

离散非线性系统开闭环 P 型迭代学习控制律及其收敛性

皮道映 孙优贤

(浙江大学工业控制技术研究所, 工业控制技术国家重点实验室 · 杭州, 310027)

摘要: 本文在讨论了一般开环与闭环迭代学习控制律的不足后, 针对一类离散非线性系统, 提出了新的开闭环 P 型迭代学习控制律, 给出了它的收敛性证明. 仿真结果表明: 开闭环 P 型迭代律优于单纯的开环或闭环 P 型迭代律.

关键词: 迭代学习控制; 非线性系统; 收敛性

1 引言

近年来, 对迭代学习控制的研究十分活跃^[1~6], 究其原因, 是它适应了控制实践的需要: 它不要求精确知道对象模型的结构、参数, 控制器所具有的学习能力使其能够根据以前的运行数据自动修正不理想的控制输入, 从而使总的控制品质随运行次数的增加而改善, 尤其适用于重复性强、非线性的场合.

经过十多年的研究, 人们已经提出了形形色色的迭代学习控制律, 包括开环迭代律和闭环迭代律. 从研究结果来看, 开环迭代律的控制性能要比闭环的差^[2]. 本文认为, 主要原因在于开环迭代律只利用了系统前次运行的信息. 而闭环迭代律在利用系统当前运行信息改善控制性能的同时舍弃了系统前次运行的信息, 因此提出新的开闭环迭代律, 它能同时利用系统前次运行和当前运行信息, 将能进一步改善控制性能.

2 开闭环 P 型迭代学习控制律及其收敛性

设被控离散非线性系统的动态模型方程为

$$\begin{cases} x(i+1) = f(i, x(i), u(i)), \\ y(i) = g(i, x(i)) + D(i)u(i). \end{cases} \quad (1)$$

式中, i ——时间; $x \in \mathbb{R}^{n \times 1}$ ——状态向量; $y \in \mathbb{R}^{m \times 1}$ ——输出向量; $u \in \mathbb{R}^{r \times 1}$ ——输入向量; f, g ——矩阵函数; D ——具有适当维数的矩阵;

模型的结构与参数未知. 在期望控制量 $u_d(i)$ 存在的前提下, 要求系统在时间区间 $[0, N]$ (N 为时间长度) 上跟踪期望输出 $y_d(i)$.

在第 k 次运行时式(1)表示为

$$\begin{cases} x_k(i+1) = f(i, x_k(i), u_k(i)), \\ y_k(i) = g(i, x_k(i)) + D(i)u_k(i). \end{cases} \quad (2)$$

式中, $x_k(i)$ ——系统在 i 时刻第 k 次运行时的状态值; $u_k(i)$ ——系统在 i 时刻第 k 次运行时的控制量; $y_k(i)$ ——系统在 i 时刻第 k 次运行时的输出值.

定义

$$\begin{cases} \delta x_k(i) = x_d(i) - x_k(i), \\ \delta u_k(i) = u_d(i) - u_k(i), \\ \delta y_k(i) = y_d(i) - y_k(i). \end{cases} \quad (3)$$

式中 $x_d(i)$ 为系统在 i 时刻的期望状态; $\delta x_k(i), \delta u_k(i), \delta y_k(i)$ 分别为第 k 次运行时 x, u, y 在 i 时刻与其期望值的偏差. 学习控制收敛等价于: $\lim_{k \rightarrow \infty} \delta x_k(i) = 0, \lim_{k \rightarrow \infty} \delta u_k(i) = 0, \lim_{k \rightarrow \infty} \delta y_k(i) = 0$.

针对系统(1), 我们所提出的开闭环 P 型迭代学习控制律为

$$u_{k+1}(i) = u_k(i) + \alpha L(i) \delta y_k(i) + (1 - \alpha) L(i) \delta y_{k+1}(i). \quad (4)$$

显然, 当 $\alpha = 0$ 时式(4)为闭环 P 型迭代律; 当 $\alpha = 1$ 时式(4)为开环 P 型迭代律. 定理 1 讨论了系统(1)采用开闭环 P 型迭代学习控制律时的收敛性问题. 为了证明的方便, 先将文[1]的一个引理转述如下(证明见原文):

引理 设 A 是 $r \times r$ 矩阵, 且 A 的谱半径 $\rho(A) < 1$, 若序列 $\{Z_k\}_{k \geq 0}, \{V_k\}_{k \geq 0} \subset \mathbb{R}^r$ 且满足:

$$1) \lim_{k \rightarrow \infty} V_k = 0, \quad 2) Z_{k+1} = AZ_k + V_k.$$

则

$$\lim_{k \rightarrow \infty} Z_k = 0.$$

定理 1 设系统(1)在时间区间 $[0, N]$ 的任一时刻 i 均满足下列条件:

① f, g 连续;

② 存在 $u_d(i)$ 使得系统的状态和输出为期望值 $x_d(i), y_d(i)$ 且 $u_d(i)$ 唯一;

③ 每次重复运行时的初始状态误差 $\{\delta x_k(0)\}_{k \geq 0}$ 是一收敛序列, 终值为零;

④ 矩阵 $[I + (1 - \alpha)L(i)D(i)]$ 是可逆的(其中: I 为适当维数的单位阵). 则对于任意的初始控制 $u_0(i)$ 及每次运行的初始状态 $x_k(0)$, 由式(4)定义的开闭环 P 型迭代学习控制律收敛的充分条件为

$$\rho([I + (1 - \alpha)L(i)D(i)]^{-1}[I - \alpha L(i)D(i)]) < 1, \quad \forall i \in [0, N]. \quad (5a)$$

必要条件为

$$\rho([I + (1 - \alpha)L(0)D(0)]^{-1}[I - \alpha L(0)D(0)]) < 1. \quad (5b)$$

若 L, D 为常数矩阵, 则式(5a)为收敛的充要条件.

证 首先证明开闭环 P 型迭代学习控制收敛的充分条件成立.

由式(2)、(3)、(4)可得

$$\begin{cases} \delta x_k(i+1) = x_d(i+1) - x_k(i+1) \\ \quad = f(i, x_d(i), u_d(i)) - f(i, x_k(i), u_k(i)), \\ \delta y_k(i) = g(i, x_d(i)) + D(i)u_d(i) - g(i, x_k(i)) - D(i)u_k(i) \\ \quad = g(i, x_d(i)) - g(i, x_k(i)) + D(i)\delta u_k(i), \\ \delta u_{k+1}(i) = \delta u_k(i) - L(i)[\alpha \delta y_k(i) + (1 - \alpha)\delta y_{k+1}(i)]. \end{cases} \quad (6)$$

定义辅助函数^[1]

$$\begin{cases} f_1(i, x, u) = f(i, x_d(i), u_d(i)) - f(i, x_d(i) - x, u_d(i) - u), \\ g_1(i, x) = g(i, x_d(i)) - g(i, x_d(i) - x). \end{cases} \quad (7)$$

由条件①有

$$\begin{cases} \lim_{\substack{x \rightarrow 0 \\ u \rightarrow 0}} f_1(i, x, u) = 0, \\ \lim_{x \rightarrow 0} g_1(i, x) = 0. \end{cases} \quad \forall i \in [0, N] \quad (8)$$

式(7)代入方程(6)有

$$\begin{aligned} \delta x_k(i+1) &= f_1(i, \delta x_k(i), \delta u_k(i)), \\ [I + (1 - \alpha)L(i)D(i)]\delta u_{k+1}(i) &= \delta u_k(i) - \alpha L(i)g_1(i, \delta x_k(i)) - (1 - \alpha)L(i)g_1(i, \delta x_{k+1}(i)). \end{aligned} \quad (9)$$

$$= [I - \alpha L(i)D(i)]\delta u_k(i) - L(i)[ag_1(i, \delta x_k(i)) + (1 - \alpha)g_1(i, \delta x_{k+1}(i))].$$

由条件④有

$$\delta u_{k+1}(i) = [I + (1 - \alpha)L(i)D(i)]^{-1}[I - \alpha L(i)D(i)]\delta u_k(i) - [I + (1 - \alpha)L(i)D(i)]^{-1}L(i)[ag_1(i, \delta x_k(i)) + (1 - \alpha)g_1(i, \delta x_{k+1}(i))]. \quad (10)$$

下面用归纳法证明.

$$\lim_{k \rightarrow \infty} \delta x_k(i) = \lim_{k \rightarrow \infty} \delta u_k(i) = \lim_{k \rightarrow \infty} \delta y_k(i) = 0, \quad \forall i \in [0, N]. \quad (11)$$

当 $i = 0$ 时, 据条件③有

$$\lim_{k \rightarrow \infty} \delta x_k(0) = 0. \quad (12)$$

由式(8), (12)有

$$\lim_{k \rightarrow \infty} g_1(0, \delta x_k(0)) = 0. \quad (13)$$

由式(5a), (10), (13)及引理有

$$\lim_{k \rightarrow \infty} \delta u_k(0) = 0. \quad (14)$$

由式(14)及条件②有

$$\lim_{k \rightarrow \infty} \delta y_k(0) = 0, \quad (15)$$

故

$$\lim_{k \rightarrow \infty} \delta x_k(0) = \lim_{k \rightarrow \infty} \delta u_k(0) = \lim_{k \rightarrow \infty} \delta y_k(0) = 0. \quad (16)$$

设当 $i = m$ 时有

$$\lim_{k \rightarrow \infty} \delta x_k(m) = \lim_{k \rightarrow \infty} \delta u_k(m) = \lim_{k \rightarrow \infty} \delta y_k(m) = 0, \quad (17)$$

则当 $i = m + 1$ 时, 由式(9)可得

$$\delta x_k(m + 1) = f_1(m, \delta x_k(m), \delta u_k(m)), \quad (18)$$

由式(8), (17), (18)有

$$\lim_{k \rightarrow \infty} \delta x_k(m + 1) = 0, \quad (19)$$

由式(8), (19)有

$$\begin{cases} \lim_{k \rightarrow \infty} ag_1(m + 1, \delta x_k(m + 1)) = 0, \\ \lim_{k \rightarrow \infty} (1 - \alpha)g_1(m + 1, \delta x_{k+1}(m + 1)) = 0. \end{cases} \quad (20)$$

由式(5a), (10), (20)及引理有

$$\lim_{k \rightarrow \infty} \delta u_k(m + 1) = 0, \quad (21)$$

由式(21)及条件②有

$$\lim_{k \rightarrow \infty} \delta y_k(m + 1) = 0, \quad (22)$$

故

$$\lim_{k \rightarrow \infty} \delta x_k(m + 1) = \lim_{k \rightarrow \infty} \delta u_k(m + 1) = \lim_{k \rightarrow \infty} \delta y_k(m + 1) = 0, \quad (23)$$

由式(16), (17), (23)知式(11)成立. 收敛的充分条件得证.

下面用反证法证明收敛的必要条件成立. 设系统(1)在开闭环 P 型迭代学习控制律(4)下收敛, 并且条件式(5b)不成立. 即

$$\lim_{k \rightarrow \infty} \delta x_k(i) = \lim_{k \rightarrow \infty} \delta u_k(i) = \lim_{k \rightarrow \infty} \delta y_k(i) = 0, \quad \forall i \in [0, N], \quad (24)$$

$$\rho([I + (1 - \alpha)L(0)D(0)]^{-1}[I - \alpha L(0)D(0)]) \geq 1. \quad (25)$$

由条件③, 取学习控制的理想初态情况: $\delta x_k(0) = 0$ ($k = 0, 1, 2, \dots$), 则由式(8)有

$$\begin{cases} g_1(0, \delta x_k(0)) = 0, \\ g_1(0, \delta x_{k+1}(0)) = 0. \end{cases} \quad (26)$$

由式(10), (26)有

$$\begin{aligned} \delta u_{k+1}(0) &= [I + (1 - \alpha)L(0)D(0)]^{-1}[I - \alpha L(0)D(0)]\delta u_k(0) \\ &= \dots \\ &= ([I + (1 - \alpha)L(0)D(0)]^{-1}[I - \alpha L(0)D(0)])^{k+1}\delta u_0(0). \end{aligned} \quad (27)$$

由式(25)和(27)可知: 当 $\delta u_0(0) \neq 0$ 时必有 $\delta u_{k+1}(0) \neq 0$ ($k = 0, 1, 2, \dots$), 这与式(24)矛盾. 故式(5b)为开闭环 P 型迭代学习控制律收敛的必要条件.

显然, 若 L, D 为常数矩阵, 则式(5a)与式(5b)是等价的. 因此此时式(5a)为开闭环 P 型迭代学习控制律收敛的充要条件. 证毕.

3 开闭环 P 型迭代学习控制仿真

设被控系统的动态模型方程为

$$\begin{cases} \begin{bmatrix} x_1(i+1) \\ x_2(i+1) \end{bmatrix} = \begin{bmatrix} -0.20 & 0 \\ -0.25 & 0.15 \end{bmatrix} \begin{bmatrix} x_1^2(i) \\ x_2(i) \end{bmatrix} + \begin{bmatrix} 0.1 & -0.1 \\ 0 & 0.1 \end{bmatrix} \begin{bmatrix} u_1(i) \\ u_2(i) \end{bmatrix}, \\ \begin{bmatrix} y_1(i) \\ y_2(i) \end{bmatrix} = \begin{bmatrix} 0.20 & 0 \\ 0 & 0.1 \end{bmatrix} \begin{bmatrix} x_1^2(i) \\ x_2(i) \end{bmatrix} + \begin{bmatrix} -0.1 & 1 \\ -0.99 & -1 \end{bmatrix} \begin{bmatrix} u_1(i) \\ u_2(i) \end{bmatrix}, \end{cases} \quad (28)$$

要求在 $i \in [0, 10]$ 内跟踪期望输出:

$$y_d(i) = \begin{bmatrix} 0.5i \\ 0.5i \end{bmatrix}.$$

取

$$\alpha = -0.5, \quad L = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}, \quad u_0(0) = \begin{bmatrix} 0 \\ 0 \end{bmatrix}, \quad x_0(0) = \begin{bmatrix} 0 \\ 0 \end{bmatrix},$$

因 $\rho([I + (1 - \alpha)LD]^{-1}[I - \alpha LD]) = 0.6331 < 1$,

开闭环 P 型迭代学习控制律将是收敛的. 仿真结果如图 1 所示, 表明系统经过 9 次迭代学习控制即可实现对期望输出的跟踪. 此时若采用开环 P 型迭代律, 则因 $\rho(I - LD) = 1.786 > 1$, 系统将发散; 若采用闭环 P 型迭代律, 则因 $\rho([I + LD]^{-1}) = 1.005 > 1$, 系统也将发散.

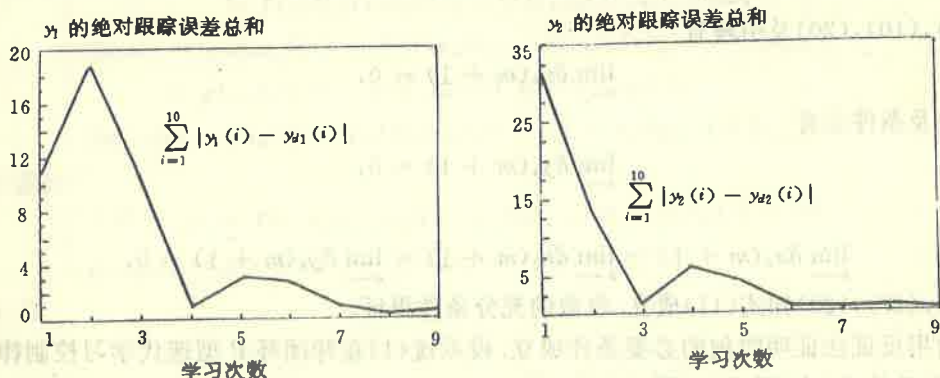


图 1 $\alpha = -0.5$ 时输出量 y_1, y_2 在每次迭代时跟踪误差绝对值总和的变化曲线

取 $\alpha = -0.95$, 其它保持不变, 此时 $\rho([I + (1 - \alpha)LD]^{-1}[I - \alpha LD]) = 0.5594$, 仿真结果(图 2)表明此时系统只需经过 8 次迭代学习控制即可跟踪期望输出.

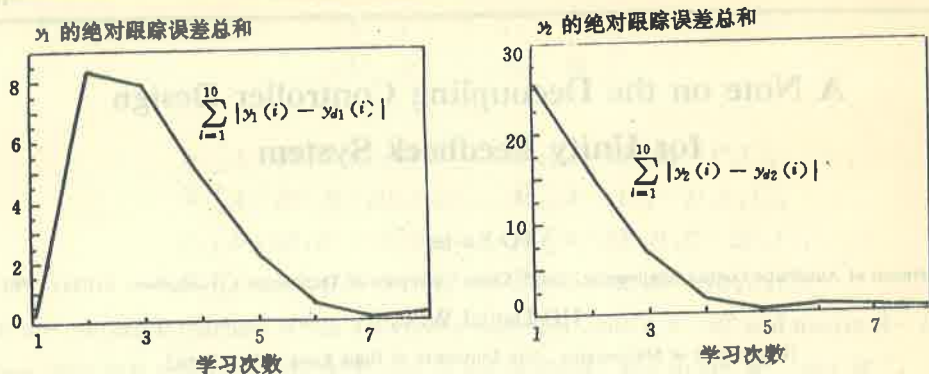


图 2 $\alpha = -0.95$ 时输出量 y_1, y_2 在每次迭代时跟踪误差绝对值总和的变化曲线

4 结 论

容易证明文[1]中的定理 4.3 与 4.7 均为本文定理 1 的特例。理论分析和计算机仿真结果表明：在相同的学习增益矩阵下，开、闭环 P 型迭代学习控制均不能收敛的系统，采用带有参数 α 的开闭环 P 型迭代学习控制后仍有可能收敛；参数 α 对学习控制收敛的速度有直接的影响。

参 考 文 献

- 1 林辉. 迭代学习控制理论研究. 西北工业大学博士学位论文, 1992
- 2 曾南, 应行仁. 非线性系统迭代学习算法. 自动化学报, 1992, 18(2): 168—176
- 3 任雪梅, 高为炳. 任意初始状态下的学习控制. 自动化学报, 1994, 20(1): 74—79
- 4 Arimoto, S.. Learning control theory for robotic motion. Int. J. Adap. Control and Signal Proc., 1990, 4: 543—564
- 5 Heinzinger, G. et al.. Stability of learning control with disturbances and uncertain initial conditions. IEEE Trans. Auto. Control, 1992, 37(1): 110—114
- 6 Sugie, T. and Ono, T.. An iterative learning control law for dynamical systems. Automatica, 1991, 27(4): 729—732

On the Convergence of Open-Closed-Loop P-Type Iterative Learning Control Scheme for Nonlinear Discrete Systems

PI Daoying and SUN Youxian

(Institute of Industrial Process Control, National Key Laboratory of Industrial Control Technology, Zhejiang University • Hangzhou, 310027, PRC)

Abstract: After analysing the weakness of common open-loop and closed-loop iterative learning control schemes, the paper develops a new open-closed-loop P-type version for a class of nonlinear discrete systems. Convergence of the scheme is proved. Simulations show that the new scheme is superior to common open-loop or closed-loop P-type iterative learning control schemes.

Key words: iterative learning control; nonlinear system; convergence

本文作者简介

皮道映 1965 年生. 1985 年和 1988 年分别获得东北大学理学学士学位和工学硕士学位, 1991 年任讲师, 1995 年获浙江大学工学博士学位后留校任教, 1996 年任副教授. 研究领域为多模型控制, 学习控制等.

孙优贤 见本刊 1997 年第 1 期第 41 页.