

Passivity-based Iterative Learning Control for Visual Feedback System

Toshiyuki Murao, Hiroyuki Kawai and Masayuki Fujita

Abstract—This paper investigates iterative learning control based on passivity for three-dimensional (3-D) visual feedback systems. Firstly, a brief summary of a visual motion observer is given. Next, a pose control error system for iterative learning control that has an output strictly passivity property is constructed. Then, iterative learning control for 3-D visual feedback systems is proposed. The transient response of the proposed control law should be improved because of the repeatability. Convergence analysis of the closed-loop system is discussed based on passivity. Finally, simulation results are shown in order to confirm the proposed method.

I. INTRODUCTION

Visual feedback control is now a very flexible and useful method in robot control [1]. Recently, Lippiello *et al.* [2] presented position-based visual servoing for a hybrid eye-in-hand/eye-to-hand multicamera system by using an extended Kalman filter and a multiarm robotic cell. Gans and Hutchinson [3] proposed hybrid switched-system control which incorporates both an image-based visual servoing controller and a position-based one. Hu *et al.* [4] proposed quaternion-based robust visual servo control in presence of camera calibration error. Swensen *et al.* [5] presented featureless kernel-based visual servoing and discussed a domain of attraction of the method through experiments. Allibert *et al.* [6] reported comparison results between two image prediction models for an image-based visual servoing scheme based on nonlinear model predictive control. The authors have been proposed passivity-based visual feedback control for three-dimensional (3-D) target tracking in a series of papers [7]–[10].

On the other hand, iterative learning control has become an attractive control method for improving the transient response and tracking performance of uncertain dynamic systems that operate repetitively [11]. Examples include not only robotics, chemical processes and biomedical applications, but also surgical assistants recently [12]. Nowadays, there are several papers related to the iterative learning control for visual feedback systems. Mansard *et al.* [13] proposed a jacobian learning method for coarsely calibrated visual feedback systems. In particular, iterative learning control approaches for visual feedback systems which address stability analysis were investigated in a number of papers [14]–[16]. In [14] and [15], a iterative learning scheme for nonlinear

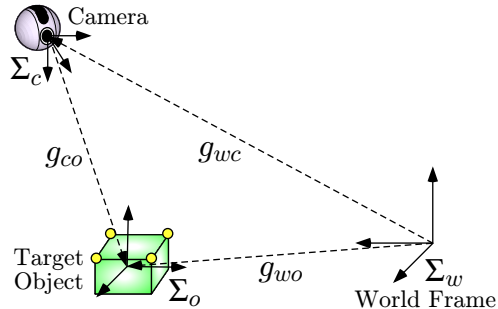


Fig. 1. Eye-in-hand visual feedback system.

visual feedback control with initial resetting error and with an unknown homography matrix were presented, respectively. Jiang *et al.* [16] reported indirect iterative learning control using neural networks without a vision camera model. Although good iterative learning control methods are reported in those papers, each control approach has the assumption that a depth along optical axis can be measured by a robot position sensor [14], a robot motion plane is parallel to an image plane [15] and selected feature points are restricted to only three points, in other words, redundant image features (over four feature points) cannot be applied [16].

In this paper, we propose iterative learning control based on passivity for 3-D visual feedback systems with an eye-in-hand configuration as shown in Fig. 1. Firstly, a brief summary of a visual motion observer in order to estimate relative rigid body motion is given. Secondly, a pose control error system for iterative learning control that has an output strictly passivity property is constructed. Next, we propose iterative learning control to ensure that an actual relative rigid body motion from a camera to a target object tracks the desired one. Our control law is based on the Arimoto-type iterative learning controller [17] that is much efficient for motion control of robot systems. Convergence analysis of the closed-loop system is discussed based on passivity. Compared with our previous work [7], the transient response can be improved through repetition to track desired trajectories. Finally, control performance of the proposed control scheme is evaluated through simulation results. In particular, the simulation under a realistic system model is performed from the more practical point of view.

II. VISUAL MOTION OBSERVER

This section mainly reviews our previous work [7] via passivity-based visual feedback control for eye-in-hand systems.

This work was not supported by any organization

T. Murao is with Master Program of Innovation for Design and Engineering Advanced Institute of Industrial Technology, Tokyo 140-0011, Japan murao-toshiyuki@aait.ac.jp

H. Kawai is with Department of Robotics, Kanazawa Institute of Technology, Ishikawa 921-8501, Japan

M. Fujita is with Department of Mechanical and Control Engineering, Tokyo Institute of Technology, Tokyo 152-8550, Japan

A. Vision Camera Model

Visual feedback systems with an eye-in-hand configuration use three coordinate frames which consist of a world frame Σ_w , a camera frame Σ_c , and an object frame Σ_o as in Fig. 1. Let $p_{co} \in \mathcal{R}^3$ and $e^{\hat{\xi}\theta_{co}} \in SO(3)$ be a position vector and a rotation matrix from the camera frame Σ_c to the object frame Σ_o . Then, a relative rigid body motion from Σ_c to Σ_o can be represented by $g_{co} = (p_{co}, e^{\hat{\xi}\theta_{co}}) \in SE(3)$ ¹. Similarly, $g_{wc} = (p_{wc}, e^{\hat{\xi}\theta_{wc}})$ and $g_{wo} = (p_{wo}, e^{\hat{\xi}\theta_{wo}})$ denote rigid body motions from the world frame Σ_w to the camera frame Σ_c and from the world frame Σ_w to the object frame Σ_o , respectively.

Position-based visual feedback control needs, in general, to derive the relative rigid body motion from image data. The relative rigid body motion from Σ_c to Σ_o can be led by using the composition rule for rigid body transformations [18] as follows:

$$g_{co} = g_{wc}^{-1} g_{wo}. \quad (1)$$

The relative rigid body motion involves a velocity of each rigid body. We define a body velocity of the camera relative to the world frame Σ_w as $V_{wc}^b = [v_{wc}^T \ \omega_{wc}^T]^T$, where v_{wc} and ω_{wc} represent a velocity of an origin and an angular velocity from Σ_w to Σ_c , respectively [18]. Differentiating Eq. (1) with respect to time, the body velocity of the relative rigid body motion g_{co} can be written as follows (See [7]):

$$V_{co}^b = -\text{Ad}_{(g_{co}^{-1})} V_{wc}^b + V_{wo}^b, \quad (2)$$

where $\text{Ad}_{(g_{ab})}$ is the adjoint transformation associated with g_{ab} and V_{wo}^b is the body velocity of the target object relative to Σ_w .

In this visual feedback system, we use a pinhole camera model with a perspective projection. Here, we consider $m(\geq 3)$ feature points on the rigid target object in this paper. Let λ be a focal length, $p_{oi} \in \mathcal{R}^3$ and $p_{ci} \in \mathcal{R}^3$ be position vectors of the target object's i -th feature point relative to Σ_o and Σ_c , respectively. Using a transformation of the coordinates, we have $p_{ci} = g_{co} p_{oi}$, where p_{ci} and p_{oi} should be regarded, with a slight abuse of notation, as $[p_{ci}^T \ 1]^T$ and $[p_{oi}^T \ 1]^T$. We assume that multiple point features on a known object p_{oi} are given. The perspective projection of the i -th feature point onto the image plane gives us the image plane coordinate $f_i := [f_{xi} \ f_{yi}]^T \in \mathcal{R}^2$ as

$$f_i = \frac{\lambda}{z_{ci}} \begin{bmatrix} x_{ci} \\ y_{ci} \end{bmatrix}, \quad (3)$$

where $p_{ci} = [x_{ci} \ y_{ci} \ z_{ci}]^T$. It is straightforward to extend this model to m image points by simply stacking the vectors of the image plane coordinate, i.e., $f := [f_1^T \ \dots \ f_m^T]^T \in \mathcal{R}^{2m}$. The image feature f only depends on the relative rigid body motion g_{co} .

B. Visual Motion Observer

In this subsection, we consider a nonlinear observer (we call the visual motion observer) in order to estimate the relative rigid body motion g_{co} from the image feature f . Using the body velocity of the relative rigid body motion g_{co} (2), we choose estimates $\bar{g}_{co} = (\bar{p}_{co}, e^{\hat{\xi}\bar{\theta}_{co}})$ and $\bar{V}_{co}^b$ of the relative rigid body motion and velocity, respectively as

$$\bar{V}_{co}^b = -\text{Ad}_{(\bar{g}_{co}^{-1})} V_{wc}^b + u_e. \quad (4)$$

The new input u_e is to be determined in order to drive the estimated values $\bar{g}_{co}$ and $\bar{V}_{co}^b$ to their actual values.

In order to establish an estimation error system, we define an estimation error between the estimated value $\bar{g}_{co}$ and the actual relative rigid body motion g_{co} as

$$g_{ee} := \bar{g}_{co}^{-1} g_{co}. \quad (5)$$

Note that $p_{co} = \bar{p}_{co}$ and $e^{\hat{\xi}\theta_{co}} = e^{\hat{\xi}\bar{\theta}_{co}}$ if and only if $g_{ee} = I_4$, i.e., $p_{ee} = 0$ and $e^{\hat{\xi}\theta_{ee}} = I_3$. We next define an error vector of the rotation matrix $e^{\hat{\xi}\theta_{ei}}$ as $r_{ei} := \text{sk}(e^{\hat{\xi}\theta_{ei}})^\vee$ where $\text{sk}(e^{\hat{\xi}\theta_{ei}})$ denotes $\frac{1}{2}(e^{\hat{\xi}\theta_{ei}} - e^{-\hat{\xi}\theta_{ei}})$. Using this notation, the vector of the estimation error is given by

$$e_e := \begin{bmatrix} p_{ee}^T & r_{ee}^T \end{bmatrix}^T. \quad (6)$$

Note that $e_e = 0$ iff $p_{ee} = 0$ and $e^{\hat{\xi}\theta_{ee}} = I_3$. Suppose an attitude estimation error θ_{ee} is small enough that we can let $e^{\hat{\xi}\theta_{ee}} \simeq I + \text{sk}(e^{\hat{\xi}\theta_{ee}})$. Therefore, using a first-order Taylor expansion approximation, the estimation error vector e_e can be obtained from the image feature f and the estimated value of the relative rigid body motion $\bar{g}_{co}$ (i.e., the measurement and the estimate) as follows:

$$e_e = J^\dagger(\bar{g}_{co})(f - \bar{f}), \quad (7)$$

where $\bar{f}$ is an estimated value of the image feature and $J(\bar{g}_{co})$ is an image Jacobian-like matrix [7]. In the same way as Eq. (2), the estimation error system can be represented by

$$V_{ee}^b = -\text{Ad}_{(g_{ee}^{-1})} u_e + V_{wo}^b. \quad (8)$$

Then, we have the following lemma relating the input u_e to the vector form of the estimation error e_e .

Lemma 1 ([7]): If $V_{wo}^b = 0$, then the following inequality holds for the estimation error system (8).

$$\int_0^T u_e^T (-e_e) dt \geq -\beta_e, \quad (9)$$

where β_e is a positive scalar.

Based on the above passivity property of the estimation error system, we consider the following control law:

$$u_e = -K_e(-e_e), \quad (10)$$

where $K_e := \text{diag}\{k_{e1}, \dots, k_{e6}\}$ is the positive gain matrix of x , y and z axes of the translation and the rotation for the estimation error.

Theorem 1 ([7]): If $V_{wo}^b = 0$, then the equilibrium point $e_e = 0$ for the closed-loop system (8) and (10) is asymptotic stable.

¹The notation of the homogeneous transform is denoted in Appendix.

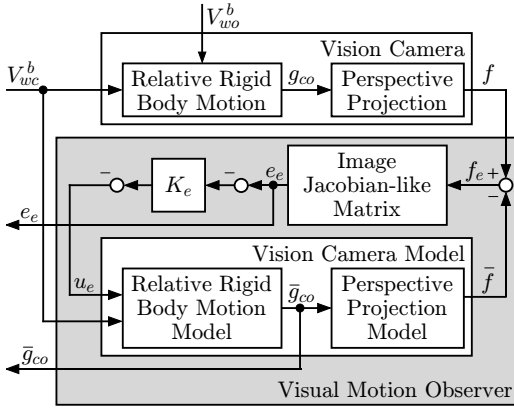


Fig. 2. Block diagram of visual motion observer.

It should be noted that if the vector of the estimation error is equal to zero, then the estimated relative rigid body motion $\bar{g}_{co}$ equals the actual one g_{co} . Fig. 2 shows a block diagram of the visual motion observer. The estimation error vector is configured by available information (i.e., the measurement and the estimate) though it is defined by unavailable one. Thanks to the proposed visual motion observer, the unmeasurable motion g_{co} can be exploited as the part of control law. Our proposed visual motion observer is composed just as Luenberger observer for linear systems.

III. ITERATIVE LEARNING CONTROL FOR PASSIVITY-BASED VISUAL FEEDBACK SYSTEM

Iterative learning control is an approach to improving the transient response performance of the system that operates repetitively [11]. In this section, an iterative learning control law for the visual feedback system is proposed based on a passivity property.

A. Pose Control Error System for Iterative Learning Control

Let us consider the dual of the estimation error system, which we call the pose control error system, in order to achieve the control objective. The control objective is to bring the actual relative rigid body motions g_{co} to a given reference one g_{cd} . First, we define the pose control error as follows:

$$g_{ec} = g_{cd}^{-1} g_{co}, \quad (11)$$

which represents the error between the relative rigid body motion g_{co} and the reference one g_{cd} . It should be remarked that g_{co} can be calculated by using the estimated relative rigid body motion $\bar{g}_{co}$ and the estimation error vector $e_e = [p_{ee}^T \ r_{ee}^T]^T$ equivalently as follows:

$$g_{co} = \bar{g}_{co} g_{ee} \quad (12)$$

$$\xi \theta_{ee} = \frac{\sin^{-1} \|r_{ee}\|}{\|r_{ee}\|} r_{ee}, \quad (13)$$

although g_{co} cannot be measured directly (see [8] for more details). Using the notation $r_{ei} = \text{sk}(e^{\xi \theta_{ei}})^\vee$, the vector of the pose control error is defined as

$$e_c := [p_{ec}^T \ r_{ec}^T]^T. \quad (14)$$

Next, we consider that the input of the vision camera V_{wc}^b has an ideal one u_i which can achieve to bring the actual relative rigid body motion g_{co} to the reference one g_{cd} . Here, we assume that the initial relative rigid body motion and the body velocity from the camera to the target object are equal to the desired ones, i.e., $g_{co}(0) = g_{cd}(0)$ and $V_{co}^b(0) = V_{cd}^b(0)$. Substituting g_{co} and V_{co}^b with g_{cd} and V_{cd}^b respectively in Eq. (2), the ideal input u_i can be represented as follows:

$$u_i = -\text{Ad}_{(g_{cd})} (V_{cd}^b - V_{wo}^b). \quad (15)$$

It should be remarked that the ideal input u_i cannot actually be implemented for the visual feedback system since the body velocity of the target object V_{wo}^b is unknown.

We now propose the camera velocity for the visual feedback system as

$$V_{wc}^b = \text{Ad}_{(g_{cd})} K_c e_c + u_l, \quad (16)$$

where $K_c := \text{diag}\{k_{c1}, \dots, k_{c6}\}$ is the positive gain matrix of x , y and z axes of the translation and the rotation for the pose control error. The new input u_l is to be determined later by an iterative learning control scheme. Substituting Eq. (16) into Eq. (2), the body velocity of the relative rigid body motion g_{co} can be obtained as follows:

$$V_{co}^b = -\text{Ad}_{(g_{co}^{-1})} (\text{Ad}_{(g_{cd})} K_c e_c + u_l) + V_{wo}^b. \quad (17)$$

It should be noted that $u_l \rightarrow u_i$ from Eq. (17) when $g_{co} \rightarrow g_{cd}$, i.e., $e_c \rightarrow 0$.

Here, we define the ideal input error as follows:

$$\begin{aligned} u_{ei} &= \text{Ad}_{(g_{cd}^{-1})} (V_{wc}^b - u_i) - K_c e_c \\ &= \text{Ad}_{(g_{cd}^{-1})} (u_l - u_i). \end{aligned} \quad (18)$$

It is obvious that the ideal input error u_{ei} can be regarded as the exact error between the actual input for the vision camera V_{wc}^b and the ideal one u_i in the case of $e_c = 0$. Differentiating Eq. (11) with respect to time, the pose control error system for the iterative learning control can be represented as

$$\begin{aligned} V_{ec}^b &= -\text{Ad}_{(g_{ec}^{-1})} V_{cd}^b + V_{co}^b \\ &= -\text{Ad}_{(g_{ec}^{-1})} (u_{ei} + K_c e_c) + (I - \text{Ad}_{(g_{ec}^{-1})}) V_{wo}^b. \end{aligned} \quad (19)$$

Moreover, we consider the following output for the input u_l using the error vector of the pose control error:

$$\nu_l = -e_c. \quad (20)$$

In the case of the ideal input u_i that can drive g_{co} to g_{cd} , it is obvious that its output is equal to zero, i.e. $\nu_i = 0$. Similar to the ideal input error u_{ei} , the following ideal output error between the actual output ν_l and the ideal one ν_i is obtained as

$$\nu_{ei} = \nu_l - \nu_i = \nu_l. \quad (21)$$

Next, we show an important relation between the input u_{ei} and the output ν_{ei} of the pose control error system for the iterative learning control.

Lemma 2: If $V_{wo}^b = 0$, then the pose control error system (19) satisfies

$$\int_0^T u_{ei}^T \nu_{ei} dt \geq -\beta + \int_0^T \nu_{ei}^T K_c \nu_{ei} dt, \quad \forall T > 0 \quad (22)$$

where β is a positive scalar.

Proof: Consider the following positive definite function

$$V = \frac{1}{2} \|p_{ec}\|^2 + \phi(e^{\hat{\theta}_{ec}}), \quad (23)$$

where $\phi(e^{\hat{\theta}_{ec}}) := \frac{1}{2} \text{tr}(I - e^{\hat{\theta}_{ec}})$ is an error function of the rotation matrix (see e.g. [19]). The positive definiteness of the function V results from the property of the error function ϕ . Evaluating the time derivative of V along the trajectories of Eq. (19) gives us

$$\begin{aligned} \dot{V} &= p_{ec}^T e^{\hat{\theta}_{ec}} e^{-\hat{\theta}_{ec}} \dot{p}_{ec} + e_R^T(e^{\hat{\theta}_{ec}}) e^{\hat{\theta}_{ec}} \omega_{ec} \\ &= e_c^T \text{Ad}_{(e^{\hat{\theta}_{ec}})} \left(-\text{Ad}_{(g_{ec}^{-1})}(u_{ei} + K_c e_c) \right) \\ &= -e_c^T \text{Ad}_{(-p_{ec})}(u_{ei} + K_c e_c). \end{aligned} \quad (24)$$

From the property of “ $\wedge$ ” (wedge), we have $p_{ec}^T \hat{p}_{ec} \omega_{uc} = -p_{ec}^T \hat{\omega}_{uc} p_{ec} = 0$, where $u_{ei} + K_c e_c = [v_{uc}^T \omega_{uc}^T]^T$. Using this fact, we can obtain

$$\begin{aligned} \dot{V} &= -e_c^T (u_{ei} + K_c e_c) \\ &= u_{ei}^T \nu_{ei} - \nu_{ei}^T K_c \nu_{ei}. \end{aligned} \quad (25)$$

Integrating Eq. (25) from 0 to T yields

$$\begin{aligned} \int_0^T u_{ei}^T \nu_{ei} dt &= V(T) - V(0) + \int_0^T \nu_{ei}^T K_c \nu_{ei} dt \\ &\geq -V(0) + \int_0^T \nu_{ei}^T K_c \nu_{ei} dt \\ &:= -\beta + \int_0^T \nu_{ei}^T K_c \nu_{ei} dt, \end{aligned} \quad (26)$$

where β is the positive scalar which only depends on the initial state of $g_{ec} = (p_{ec}, e^{\hat{\theta}_{ec}})$. ■

Lemma 2 implies that the pose control error system (19) is *output strictly passive* from the input u_{ei} to the output ν_{ei} as in the definition in [20]. It should be noted that the pose control error system (19) is different from that proposed in [7], in both points that the proposed camera velocity has already been integrated into it and that the input of the pose control error system is the error between the actual camera velocity and the ideal one. One of the main contributions of this paper is that we construct the pose control error system for the iterative learning control which has an output strictly passivity.

B. Iterative Learning Control for Passivity-based Visual Feedback System

In this subsection, we present an iterative learning control law for the visual feedback system. The problem of iterative learning control is, in general, to find a recursive form of a learning control law $u_l^{k+1} = F(u_l^k, \nu_{ei}^k)$ in trial number k that eventually realizes the convergence $\nu_{ei} \rightarrow 0$ as $k \rightarrow \infty$ [17]. In this paper, we tackle the same problem as mentioned in [17] for the visual feedback system. We now propose the learning control update law for the visual feedback system as follows:

$$u_l^{k+1} = u_l^k - \text{Ad}_{(g_{cd})} K_l \nu_{ei}^k, \quad (27)$$

where $K_l := \text{diag}\{k_{l1}, \dots, k_{l6}\}$ is the positive gain matrix of x , y and z axes of the translation and the rotation for the

iterative learning. From Eqs. (18), (21) and (27), it can be easily derived the following relationship

$$u_{ei}^{k+1} = u_{ei}^k - K_l \nu_{ei}^k. \quad (28)$$

Suppose that the target object is static, the following theorem concerning the convergence of the iterative learning control for the visual feedback system holds.

Theorem 2: Suppose that $V_{wo}^b = 0$ and $0 < K_l < 2K_c$, then the iterative learning control law (27) for the pose control error system (19) guarantees the convergence of $e_c^k = 0$ in the sense of a $L_2[0, T]$ norm.

Proof: For Eq. (28), multiplying both sides by the positive definite symmetric matrix K_l^{-1} can be transformed into

$$K_l^{-1} u_{ei}^{k+1} = K_l^{-1} u_{ei}^k - \nu_{ei}^k. \quad (29)$$

Inner products of both sides of Eqs. (28) and (29) can be obtained

$$\begin{aligned} (u_{ei}^{k+1})^T K_l^{-1} u_{ei}^{k+1} &= (u_{ei}^k)^T K_l^{-1} u_{ei}^k + (\nu_{ei}^k)^T K_l \nu_{ei}^k - 2(u_{ei}^k)^T \nu_{ei}^k. \end{aligned} \quad (30)$$

Integrating of both sides Eq. (30) from 0 to T yields

$$\begin{aligned} \int_0^T (u_{ei}^{k+1})^T K_l^{-1} u_{ei}^{k+1} dt &= \int_0^T (u_{ei}^k)^T K_l^{-1} u_{ei}^k dt + \int_0^T (\nu_{ei}^k)^T K_l \nu_{ei}^k dt \\ &\quad - 2 \int_0^T (u_{ei}^k)^T \nu_{ei}^k dt. \end{aligned} \quad (31)$$

This equality can be represented in terms of a L_2 norm as follows:

$$\|u_{ei}^{k+1}\|_{K_l^{-1}}^2 = \|u_{ei}^k\|_{K_l^{-1}}^2 + \|\nu_{ei}^k\|_{K_l}^2 - 2 \int_0^T (u_{ei}^k)^T \nu_{ei}^k dt, \quad (32)$$

where $\|\cdot\|_{\{\cdot\}}$ is defined as

$$\|u_{ei}^k\|_{K_l^{-1}}^2 := \int_0^T (u_{ei}^k)^T K_l^{-1} u_{ei}^k dt. \quad (33)$$

In the proof of Lemma 2, we have already shown the following relationship:

$$\int_0^T (u_{ei}^k)^T \nu_{ei}^k dt = V^{k+1} - V^k + \|\nu_{ei}^k\|_{K_c}^2, \quad (34)$$

where we have regarded $V(T)$ and $V(0)$ as V^{k+1} and V^k , respectively. From Eqs. (32) and (34), it can be easily shown that

$$\|u_{ei}^{k+1}\|_{K_l^{-1}}^2 + 2V^{k+1} = \|u_{ei}^k\|_{K_l^{-1}}^2 + 2V^k - \|\nu_{ei}^k\|_{(2K_c - K_l)}^2. \quad (35)$$

From Eq. (35), the fact that the function $\{\|u_{ei}^k\|_{K_l^{-1}}^2 + 2V^k\}$ is a monotonically non-increasing function and bounded below, and the assumption $0 < K_l < 2K_c$, it is clear that it converges to a non-negative value when $k \rightarrow \infty$. Then, the output error $\|\nu_{ei}^k\|_{(2K_c - K_l)}^2$ tends to zero as $k \rightarrow \infty$, i.e., ν_{ei}^k converges to zero in the sense of a $L_2[0, T]$ norm. Since

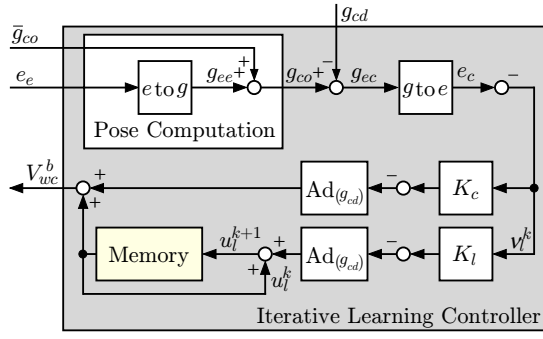


Fig. 3. Block diagram of iterative learning control.

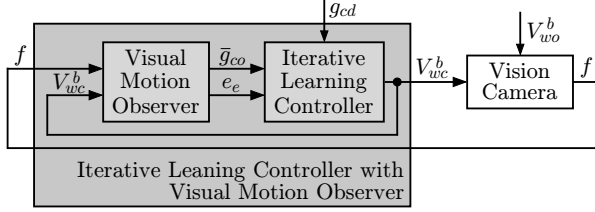


Fig. 4. Block diagram of iterative learning control with visual motion observer.

$\nu_l^k = -e_c^k$, it can be concluded that the pose control error e_c in trial number k converges to zero. ■

Theorem 2 guarantees the convergence of the relative rigid body motion g_{co} in trial number k to the desired one g_{cd} using the property that the pose control error system (19) is output strictly passive. From Theorem 2, it can be also shown that the input error $\|u_{ei}^k\|_{K_l^{-1}}^2$ tends to zero as $k \rightarrow \infty$.

Since $e_c^k = 0$ as $k \rightarrow \infty$ in Eq. (18), the body velocity of the vision camera V_{wc}^b converges to the ideal input u_i in the sense of a $L_2[0, T]$ norm. Fig. 3 shows a block diagram of the iterative learning controller. It should be noted that the input for the vision camera V_{wc}^b does not need the ideal input u_i in practice as shown in Fig. 3, though it is assumed its existence using the unknown target object velocity V_{wo}^b theoretically. Fig. 4 shows a block diagram of the closed-loop system which consists of the vision camera and the passivity-based iterative learning controller with visual motion observer.

Since the iterative learning control approach is generally to improve the transient response and tracking performance for the repeatability, the control performance of the proposed control law for the 3-D visual feedback system should be improved compared to the previous one in [7]. This is one of the main advantages of this work. In addition, it is a fact that the image Jacobian-like matrix $J(\cdot)$ of the pinhole camera model is exactly the same form as that of the panoramic camera model [9]. Therefore, the proposed iterative learning control approach can be applied to visual feedback systems with a panoramic camera using the image Jacobian-like matrix $J(\cdot)$ in [9]. This allows us to extend the technological application area.

IV. SIMULATION RESULTS

In this section, we present simulation results for the iterative learning control with a static target object. The desired relative rigid body motion g_{cd} is generated by using

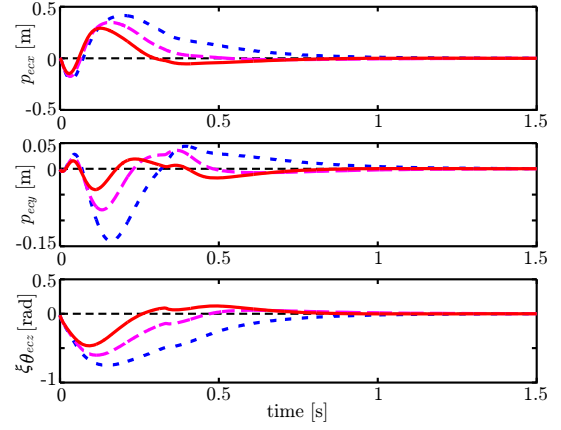


Fig. 5. Pose control error e_c : dotted; with $k = 1$: dashed; with $k = 5$: solid; with $k = 10$.

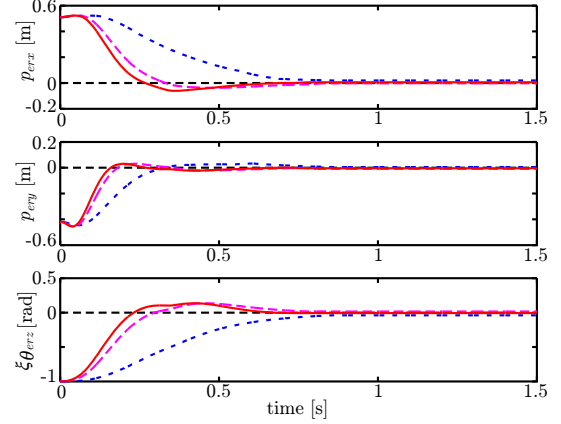


Fig. 6. Pose control error e_r : dotted; with $k = 1$: dashed; with $k = 10$: solid; with $k = 15$.

the image space navigation function-based path planner [10]. The control objective is that the vision camera tracks the static target object to keep all target feature points (four points) inside the camera field of view. In other words, it is bring the actual relative rigid body motion g_{co} to the given time-varying reference one g_{cd} , i.e., it can be achieved to make the pose control error zero.

The simulation is carried out with the initial condition $p_{co} = [0.2 \ 0.2 \ -1.35]^T$ m, $\xi\theta_{co} = [0 \ 0 \ 0]^T$ rad. The final desired relative rigid body motion is $p_{cd_f} = [-0.23 \ -0.3 \ -1.35]^T$ m, $\xi\theta_{cd_f} = [0 \ 0 \ \pi/2]^T$ rad. The gains were selected as $K_e = 5I$, $K_c = 5I$ and $K_l = I$ satisfy $0 < K_l < 2K_c$.

The simulation result is presented in Fig. 5. Fig. 5 shows the pose control error e_c . The dotted, dashed and solid lines denote the errors applying the update control law in trial number $k = 1$, $k = 5$ and $k = 10$, respectively. We focus on the errors of the translations of x and y and the rotation of z . In Fig. 5, the asymptotic stability can be confirmed by steady state performance. Moreover, the rise time applying the iterative learning control is shorter with increasing repetition. Compared to [7] that is equivalent to the iterative learning control approach in the case of $k = 1$, it can be easily verified that the control performance is improved.

Next, we consider a more realistic situation which has

TABLE I
NORM OF THE POSE CONTROL ERROR e_r AT 1.5 S.

Trial Number	Norm
$k = 1$	0.0482
$k = 10$	0.0245
$k = 15$	0.0096

the influence of the friction for the vision camera motion. In particular, if the elements of the camera body velocity input is calculated as $|v_{wci}| < 0.1$ and $|\omega_{wci}| < 0.1$, where $i = x, y, z$, then we regards it as zero, i.e., the friction makes the vision camera immovable in the case of $|v_{wci}| < 0.1$ and $|\omega_{wci}| < 0.1$. Fig. 6 shows the actual control error e_r , which is the error vector between the current relative rigid body motion g_{co} and the final desired one g_{cd_f} , instead of the time-varying desired one g_{cd} . The norm of the pose control error e_r at 1.5s is also shown in Table I. From Fig. 6 and Table I, the slight error remains at the steady state because of the friction force. However, it can be confirmed that it decreases with increasing repetition. This simulation results indicates that the steady state error in trial number $k - 1$ influences directly generation of the control input in the next trial. Since the control input increases, the steady state error is reduced, consequently. This is the great merit of the iterative learning control for the more realistic visual feedback system.

V. CONCLUSIONS

This paper proposes iterative learning control based on passivity for 3-D visual feedback systems. The main contribution of this paper is to show that the iterative learning control law which can improve the transient response and tracking performance for the repeatability is designed for the visual feedback system. Convergence analysis of the closed-loop system is discussed based on passivity. Simulation results are presented to verify the control performance of the proposed control scheme and the previous one [7].

REFERENCES

- [1] F. Chaumette and S. A. Hutchinson, "Visual Servoing and Visual Tracking," In: B. Siciliano and O. Khatib (Eds.), *Springer Handbook of Robotics*, Springer-Verlag, pp. 563–583, 2008.
- [2] V. Lippiello, B. Siciliano and L. Villani, "Position-based Visual Servoing in Industrial Multirobot Cells using a Hybrid Camera Configuration," *IEEE Trans. on Robotics*, Vol. 23, No. 1, pp. 73–86, 2007.
- [3] N. R. Gans and S. A. Hutchinson, "Stable Visual Servoing through Hybrid Switched-System Control," *IEEE Trans. on Robotics*, Vol. 23, No. 3, pp. 530–540, 2007.
- [4] G. Hu, N. Gans and W. E. Dixon, "Quaternion-based Visual Servo Control in the Presence of Camera Calibration Error," *International Journal of Robust and Nonlinear Control*, Vol. 20, No. 5, pp. 489–503, 2010.
- [5] J. P. Swensen, V. Kallem and N. J. Cowan, "Empirical Characterization of Convergence Properties for Kernel-based Visual Servoing," In: G. Chesi and K. Hashimoto (Eds.), *Visual Servoing via Advanced Numerical Methods*, Springer-Verlag, pp. 22–38, 2010.
- [6] G. Allibert, E. Courtial and F. Chaumette, "Predictive Control for Constrained Image-based Visual Servoing," *IEEE Trans. on Robotics*, Vol. 26, No. 5, pp. 933–939, 2010.
- [7] M. Fujita, H. Kawai and M. W. Spong, "Passivity-based Dynamic Visual Feedback Control for Three Dimensional Target Tracking: Stability and L_2 -gain Performance Analysis," *IEEE Trans. on Control Systems Technology*, Vol. 15, No. 1, pp. 40–52, 2007.

- [8] T. Murao, H. Kawai and M. Fujita, "Predictive Visual Feedback Control with Eye-in/to-Hand Configuration via Stabilizing Receding Horizon Approach," *Proc. of the 17th IFAC World Congress on Automatic Control*, pp. 5341–5346, 2008.
- [9] H. Kawai, T. Murao and M. Fujita, "Visual Motion Observer-based Pose Control with Panoramic Camera via Passivity Approach," *Proc. of the 2010 American Control Conference*, pp. 4534–4539, 2010.
- [10] T. Murao, H. Kawai and M. Fujita, "Visual Motion Observer-based Stabilizing Receding Horizon Control via Image Space Navigation Function," *Proc. of the 2010 IEEE Multi-conference on Systems and Control*, pp. 1648–1653, 2010.
- [11] H.-S. Ahn, Y. Q. Chen and K. L. Moore, "Iterative Learning Control: Brief Survey and Categorization," *IEEE Trans. on Systems, Man, and Cybernetics-Part C: Applications and Reviews*, Vol. 37, No. 6, pp. 1099–1121, 2007.
- [12] J. Berg, S. Miller, D. Duckworth, H. Hu, A. Wan, X.-Y. Fu, K. Goldberg and P. Abbeel, "Superhuman Performance of Surgical Tasks by Robots using Iterative Learning from Human-Guided Demonstrations," *Proc. of the 2010 IEEE International Conference on Robotics and Automation*, pp. 2074–2081, 2010.
- [13] N. Mansard, M. Lopes, J. Santos-Victor and F. Chaumette, "Jacobian Learning Methods for Tasks Sequencing in Visual Servoing," *Proc. of the 2006 IEEE/RSJ International Conference on Intelligent Robots and Systems*, pp. 4284–4290, 2006.
- [14] P. Jiang and R. Unbehauen, "Robot Visual Servoing with Iterative Learning Control," *IEEE Trans. on Systems, Man, and Cybernetics-Part A: Systems and Humans*, Vol. 32, No. 2, pp. 281–287, 2002.
- [15] D. Yi, J. Wu and P. Jiang, "Iterative Learning Control for Visual Servoing with Unknown Homography Matrix," *Proc. of the 2007 IEEE International Conference on Control and Automation*, pp. 2791–2796, 2007.
- [16] P. Jiang, L. C. A. Bamforth, Z. Feng and J. E. F. Baruch, "Indirect Iterative Learning Control for a Discrete Visual Servo without a Camera-Robot Model," *IEEE Trans. on Systems, Man, and Cybernetics-Part B: Cybernetics*, Vol. 37, No. 4, pp. 863–876, 2007.
- [17] S. Arimoto and T. Naniwa, "Learnability and Adaptability from the Viewpoint of Passivity Analysis," *Intelligent Automation and Soft Computing*, Vol. 8, No. 2, pp. 71–94, 2002.
- [18] R. Murray, Z. Li and S. S. Sastry, *A Mathematical Introduction to Robotic Manipulation*, CRC Press, 1994.
- [19] F. Bullo and A. D. Lewis, *Geometric Control of Mechanical Systems*, Springer-Verlag, 2004.
- [20] B. Brogliato, R. Lozano, B. Maschke and O. Egeiland, *Dissipative Systems Analysis and Control: Theory and Applications* (2nd ed.), Springer-Verlag, 2007.

APPENDIX

In this paper, we use the notation $e^{\hat{\xi}\theta_{ab}} \in \mathcal{R}^{3 \times 3}$ to represent the change of the principle axes of a frame Σ_b relative to a frame Σ_a . $\xi_{ab} \in \mathcal{R}^3$ specifies the direction of rotation and $\theta_{ab} \in \mathcal{R}$ is the angle of rotation. For simplicity we use $\hat{\xi}\theta_{ab}$ to denote $\xi_{ab}\theta_{ab}$. The notation ' $\wedge$ ' (wedge) is the skew-symmetric operator such that $\hat{\xi}\theta_{ab} = \xi_{ab} \times \theta_{ab}$ for the vector cross-product $\times$ and any vector $\theta \in \mathcal{R}^3$. The notation ' $\vee$ ' (vee) denotes the inverse operator to ' $\wedge$ ', i.e., $so(3) \rightarrow \mathcal{R}^3$. Recall that a skew-symmetric matrix corresponds to an axis of rotation (via the mapping $a \mapsto \hat{a}$). We use the 4×4 matrix

$$g_{ab} = \begin{bmatrix} e^{\hat{\xi}\theta_{ab}} & p_{ab} \\ 0 & 1 \end{bmatrix} \quad (36)$$

as the homogeneous representation of $g_{ab} = (p_{ab}, e^{\hat{\xi}\theta_{ab}}) \in SE(3)$ describing the configuration of a frame Σ_b relative to a frame Σ_a . The adjoint transformation associated with g_{ab} , written $Ad_{(g_{ab})}$, is given by

$$Ad_{(g_{ab})} = \begin{bmatrix} e^{\hat{\xi}\theta_{ab}} & \hat{p}_{ab}e^{\hat{\xi}\theta_{ab}} \\ 0 & e^{\hat{\xi}\theta_{ab}} \end{bmatrix}. \quad (37)$$