dynamics is described by a chain of r integrators. $q : \mathbb{R}^r \times \mathbb{R}^{n-r} \rightarrow \mathbb{R}^{n-r}$ is the internal dynamics of the system with state vector $\eta \in \mathbb{R}^{n-r}$. Note that the time derivatives of the output fulfill $y^{(i)} = (\xi)_{i+1}$ for all $i \in \{0, \dots, r-1\}$. Finally, for the sake of conciseness of notation, we take $x \triangleq [\eta^T, \xi^T]^T \in \mathbb{R}^n$ as the state vector, and we re-write (1) in the compact form

$$\dot{x}(t) = F(x(t)) + G(x(t))u(t), \quad y(t) = (x)_{n-r+1}(t). \quad (2)$$

where $F : \mathbb{R}^n \rightarrow \mathbb{R}^n$ is the drift vector and $G = [0 \dots 0 \ g]^T \in \mathbb{R}^n$ is the control vector field, with $g : \mathbb{R} \rightarrow \mathbb{R}$. The drift F is assumed to be globally uniformly Lipschitz, i.e. a finite $\bar{F} > 0$ exists such that

$$\|F(\zeta) - F(\chi)\| \leq \bar{F} \|\zeta - \chi\| \quad \forall \zeta, \chi \in \mathbb{R}^n. \quad (3)$$

Remark 1. System (4) could also be extended considering a chain of $\sigma \in \mathbb{N}$ integrators of the output y on top of ξ , still maintaining the same structure as in (1).

B. Iterations

In this work, we consider the improvement of tracking performance that can be achieved by repeating multiple times the same task. We therefore introduce the iteration index $k \in \mathbb{N}/\{0\}$, describing the trial-by-trial evolution. For each trial we focus on the system evolution in $t \in [0, t_f]$. The system (2) can be rewritten as

$$\dot{x}_k(t) = F(x_k(t)) + G(x_k(t))u_k(t), \quad y_k(t) = (x_k)_{n-r+1}(t). \quad (4)$$

C. Problem statement

We call $y_d : [0, t_f] \rightarrow \mathbb{R}$ the desired output evolution. This univocally defines a desired evolution of the augmented (i.e. including the σ integrators) ξ as $\begin{bmatrix} y_d^{(-\sigma)} & \dots & y_d^{(-1)} & y_d & y_d^{(1)} & \dots & y_d^{(r-1)} \end{bmatrix}^T(t)$, where $y_d^{(-i)} = \int_0^t \dots \int_0^{t_{i-1}} y_d(t_i) dt_i \dots dt_1$. Our goal is to find an iterative rule in the form

$$u_{k+1} = L(u_k, u_0, e_k), \quad (5)$$

such that

$$\lim_{k \rightarrow \infty} u_k = u_d, \quad (6)$$

where u_0 is an initial guess, and u_d is the input producing the output evolution y_d , i.e such that $\dot{x}_d = F(x_d) + G(x_d)u_d$, and $y_d = (x_d)_{n-r+1}$.

D. Further remarks on notation

It is instrumental for the derivation of the proposed approach to define the lambda-norm $\|\cdot\|_\lambda \triangleq \sup_t e^{-\lambda t} \|\cdot\|$ and the operator $\mathcal{C}_a^b(x) \triangleq \begin{bmatrix} \frac{t^{b-1}}{(b-1)!} & \dots & \frac{t^{a-1}}{(a-1)!} \end{bmatrix}^T \in \mathbb{R}^{b-a}$, for $b > a > 0$ and $t \in \mathbb{R}$. Note that no element in $\mathcal{C}_a^b(t)$ is equal to zero if $t \neq 0$.

III. MAIN RESULTS

We proposed the following linear instance of (5) for solving the problem

$$u_{k+1} = \underbrace{Qu_k + (1-Q)u_0}_{u_{k+1}^{\text{ff}}} + \underbrace{K^{\text{ff}}\epsilon_k + K^{\text{fb}}\epsilon_{k+1}}_{u_{k+1}^{\text{fb}}}, \quad (7)$$

where u_0 is the initial guess, $Q \in (0, 1]$ is the forgetting factor. Note that $K^{\text{ff}}, K^{\text{fb}} \in \mathbb{R}^{1 \times p}$ are feed-forward and feedback gains, respectively. As highlighted in Remark 1, we consider a system with σ integrators. Thus, the error vector is defined as $\epsilon_k(t) = [e^{(-\sigma)} \dots e^{(-1)} \ e \ e^{(1)} \dots e^{(r-1)}]^T(t) \in \mathbb{R}^p$, where $p = \sigma + r$, $e(t) \triangleq y_d(t) - y_k(t)$ output error relative to desired output y_d , $e^{(-i)} = \int_0^t \dots \int_0^{t_{i-1}} e(t_i) dt_i \dots dt_1$, and $e^{(i)} = \frac{d^i}{dt^i} e(t)$. Finally, u^{ff} and u^{fb} are the feedforward and the feedback component of the control action, respectively.

A. Main theorem

Theorem 1. Consider the system (4) with g constant in $\mathbb{R}^n$, controlled by (7). If

$$\|Q - (K^{\text{ff}} + QK^{\text{fb}})B\| < 1, \quad (8)$$

then a $\lambda > 0$ and a $a > 0$ exist such that

$$\lim_{k \rightarrow \infty} \|y_d - y_k\|_\lambda \leq (1-Q)a \|u_d - u_0\|_\lambda, \quad (9)$$

where $B \triangleq g e^{-\lambda t} \int_0^t \mathcal{C}_1^{p+1}(t-s) ds$.

Proof: For readability purpose, continuous-time variable is omitted where implicit. Let us define the the initial control error as $\Delta u_0 \triangleq u_d - u_0$ and the control error at the k -th trial as $\Delta u_{k+1} \triangleq u_d - u_{k+1}$. Then, we can write

$$\Delta u_{k+1} = Q\Delta u_k + (1-Q)\Delta u_0 - K^{\text{ff}}\epsilon_k - K^{\text{fb}}\epsilon_{k+1} \quad (10)$$

$$= \Delta u_{k+1}^{\text{ff}} - K^{\text{fb}}\epsilon_{k+1}, \quad (11)$$

then pure feed-forward learning rule is

$$\begin{aligned} \Delta u_{k+1}^{\text{ff}} &= Q\Delta u_k + (1-Q)\Delta u_0 - K^{\text{ff}}\epsilon_k \\ &= Q\Delta u_k^{\text{ff}} + (1-Q)\Delta u_0 - K\epsilon_k, \end{aligned}$$

where $K \triangleq K^{\text{ff}} + QK^{\text{fb}}$. Adding and subtracting $KB\Delta u_k^{\text{ff}}$ leads to

$$= (Q - KB)\Delta u_k^{\text{ff}} + (1-Q)\Delta u_0 - K(\epsilon_k - B\Delta u_k^{\text{ff}}). \quad (12)$$

The error vector ϵ_k can be written as a function of the r -th Lie-derivative of system output error e_k , using Cauchy's formula for repeated integrals

$$\epsilon_k = \int_0^t \mathcal{C}_1^{p+1}(t-s) [\Delta f_k - gK^{\text{fb}}\epsilon_k + g\Delta u_k^{\text{ff}}] ds, \quad (13)$$

where f and g are defined as in (1), and $\Delta f_k(s) \triangleq f(x_d(s)) - f(x_k(s))$. Take norm of both sides

$$\|\epsilon_k\| \leq \int_0^t \left\| \mathcal{C}_1^{p+1}(t-s) \right\| \left(|\Delta f_k| + |gK^{\text{fb}}|\epsilon_k + |g|u_k^{\text{ff}} \right) ds. \quad (14)$$

Note that f is continuous and Lipschitz in y , then exists $\bar{f}$ positive constant such that $|f(y_d) - f(y_k)| \leq \bar{f}|y_d - y_k| \leq$

$\bar{f} \|\epsilon_k\|$. Substituting it in (14) and using Grönwall-Bellman theorem ([23], Corollary 2) leads to

$$\|\epsilon_k\| \leq |g| \int_0^t \exp \left((\bar{f} + |gK^{\text{fb}}|) \int_s^t \left\| \mathcal{C}_1^{p+1}(t-h) \right\| dh \right) \cdot \left\| \mathcal{C}_1^{p+1}(t-s) \right\| |\Delta u_k^{\text{ff}}| ds. \quad (15)$$

An upper bound of $\|\mathcal{C}(t-s)\|$ can be set using the ∞ -norm and the exponential series expansion: $\|\mathcal{C}(t-s)\| \leq e^{t-s}$. Furthermore, since $t \in [0, t_f]$, $s \in [0, t]$, and $h \in [s, t]$, we can set $\|\mathcal{C}(t-h)\| \leq \|\mathcal{C}(t_f)\|$. Substituting in (15) leads to

$$\|\epsilon_k\| \leq |g| \int_0^t e^{(\bar{f} + |gK^{\text{fb}}|) \|\mathcal{C}_1^{p+1}(t_f)\| (t-s)} e^{t-s} |\Delta u_k^{\text{ff}}| ds. \quad (16)$$

Let $b \triangleq (\bar{f} + |gK^{\text{fb}}|) \|\mathcal{C}_1^{p+1}(t_f)\|$, and define the upper bound $|\Delta u_k^{\text{ff}}(s)| = e^{\lambda s} e^{-\lambda s} |\Delta u_k^{\text{ff}}(s)| \leq e^{\lambda s} \|\Delta u_k^{\text{ff}}(s)\|_{\lambda}$, then

$$\begin{aligned} \|\epsilon_k\| &\leq |g| \int_0^t e^{(b+1)(t-s)} e^{\lambda s} ds \|\Delta u_k^{\text{ff}}\|_{\lambda} \\ &= |g| \int_0^t e^{(b+1)t + (\lambda - (b+1))s} ds \|\Delta u_k^{\text{ff}}\|_{\lambda} \\ &= |g| \frac{e^{\lambda t} - e^{(b+1)t}}{\lambda - (b+1)} \|\Delta u_k^{\text{ff}}\|_{\lambda}. \end{aligned} \quad (17)$$

Multiplying both sides by positive values $e^{-\lambda t}$ and taking the supremum we obtain

$$\|\epsilon_k\|_{\lambda} \leq |g| \sup_t \frac{1 - e^{((b+1)-\lambda)t}}{\lambda - (b+1)} \|\Delta u_k^{\text{ff}}\|_{\lambda}. \quad (18)$$

Computing the norm of both sides of (12), and substituting in (18) leads to

$$\begin{aligned} |\Delta u_{k+1}^{\text{ff}}| &\leq |Q - KB| |\Delta u_k^{\text{ff}}| + (1-Q) |\Delta u_0| \\ &\quad + \|K\| \|\epsilon_k\| + \|K\| \|B\| |\Delta u_k^{\text{ff}}| \\ &\leq |Q - KB| |\Delta u_k^{\text{ff}}| + (1-Q) |\Delta u_0| \\ &\quad + \|K\| |g| \frac{e^{\lambda t} - e^{(b+1)t}}{\lambda - (b+1)} \|\Delta u_k^{\text{ff}}\|_{\lambda} \\ &\quad + \|K\| \|B\| |\Delta u_k^{\text{ff}}|. \end{aligned} \quad (19)$$

Multiplying both sides by positive values $e^{-\lambda t}$ and taking the supremum we obtain

$$\begin{aligned} |\Delta u_{k+1}^{\text{ff}}|_{\lambda} &\leq \sup_t |Q - KB| \|\Delta u_k^{\text{ff}}\|_{\lambda} + (1-Q) \|\Delta u_0\|_{\lambda} \\ &\quad + \|K\| |g| \sup_t \frac{1 - e^{((b+1)-\lambda)t}}{\lambda - (b+1)} \|\Delta u_k^{\text{ff}}\|_{\lambda} \\ &\quad + \|K\| \sup_t \|B\| \|\Delta u_k^{\text{ff}}\|_{\lambda}. \end{aligned} \quad (20)$$

At this point, we evaluate an upper bound of $\sup_t \|B\|$

$$\begin{aligned} \sup_t \|B\| &\leq |g| \sup_t e^{-\lambda t} \int_0^t \mathcal{C}_1^{p+1}(t-s) ds \\ &\leq |g| \sup_t e^{-\lambda t} \mathcal{C}_2^{p+2}(t) \leq |g| \frac{1}{\lambda e}. \end{aligned} \quad (21)$$

Substitute (21) in (20), and let us define

$$\rho \triangleq \sup_t |Q - KB| + \|K\| |g| \left(\frac{1}{\lambda - (b+1)} + \frac{1}{\lambda e} \right), \quad (22)$$

if $\lambda > b+1$, then (20) can be rewritten as

$$\begin{aligned} |\Delta u_{k+1}^{\text{ff}}|_{\lambda} &\leq \rho \|\Delta u_k^{\text{ff}}\|_{\lambda} + (1-Q) \|\Delta u_0\|_{\lambda} \\ &\leq \rho^{k+1} \|\Delta u_0^{\text{ff}}\|_{\lambda} + (1-Q) \|\Delta u_0\|_{\lambda} \sum_{i=0}^k \rho^i. \end{aligned} \quad (23)$$

If $\rho < 1$, for $k \rightarrow \infty$ we have

$$\lim_{k \rightarrow \infty} |\Delta u_{k+1}^{\text{ff}}|_{\lambda} \leq \frac{(1-Q)}{1-\rho} \|\Delta u_0\|_{\lambda}. \quad (24)$$

At this point, we find λ such that $\rho < 1$

$$\begin{aligned} \rho &\leq \sup_t |Q - KB| + \|K\| |g| \left(\frac{1}{\lambda - (b+1)} + \frac{1}{\lambda e} \right) < 1 \\ &\quad \frac{1}{\lambda e(\lambda - (b+1))} \left[\sup_t |Q - KB| \lambda e(\lambda - (b+1)) \right. \\ &\quad \left. + \|K\| |g| (\lambda(e+1) - (b+1)) \right] < 1 \\ &\quad \sup_t |Q - KB| \lambda + \left(1 + \frac{1}{e} \right) \|K\| |g| < \lambda - (b+1) \\ &\quad \left(1 + \frac{1}{e} \right) \|K\| |g| + b+1 < \lambda(1 - \sup_t |Q - KB|). \end{aligned} \quad (25)$$

If (8) holds, also $\sup_t |Q - KB| < 1$ holds true, then

$$\frac{(1 + \frac{1}{e}) \|K\| |g| + b+1}{1 - \sup_t |Q - KB|} < \lambda. \quad (26)$$

An upper estimation $\bar{Q}$ of $\sup_t |Q - KB|$ can be found as

$$\sup_t |Q - KB| \leq \max(Q - K^-, K^+ - Q) \triangleq \bar{Q}, \quad (27)$$

where

$$\begin{aligned} K^- &\triangleq \frac{g}{(b+1)e} \sum_{K_i < 0} K_i < \inf_t KB, \\ K^+ &\triangleq \frac{g}{(b+1)e} \sum_{K_i > 0} K_i > \sup_t KB. \end{aligned}$$

Therefore, in order to have $\rho < 1$, λ must be such that

$$\frac{(1 + \frac{1}{e}) \|K\| |g| + b+1}{1 - \bar{Q}} < \lambda.$$

At this point, we focus on the state error $\Delta x_k \triangleq x_d - x_k$. Taking its norm leads to

$$\|\Delta x_k\| \leq \int_0^t \|\Delta F_k\| + \|gK^{\text{fb}}\| \|\epsilon_k\| + |g| |\Delta u_k^{\text{ff}}| ds, \quad (28)$$

where $\Delta F_k \triangleq F(x_d) - F(x_k)$. From (3), F is Lipschitz, furthermore since e_k is continuous and derivable up to $r-1$ times, also ϵ_k (function of e_k) is Lipschitz in x . Therefore, a positive constant $\bar{E}$ exists such that $\|\epsilon_k\| \leq \bar{E} \|x_d - x_k\|$; let $c \triangleq \bar{F} + \|gK^{\text{fb}}\| \bar{E}$. Thus

$$\|\Delta x_k\| \leq \int_0^t c \|\Delta x_k\| + |g| |\Delta u_k^{\text{ff}}| ds. \quad (29)$$

Applying Grönwall-Bellman theorem

$$\begin{aligned}\|\Delta x_k\| &\leq |g| \int_0^t e^{c(t-s)} |\Delta u_k^{\text{ff}}| ds \\ &\leq |g| \int_0^t e^{c(t-s)} e^{\lambda s} ds \|\Delta u_k^{\text{ff}}\|_\lambda.\end{aligned}$$

An upper bound of the state error is

$$\|\Delta x_k\| \leq |g| e^{\lambda t} \frac{1 - e^{(c-\lambda)t}}{\lambda - c} \|\Delta u_k^{\text{ff}}\|_\lambda. \quad (30)$$

Substituting (24) into (30) and taking the supremum

$$\lim_{k \rightarrow \infty} \|\Delta x_k\|_\infty \leq |g| \sup_t e^{\lambda t} \frac{1 - e^{(c-\lambda)t}}{\lambda - c} \frac{1 - Q}{1 - \rho} \|\Delta u_0\|_\lambda. \quad (31)$$

Given the expression of y_k from (4) - i.e. extraction of a single element from x_k - we can affirm that $|y_d - y_k| \leq \|x_d - x_k\|$. Therefore

$$\begin{aligned}\lim_{k \rightarrow \infty} \|y_d - y_k\|_\infty &\leq \|x_d - x_k\|_\infty \\ &\leq |g| \sup_t e^{\lambda t} \frac{1 - e^{(c-\lambda)t}}{\lambda - c} \frac{1 - Q}{1 - \rho} \|\Delta u_0\|_\lambda.\end{aligned} \quad (32)$$

Multiplying both sides of (30) by $e^{-\lambda t}$ and taking the supremum

$$\|\Delta x_k\|_\lambda \leq |g| \sup_t \frac{1 - e^{(c-\lambda)t}}{\lambda - c} \|\Delta u_k^{\text{ff}}\|_\lambda. \quad (33)$$

Substituting (24) into (33)

$$\lim_{k \rightarrow \infty} \|\Delta x_k\|_\lambda \leq |g| \sup_t \frac{1 - e^{(c-\lambda)t}}{\lambda - c} \frac{1 - Q}{1 - \rho} \|\Delta u_0\|_\lambda, \quad (34)$$

therefore

$$\begin{aligned}\lim_{k \rightarrow \infty} \|y_d - y_k\|_\lambda &\leq \|x_d - x_k\|_\lambda \\ &\leq |g| \sup_t \frac{1 - e^{(c-\lambda)t}}{\lambda - c} \frac{1 - Q}{1 - \rho} \|\Delta u_0\|_\lambda.\end{aligned} \quad (35)$$

Defining $a \triangleq |g| \sup_t \frac{1 - e^{(c-\lambda)t}}{(\lambda - c)(1 - \rho)}$, then (9) follows. ■

Remark 2. The value a is such that $\lim_{\lambda \rightarrow \infty} a = 0$. This comes directly from the definition of a .

B. Non constant input field

Assumption of g constant in Theorem 1 can be relaxed to non constant input field thanks to the following Corollary.

Corollary 1. Suppose that all the hypothesis of Theorem 1 are verified, but the one on g being constant. Then, the thesis of Theorem 1 is still valid when using the actions

$$u_k = (\bar{g}/g(x_k))v_k, \quad (36)$$

and with $B \triangleq \bar{g} e^{-\lambda t} \int_0^t \mathcal{C}_1^{p+1}(t-s) ds$, where v_k is evaluated as in (7), and $\bar{g} \in \mathbb{R}$.

Proof: The corollary follows directly by noticing that (36) applied to (1) generates a system fulfilling the hypotheses of Theorem 1. It is worth noticing here that since

we hypothesized constant relative degree in Sec. II-A, then $g(x_k) \neq 0$, and (36) is always well defined. ■

Remark 3. Contrary to (7) - which is model-free net of having access to ξ - (36) introduces a model based component which is function of the whole state.

C. Further remarks

If the forgetting factor is $Q = 1$, i.e. learning scheme has perfect memory of the previous trial control, then if condition (8) is satisfied, limit (9) becomes

$$\lim_{k \rightarrow \infty} \|y_d - y_k\|_\lambda = 0, \quad (37)$$

and asymptotically convergent tracking is met.

Furthermore, (32) becomes

$$\lim_{k \rightarrow \infty} \|y_d - y_k\|_\infty = 0, \quad (38)$$

ensuring convergence also in ∞ -norm. However, there is no proof about monotonicity, since the upper bound (30) may initially grow, and together with it the tracking error [5]. An example of this behavior is provided in Sec. V-B.

Note that condition (8) is satisfied for $Q = 1$, if $0 < KB < 2$, i.e.

$$0 < g e^{-\lambda t} (K^{\text{fb}} + K^{\text{ff}})^{\frac{p+2}{2}}(t) < 2, \quad (39)$$

which decreases when λ increases. This equation can be further reformulated as

$$0 < e^{-\lambda t} \sum_{i=1}^p (K^{\text{fb}} + K^{\text{ff}})_i \frac{t^{p+1-i}}{(p+1-i)!} \leq g \frac{1}{\lambda e} \sum_{i=1}^p (K^{\text{fb}} + K^{\text{ff}})_i < 2. \quad (40)$$

Note that u_0 is the initial control guess derived by a priori information about the system. If no initial control guess is provided, i.e. $u_0 \equiv 0$ the control learning rule (7) becomes

$$\Delta u_{k+1} = Q \Delta u_k + K^{\text{fb}} \epsilon_k + K^{\text{ff}} \epsilon_{k+1}, \quad (41)$$

similar to ILC schemes in literature [24], where first trial behavior is entirely guided by feedback control.

IV. DELAYED SYSTEM

Theorem 1 can be extended considering a constant delay $\delta > 0$ on the feedback component, which modifies (4) in

$$u_{k+1}(t) = Q u_k(t) + (1 - Q) u_0(t) + K^{\text{ff}} \epsilon_k(t) + K^{\text{fb}} \epsilon_k(t - \delta), \quad (42)$$

where we define the error vector to be zero for negative time values, i.e. $\epsilon_k \equiv 0 \quad \forall t < 0$.

Remark 4. This rule is the generalization to generic n and σ , of an auto-regressive mathematical model that has proven to be a robust description of human motor adaptation during object manipulation [25].

Corollary 2. Theorem 1 still holds, when considering (42) instead of (7).

Proof: The idea is to follow the same steps of Theorem 1 proof except for some intermediate steps where we take care of delay that affects the feedback error.

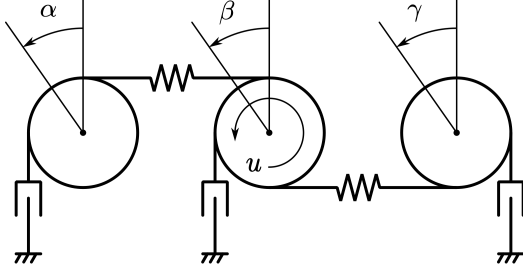


Fig. 1: Diagram of the system (47) used in Simulations 1–4.

Since by definition $u_k(t) = 0$ and $\epsilon_k = 0$ for all $t < 0$, the following upper bounds hold

$$\int_0^t \|\epsilon_k(s-\delta)\| ds = \int_0^{t-\delta} \|\epsilon_k(s)\| ds \leq \int_0^t \|\epsilon_k(s)\| ds, \quad (43)$$

$$\int_0^t |\Delta u_k^{\text{ff}}(s-\delta)| ds = \int_0^{t-\delta} |\Delta u_k^{\text{ff}}(s)| ds \leq \int_0^t |\Delta u_k^{\text{ff}}(s)| ds. \quad (44)$$

Then we can find an upper bound to (14) modified to take into account delays

$$\begin{aligned} \|\epsilon_k(t-\phi)\| &\leq \int_0^{t-\phi} \left\| \bar{C}_1^{p+1}(t-s) \right\| (|\Delta f_k| \\ &\quad + \|gK^{\text{fb}}\| \|\epsilon_k(s-\delta)\| + |g||u_k^{\text{ff}}|) ds \\ &\leq \int_0^t \left\| \bar{C}_1^{p+1}(t-s) \right\| (\bar{f} \|\epsilon_k\| \\ &\quad + \|gK^{\text{fb}}\| \|\epsilon_k(s-\delta)\| + |g||u_k^{\text{ff}}|) ds, \end{aligned} \quad (45)$$

where $\phi \in \{0, \delta\}$. Note that

$$\begin{aligned} \int_0^t \left\| \bar{C}_1^{p+1}(t-s) \right\| \|\epsilon_k(s-\delta)\| ds \\ \leq \int_0^t \left\| \bar{C}_1^{p+1}(t-s-\delta) \right\| \|\epsilon_k(s)\| ds. \end{aligned} \quad (46)$$

Furthermore, since $t-s-\delta < t-s < t_f$, then $\|\bar{C}(t-s-\delta)\| \leq \|\bar{C}(t_f)\|$, and applying Grönwall-Bellman theorem to (45) we obtain (16).

Taking into account previous upper bounds, the original proof follows straight-forward. ■

V. SIMULATIONS AND RESULTS

To show the effectiveness of the proposed methods, we simulate two different systems controlled by ILC schemes (7) and (42). The two set of simulations are discussed in two separate subsections. We choose the control gains K^{fb} and K^{ff} so to satisfy condition (8) in Theorem 1. Then, we check if the output transient response shows a convergent behavior trial after trial.

A. Mechanical system with $r = 4$

Simulation Setup. Consider the flexible mechanism [26] depicted in Fig. 1. Three rotating masses coupled by springs are controlled by an input u torque applied on the second mass. $\alpha, \beta, \gamma \in \mathbb{R}$ are the absolute rotation angles. Some of the passive elements have nonlinear characteristics. We

want to control the mass on the left. The system state-space representation is

$$\begin{aligned} \ddot{\alpha} &= 10(\beta - \alpha) - \dot{\alpha}, \\ \ddot{\beta} &= 10[(\gamma - \beta)^3 + (\gamma - \beta) - (\beta - \alpha)] - \dot{\beta} + u, \\ \ddot{\gamma} &= -10[(\gamma - \beta)^3 + (\gamma - \beta)] - \frac{1}{1 + \exp(-\dot{\gamma})}, \\ y &= \alpha. \end{aligned} \quad (47)$$

The normal form (1) is

$$\begin{aligned} (x)_3 &= (\xi)_1 = y = \alpha, \\ (x)_4 &= (\xi)_2 = y^{(1)} = \dot{\alpha}, \\ (x)_5 &= (\xi)_3 = y^{(2)} = \ddot{\alpha} = 10(\beta - y) - y^{(1)}, \\ (x)_6 &= (\xi)_4 = y^{(3)} = 10(\dot{\beta} - y^{(1)}) - y^{(2)}, \\ y^{(4)} &= 10(\ddot{\beta} - y^{(2)}) - y^{(3)} \\ &= 100[(\gamma - \beta)^3 + (\gamma - \beta) - (\beta - \alpha)] \\ &\quad - 10(\dot{\beta} + y^{(2)}) - y^{(3)} + 10u. \end{aligned}$$

Therefore the relative degree is $r = 4$. It is worth noting that a zero dynamics is also present: $x_1 = \eta_1 = \gamma$, $x_2 = \eta_2 = \dot{\gamma}$.

The reference trajectory is

$$y_d = e^{-4t} t^{10}, \quad (48)$$

with terminal time $t_f = 5$ s. The initial conditions are

$$\alpha(0) = \beta(0) = \gamma(0) = \dot{\alpha}(0) = \dot{\beta}(0) = \dot{\gamma}(0) = 0.$$

We employ the control law (47) with the ILC scheme (42) where

$$\begin{aligned} e &= y_d - y, \\ \epsilon &= [e \quad e^{(1)} \quad e^{(2)} \quad e^{(3)}]^T \in \mathbb{R}^4, \\ K^{\text{fb}} &= \begin{bmatrix} 0 & 1 & 0 & 0 \end{bmatrix} \in \mathbb{R}^{1 \times 4}, \\ K^{\text{ff}} &= \begin{bmatrix} 0 & 1 & 0 & 0.5 \end{bmatrix} \in \mathbb{R}^{1 \times 4}, \\ \delta &= 0.01. \end{aligned} \quad (49)$$

Since there is a delay $\delta \neq 0$, in the following we rely on Corollary 2 to apply the results of Theorem 1. We test this system with several choices of the forgetting factor Q and the initial guess u_0 .

Simulation 1. Consider the parameters

$$\begin{aligned} Q &= 0.95, \\ u_0(t) &= 0, \quad \forall t \in [0, t_f]. \end{aligned} \quad (50)$$

Given (49) and (50), we have that $\|Q - (K^{\text{ff}} + QK^{\text{fb}})B\| \approx 0.95$. Therefore, the convergence condition $\|Q - (K^{\text{ff}} + QK^{\text{fb}})B\| < 1$ in Theorem 1 is satisfied. Thus, we expect this system to converge but with non-zero steady state error, as described by (9).

Figs 2a–2b show the output and input responses at 1–st, 4–th, and 30–th trials, respectively.

Simulation 2. Consider the parameters

$$\begin{aligned} Q &= 0.95, \\ u_0(t) &= \cos(2t), \quad \forall t \in [0, t_f]. \end{aligned} \quad (51)$$

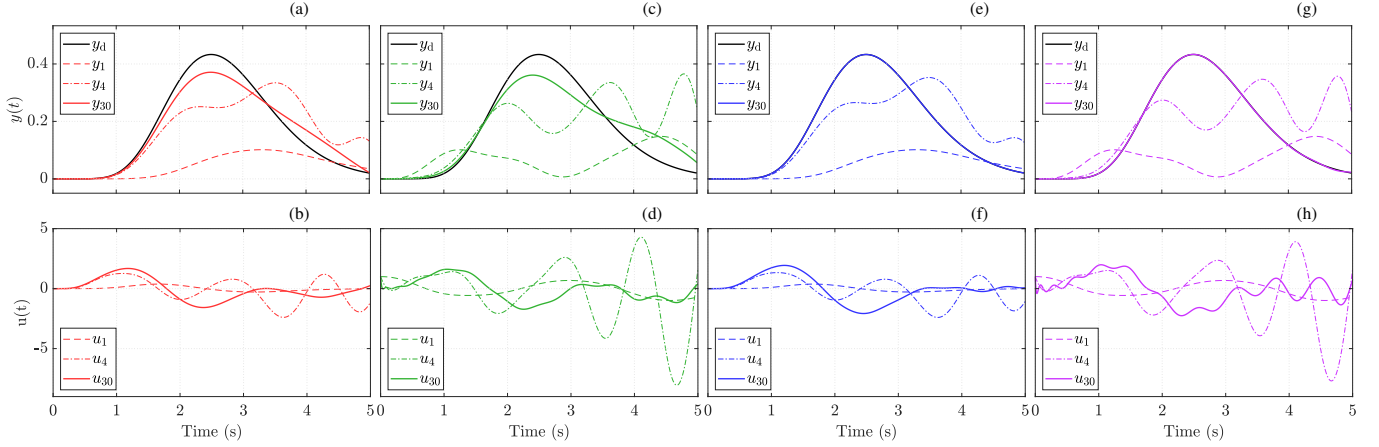


Fig. 2: Output and input signals for 1-st, 4-th, and 30-th trials. (a)-(b) are relative to Sim. 1 ($Q < 1$, $u_0(t) \equiv 0$), (c)-(d) to Sim. 2 ($Q < 1$, $u_0(t) \neq 0$), (e)-(f) to Sim. 3 ($Q = 1$, $u_0(t) \equiv 0$), and finally (g)-(h) to Sim. 4 ($Q = 1$, $u_0(t) \neq 0$).

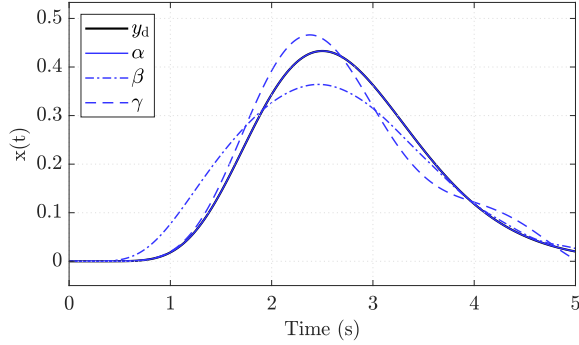


Fig. 3: System states for Simulation 3, 30-th iteration.

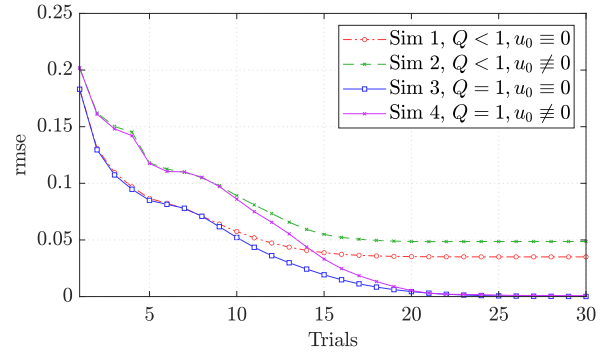


Fig. 4: Tracking rmse per trial.

As in Simulation 1, given (49) and (50), we have that $\|Q - (K^f + QK^b)B\| \approx 0.95$. Therefore, the convergence condition $\|Q - (K^f + QK^b)B\| < 1$ in Theorem 1 is satisfied. Thus, we expect this system to converge but with non-zero steady state error. However, in this case the choice of $u_0(t)$ is arbitrary, then we likely expect an error greater than Simulation 1 error, as described by (9).

Figs 2c–2d show the output and input responses at 1–st, 4–th, and 30–th trials, respectively.

Simulation 3. We employ the parameters

$$\begin{aligned} Q &= 1, \\ u_0(t) &= 0, \quad \forall t \in [0, t_f]. \end{aligned} \quad (52)$$

Given (49) and (52), convergence condition (40) is satisfied for $\lambda > 813$. Therefore, we expect this system to converge with zero steady state error, as described by (38)).

Figs 2e–2f show the output and input responses at 1–st, 4–th, and 30–th trials, respectively. Fig. 3 shows internal state behavior for the stable 30-th iteration.

Simulation 4. We employ the parameters

$$\begin{aligned} Q &= 1, \\ u_0(t) &= \cos(2t), \quad \forall t \in [0, t_f]. \end{aligned} \quad (53)$$

Given (49) and (52), convergence condition (40) is sat-

isfied for $\lambda > 813$. Therefore, we expect this system to converge with zero steady state error, as described by (38)).

Figs 2g–2h show the output and input responses at 1–st, 4–th, and 30–th trials, respectively.

B. 2D soft inverted pendulum

Simulation Setup. We consider the 2D soft inverted pendulum in Fig. 5 as introduced in [27]. This is a mechan-

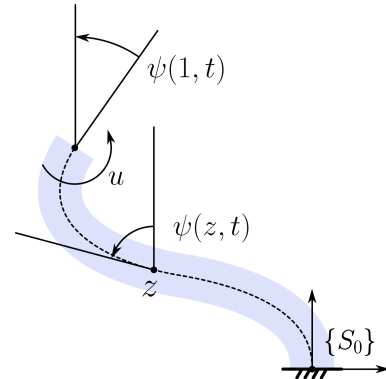


Fig. 5: Diagram of soft inverted pendulum, $z \in [0, 1]$ is the curvilinear coordinate along the axis. $\{S_0\}$ is the base frame, and $\psi(z, t)$ measures the orientation of each point along the robot w.r.t. that frame.

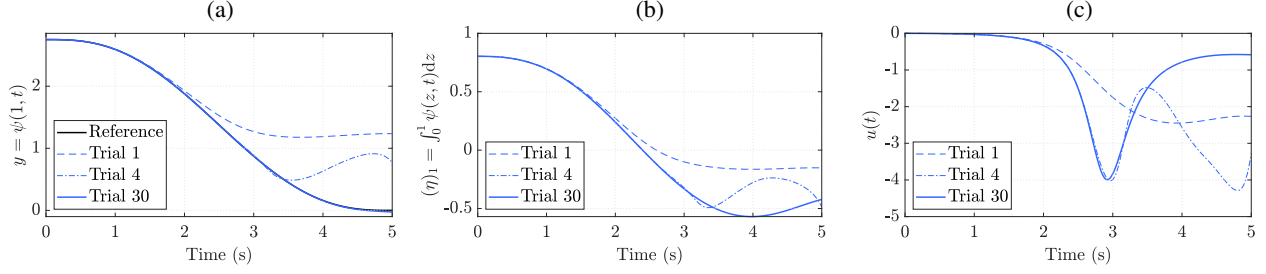


Fig. 6: State and control signals for 1-st, 4-th, and 30-th trials of Simulation 5. (a)-(b) show the two configuration variables, (c) shows the input torque u .

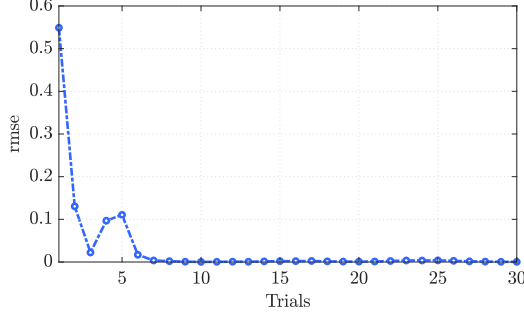


Fig. 7: Tracking root mean square error for Sim. 5.

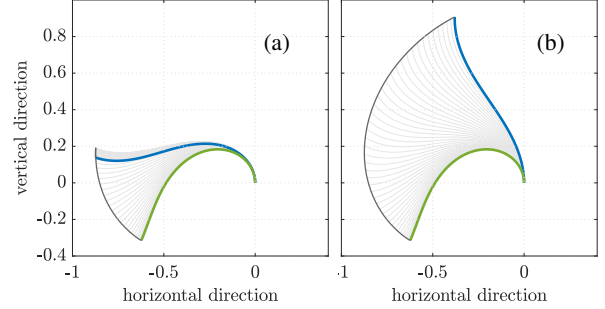


Fig. 8: Pendulum behavior for 1-st (a) and 30-th (b) iterations.

ical system made of continuously deformable soft material, which is described with just two degrees of freedom thanks to an affine curvature hypothesis (please refer to the paper for more details). Therefore $n = 4$. The system is under actuated through a single torque applied at the tip, as in figure. Consider $\psi : [0, 1] \times \mathbb{R} \rightarrow \mathbb{R}$ as the function associating to all points along the pendulum (parametrized in $[0, 1]$) their relative orientation w.r.t. the base, at any given time t . In [27] it is shown that the following can be taken as state vector for the robot.

$$\begin{aligned} (x)_1 &= (\eta)_1 = \int_0^1 \psi(z, t) dz, \\ (x)_2 &= (\eta)_2 = \left(\int_0^1 \psi(z, t) \right)^{(1)}(1, t), \\ (x)_3 &= (\xi)_1 = \psi(1, t), \\ (x)_4 &= (\xi)_2 = (\psi(1, t))^{(1)}. \end{aligned} \quad (54)$$

Note that this is in normal form (1), when $y = \psi(1, t)$. We will consider this output function in the rest of the section.

To perform simulations, we use the same mechanical characteristics as the simulations in [27]. Among all the stiffnesses considered there, we choose $\kappa = 1$. Thus the pendulum has three equilibrium points - vertical, bottom-left, bottom-right. The first being unstable, and the other two stable. We want to improve the swing up task, e.g. moving the pendulum from a stable equilibrium (bottom-left here w.l.o.g.) to a neighborhood of the unstable vertical configuration $x = 0$.

We apply Corollary 1 with $\bar{g} = g(0)$ so to scale the input field around the vertical unstable equilibrium to one. We want the tip orientation y to track the following swing up

behavior

$$y_d = (x)_3(0) \left[1 - 10 \left(\frac{t}{t_f} \right)^3 + 15 \left(\frac{t}{t_f} \right)^4 - 6 \left(\frac{t}{t_f} \right)^5 \right], \quad (55)$$

where the initial state of the system is on bottom-left equilibrium $x(0) = (0.8031, 0, 2.7513, 0)$. The other parameters are

$$\begin{aligned} \epsilon &= [e^{(-3)} \quad e^{(-2)} \quad e^{(-1)} \quad e \quad e^{(1)}]^T \in \mathbb{R}^5, \\ K^b &= \begin{bmatrix} 0 & 0 & 0.5 & 1 & 2 \end{bmatrix} \in \mathbb{R}^{1 \times 5}, \\ K^f &= \begin{bmatrix} 0.1 & 0 & 1 & 2 & 0 \end{bmatrix} \in \mathbb{R}^{1 \times 5}, \\ Q &= 1, \\ \delta &= 0, \\ u_0(t) &= 0, \quad \forall t \in [0, t_f]. \end{aligned} \quad (56)$$

Note that we use up to three integrals, i.e., $\sigma = 3$.

Simulation 5. We show in Fig. 6 the states and control signals, in Fig. 7 the behavior of the Root Mean Square Error (RMSE) per trial for the output y , in Fig. 8 the pendulum behavior for the first and last iterations, respectively.

C. Discussion

Figs 2a–2b show trials progression of Sim. 1. We can see how the output converges to desired output but with an error. The remaining gap between the signals does not improve significantly from 30-th iteration forward, this behavior was predicted by Theorem 1, since the tracking error upper bound is not zero, as described by (9) and (32). The input signal displays oscillations since early 4-th iteration.

Figs 2c–2d show trials progression of Sim. 2. We can see how the system converges even if we choose a detrimental

initial guess (51) that affects input in every iteration. The final outcome does not improve significantly from 30-th iteration forward, and the tracking is worse than Sim. 1; steady state tracking error was predicted by (9) where we changed $|\Delta u_0|$. Note that we can also improve steady state tracking error with a wiser choice of $u_0(t)$, when $Q < 1$.

Fig.s 2e–2h show trials progression of Sim.s 3 and 4. In these cases the forgetting factor $Q = 1$ lets the system converge with negligible tracking error in only 30 iterations, getting closer to zero as predicted by (38). This holds true even in Sim. 4, where the choice of u_0 lowers the performance. In Sim. 3, the input at early 4-th trial displays an oscillatory behavior that is lost when the system reaches the 30-th iteration. Furthermore, Fig. 3 show internal state for the last iteration of Sim. 3. All the three masses have similar behaviors, since their movements are coupled by springs.

Fig. 4 summarizes the trial-by-trial performances of Sim.s 1–4. The performance index is the RMSE.

Fig.s 6a–6c show trials progression of Sim. 5. We can see how $\psi(1, t)$ tracks y_d (55) with negligible error starting at least from the 10-th iteration (Fig. 7), moving the tip in a vertical position after 5 seconds. Zero error behavior is expected since forgetting factor is unitary ($Q = 1$). Finally, Fig. 8 depicts the motion of the pendulum during 1-st and 30-th iteration. Fig. 8a shows that the feedback alone is not capable to complete the swing up movement in only 5 seconds.

VI. CONCLUSIONS

This paper introduced an online-offline ILC scheme of PI^σ - PI^σ -type. This learning rule can include an arbitrary number σ of integrators on feedback action and offline update. Its convergence is proven for a large class of nonlinear systems - including arbitrary relative degree $r \geq 1$ (e.g. mechanical systems relying on position output), delayed feedback loop, non-unitary forgetting factor. Future work will be devoted to experimental validation of the theory, with specific attention to the soft robotic application [28].

REFERENCES

- [1] G. Casalino and G. Bartolini, "A learning procedure for the control of movements of robotic manipulators," in *IASTED symposium on robotics and automation*, pp. 108–111, 1984.
- [2] S. Arimoto, S. Kawamura, and F. Miyazaki, "Bettering operation of dynamic systems by learning: A new control theory for servomechanism or mechatronics systems," in *The 23rd IEEE Conference on Decision and Control*, pp. 1064–1069, IEEE, 1984.
- [3] D. A. Bristow and J. R. Singler, "Analysis of transient growth in iterative learning control using pseudospectra," 2009.
- [4] D. Shen, "Iterative learning control with incomplete information: A survey," *IEEE/CAA Journal of Automatica Sinica*, vol. 5, no. 5, pp. 885–901, 2018.
- [5] H.-S. Lee and Z. Bien, "A note on convergence property of iterative learning controller with respect to sup norm," *Automatica*, vol. 33, no. 8, pp. 1591–1593, 1997.
- [6] T. Hashikawa and Y. Fujisaki, "Convergence conditions of iterative learning control revisited: A unified viewpoint to continuous-time and discrete-time cases," in *2013 IEEE International Symposium on Intelligent Control (ISIC)*, pp. 31–34, IEEE, 2013.
- [7] H.-S. Ahn, C.-H. Choi, and K.-b. Kim, "Iterative learning control for a class of nonlinear systems," *Automatica*, vol. 29, no. 6, pp. 1575–1578, 1993.
- [8] A. Isidori, "Elementary theory of nonlinear feedback for single-input single-output systems," in *Nonlinear Control Systems*, pp. 137–217, Springer, 1995.
- [9] M. Sun and D. Wang, "Anticipatory iterative learning control for nonlinear systems with arbitrary relative degree," *IEEE Transactions on Automatic Control*, vol. 46, no. 5, pp. 783–788, 2001.
- [10] M. Sun, D. Wang, and Y. Wang, "Sampled-data iterative learning control with well-defined relative degree," *International Journal of Robust and Nonlinear Control: IFAC-Affiliated Journal*, vol. 14, no. 8, pp. 719–739, 2004.
- [11] M. Pierallini, F. Angelini, R. Mengacci, A. Palleschi, A. Bicchi, and M. Garabini, "Trajectory tracking of a one-link flexible arm via iterative learning control," in *2020 IEEE/RSJ International Conference on Intelligent Robots and Systems, IROS 2020*, Institute of Electrical and Electronics Engineers Inc., 2020.
- [12] C. Tan, S. Wang, and J. Wang, "Robust iterative learning control for iteration-and time-varying disturbance rejection," *International Journal of Systems Science*, vol. 51, no. 3, pp. 461–472, 2020.
- [13] Y. Yu, J. Wan, and H. Bi, "Suboptimal learning control for nonlinearly parametric time-delay systems under alignment condition," *IEEE Access*, vol. 6, pp. 2922–2929, 2017.
- [14] B. E. Helfrich, C. Lee, D. A. Bristow, X. Xiao, J. Dong, A. G. Alleyne, S. M. Salapaka, and P. M. Ferreira, "Combined h_∞ -feedback control and iterative learning control design with application to nanopositioning systems," *IEEE Transactions on Control Systems Technology*, vol. 18, no. 2, pp. 336–351, 2010.
- [15] Z.-B. Wei, Q. Quan, and K.-Y. Cai, "Output feedback ilc for a class of nonminimum phase nonlinear systems with input saturation: An additive-state-decomposition-based method," *IEEE Transactions on Automatic Control*, vol. 62, no. 1, pp. 502–508, 2016.
- [16] K. Pereida, D. Kooijman, R. R. Duivenvoorden, and A. P. Schoellig, "Transfer learning for high-precision trajectory tracking through adaptive feedback and iterative learning," *International Journal of Adaptive Control and Signal Processing*, vol. 33, no. 2, pp. 388–409, 2019.
- [17] J.-X. Xu, T. H. Lee, and H.-W. Zhang, "Analysis and comparison of iterative learning control schemes," *Engineering Applications of Artificial Intelligence*, vol. 17, no. 6, pp. 675–686, 2004.
- [18] P. Ouyang, "Pd-pd type learning control for uncertain nonlinear systems," in *International Design Engineering Technical Conferences and Computers and Information in Engineering Conference*, vol. 49019, pp. 699–707, 2009.
- [19] Z. Feng, Z. Zhang, and D. Pi, "Open-closed-loop pd-type iterative learning controller for nonlinear systems and its convergence," in *Fifth World Congress on Intelligent Control and Automation (IEEE Cat. No. 04EX788)*, vol. 2, pp. 1241–1245, IEEE, 2004.
- [20] J. Shou, D. Pi, and W. Wang, "Sufficient conditions for the convergence of open-closed-loop pid-type iterative learning control for nonlinear time-varying systems," in *SMC'03 Conference Proceedings. 2003 IEEE International Conference on Systems, Man and Cybernetics. Conference Theme-System Security and Assurance (Cat. No. 03CH37483)*, vol. 3, pp. 2557–2562, IEEE, 2003.
- [21] J. Dong, B. He, C. Zhang, and G. Li, "Open-closed-loop pd iterative learning control with a variable forgetting factor for a two-wheeled self-balancing mobile robot," *Complexity*, vol. 2019, 2019.
- [22] X. Shan-hai, Z. Zhong, and Z. Xin, "Pd-type open-closed-loop iterative learning control in the networked control system," in *2016 Chinese Control and Decision Conference (CCDC)*, pp. 5738–5744, IEEE, 2016.
- [23] S. S. Dragomir, "Some gronwall type inequalities and applications," URL: <http://rgmia.vu.edu.au/SSDragomirWeb.html>, 2002.
- [24] F. Angelini, C. Della Santina, M. Garabini, M. Bianchi, G. M. Gasparri, G. Grioli, M. G. Catalano, and A. Bicchi, "Decentralized trajectory tracking control for soft robots interacting with the environment," *IEEE Transactions on Robotics*, 2018.
- [25] L. Cenceschi, C. Della Santina, G. Averta, M. Garabini, Q. Fu, M. Santello, M. Bianchi, and A. Bicchi, "Modeling previous trial effect in human manipulation through iterative learning control," *Advanced Intelligent Systems*, vol. 2, no. 9, p. 1900074, 2020.
- [26] C. Della Santina, *Flexible Manipulators*. Berlin, Heidelberg: Springer Berlin Heidelberg, 2021.
- [27] C. Della Santina, "The soft inverted pendulum with affine curvature," in *2020 59th IEEE Conference on Decision and Control (CDC)*, pp. 4135–4142, IEEE, 2020.
- [28] C. Della Santina, M. G. Catalano, and A. Bicchi, *Soft Robots*, pp. 1–15. Berlin, Heidelberg: Springer Berlin Heidelberg, 2020.