

基于稳定逆的非最小相位系统的迭代学习控制

刘 山¹, 吴铁军²

(1. 浙江大学 工业控制技术国家重点实验室, 浙江 杭州 310027; 2. 浙江大学 智能系统与决策研究所, 浙江 杭州 310027)

摘要: 针对迭代学习控制在非最小相位系统上应用效果差的缺点, 根据最优化性能指标和非因果的稳定逆理论, 提出了一种基于稳定逆的最优开闭环综合迭代学习控制, 分析了学习律的收敛性并给出了此种非因果的学习律在实际应用中的运用方式。

关键词: 迭代学习控制; 非最小相位; 非因果稳定逆; 优化性能指标

中图分类号: TP273 **文献标识码:** A

Stable-inversion based iterative learning control for non-minimum phase plants

LIU Shan¹, WU Tie-jun²

(1. National Laboratory of Industrial Control Technology, Zhejiang University, Zhejiang Hangzhou 310027, China;

2. Institute of Intelligent Systems and Decision Making, Zhejiang University, Zhejiang Hangzhou 310027, China)

Abstract: In order to deal with the poor tracking effect which occurs when the iterative learning control (ILC) was applied to the non-minimum phase system, an optimal ILC scheme with current feedback was presented for linear non-minimum phase plants based on an optimality criterion and noncausal stable inversion. The convergence of this scheme was analyzed and the utility mode of using the noncausal algorithm was proposed for the practical application.

Key words: iterative learning control; non-minimum phase; noncausal stable inversion; optimality criterion

1 引言(Introduction)

迭代学习控制针对具有重复运行性质的系统, 利用以前运行的信息, 通过迭代的方式修正控制信号, 希望最终实现从系统运行的初始时刻开始, 系统输出就精确跟踪上期望轨迹, 追求对系统的瞬态响应的完全控制, 其迭代学习过程实质上是对系统的逆的逼近过程^[1]。然而, 严格正则系统不存在因果的逆, 非最小相位系统不存在稳定的逆, 对于这两类系统, 采用因果的迭代学习控制会发生困难。同时, 由于滞后的相位, 无论怎样调节控制作用, 都将在控制过程的初始阶段存在跟踪误差。

为了克服系统的逆不存在或不稳定对迭代学习控制造成的困难, 可以在系统的跟踪精度和迭代学习过程的收敛性之间寻求折中。对于最小相位的严格正则系统, 在牺牲对高频信号的跟踪能力的前提下, 引入遗忘因子或者低通滤波器可以较满意地解决逆不存在的问题。但对于非最小相位系统, 由于引起非最小相位系统的逆系统输出发散的信号并不只

是高频信号, 因此, 仅靠滤去高频信号并不能得到令人满意的结果。Roh 等人^[2]针对非最小相位系统, 通过频域优化技术在迭代学习律中引入频段滤波器, 控制效果明显好于仅引入低通滤波器的方法, 但对于因相位滞后而导致的初始阶段的误差则无能为力。实际上, 很多迭代学习律对于非最小相位系统或者收敛速度很缓慢, 或者最终的收敛误差很大, 或者干脆发散^[2~4]。Amann 等人^[3]分析了梯度型最优迭代学习控制对非最小相位系统效果较差的原因, 指出这主要与系统的零动态有关。系统的不可观部分对系统输出的影响不可能通过因果的迭代学习控制消除, 因此总存在较大的收敛误差, 如果系统的不可观部分是不稳定的, 迭代学习控制就失效了。

由频域分析可知, 非最小相位系统存在两种逆, 一种是因果的但不稳定, 另一种则是稳定的但非因果^[5,6]。Devasia 等人^[5]和 Hunt 等人^[6]正是基于此种性质, 提出了系统的稳定逆概念, 并给出了非最小相位的线性和非线性系统的稳定非因果逆的计算方法。

但由于这些算法是非因果的,在反馈控制中只能由预测技术给以实现.

在迭代学习控制的过程中,下一次迭代运行前,系统上一次运行所有时刻的信息都是已知的,为构造系统的非因果稳定逆提供了条件,所以,非因果逆在迭代学习控制中有良好的应用前景. Ghosh 等人^[7]和 Sogo 等人^[8]正是根据稳定的非因果逆的计算方法分别提出了非最小相位系统的多种迭代学习律. Jeong 等人^[9]针对非最小相位离散线性系统,根据以前运行的数据构造系统的非因果逆,提出了一种迭代学习律. 这些学习律均能够使得非最小相位系统实现精确跟踪期望轨迹的任务,但算法都比较复杂且均为开环学习律.

本文根据非最小相位系统的非因果稳定逆理论,依据最优化性能指标,提出了一种基于稳定逆的最优开闭环综合迭代学习控制,然后分析了学习律的收敛性并给出了此种非因果的学习律在实际应用中的运用方式. 最后,给出了针对非最小相位系统的仿真例子.

2 问题描述 (Problem statement)

线性系统 P 第 k 次运行过程的状态方程为

$$\begin{cases} \dot{x}_k(t) = Ax_k(t) + Bu_k(t), \\ y_k(t) = Cx_k(t). \end{cases} \quad (1)$$

其中, $x_k(t) \in \mathbb{R}^n$ 为状态, $u_k(t) \in \mathbb{R}$ 为输入, $y_k(t) \in \mathbb{R}$ 为输出. 系统 P 的传递函数为

$$P(s) = C(sI - A)^{-1}B. \quad (2)$$

若 $P(s)$ 在右半平面存在零点,则 P 是非最小相位系统.

迭代学习控制的目标为: 给定有限时间区间 $[0, T]$ 上系统的期望轨迹为 $y_d(t)$, 假设在每次迭代运行过程中, 系统的初始条件重复, 希望寻找一系列的控制输入 $\{u_k(t)\}$, 在该系列控制输入作用下系统的系列输出为 $\{y_k(t)\}$, 使得随着 k 的增大, $y_k(t)$ 能够在区间 $[0, T]$ 上逐渐逼近 $y_d(t)$.

迭代学习控制过程是一个逐渐逼近系统的逆的过程. 对于非最小相位系统, 如果需要保证因果性, 则其逆是不稳定的, 因此, 在这种情况下, 因果的迭代学习过程的收敛性及稳定性均不能保证.

若放松对系统逆的因果性的要求, 可以得到系统的稳定的逆.

定义 1^[6] (稳定逆) 对于系统(1), 给定 $[0, T]$ 上的期望轨迹 $y_d(t)$, 若存在输入 $u_c(t)$, 状态 $x_c(t)$ 和输出 $y_c(t)$ 满足

1) u_c, x_c 和 y_c 在 $t \in [-\infty, +\infty]$ 上均满足系

统(1)的状态方程, 即

$$\begin{cases} \dot{x}_c(t) = Ax_c(t) + Bu_c(t), \\ y_c(t) = Cx_c(t), \end{cases} \quad t \in [-\infty, +\infty]. \quad (3)$$

2) y_c 在 $[0, T]$ 上精确跟踪期望轨迹 $y_d(t)$, 即

$$y_c(t) = y_d(t), \quad t \in [0, T]. \quad (4)$$

3) u_c, x_c 和 y_c 在 $t \in [-\infty, +\infty]$ 上有界, 且当 $t \rightarrow \pm\infty$ 时, 满足

$$u_c(t) \rightarrow 0, \quad x_c(t) \rightarrow 0, \quad y_c(t) \rightarrow 0, \quad (5)$$

则称 $u_c(t)$ 为系统(1)针对期望轨迹 $y_d(t)$ 的稳定逆.

对于最小相位系统, 稳定逆等价于一般意义上的系统逆, 对于非最小相位系统, 稳定逆则是非因果但保证稳定的逆.

假设 1 系统的期望轨迹充分光滑, 且 $y_d(t) \in L_\infty$ 和 $\frac{d^r}{dt^r}(y_d(t)) \in L_\infty$, 这里, r 为系统 P 的相对阶.

若系统的期望轨迹满足假设 1, 则能保证系统(1)针对期望轨迹 $y_d(t)$ 的稳定逆存在^[6].

非因果的稳定逆在保证系统输出与期望轨迹完全一致的前提下, 同时要求保证逆的稳定性, 达到这一要求的代价是放松了对系统在零时刻的初始状态条件的约束, 它不要求系统的初始条件固定, 相反, 系统在零时刻的初始状态条件是由逆 $u_c(t)$ 在 $(-\infty, 0)$ 时段对系统作用后自动决定的, 当然, 这在实际过程中是不可实现的, 因此是非因果的.

稳定逆实际上与系统的零动态有关. 稳定逆是在 $t < 0$ 的时间内通过输入作用在零动态过程中调整系统的状态, 直到在零时刻时, 系统的状态能够被调整到满足稳定逆的要求, 而后, 系统的输入作用开始使得系统的输出与期望轨迹一致^[6]. 对于任意的期望轨迹, 系统稳定逆的计算是比较复杂的, 需要利用从系统运行的终端时刻以后(最好是 $+\infty$)至系统运行的初始时刻以前(最好是 $-\infty$)的反向积分. 一般要采用迭代法进行计算, 而且要注意计算的稳定性^[5].

下面, 本文考虑用迭代学习控制的方法来得到系统的稳定逆.

首先, 将期望轨迹 y_d 在 $(-\infty, +\infty)$ 上充分光滑地延拓为 y_r , 即选取 $t_0 < 0$ 和在 $[t_0, 0]$ 上充分光滑的 $r_1(t)$, 满足 $r_1(t_0) = 0, r_1(0) = y_d(0)$, 选取 $t_f > T$ 和在 $[T, t_f]$ 上充分光滑的 $r_2(t)$, 满足 $r_2(T) = y_d(T), r_2(t_f) = 0$, 令 $y_r(t)$ 为

$$y_r(t) = \begin{cases} 0, & t \in (-\infty, t_0], \\ r_1(t), & t \in [t_0, 0], \\ y_d(t), & t \in [0, T], \\ r_2(t), & t \in [0, t_f], \\ 0, & t \in [t_f, +\infty). \end{cases} \quad (6)$$

此时,求系统稳定逆的问题转化为求如下最优跟踪问题的解

$$\min_u \{J(u) = \frac{1}{2} \int_{-\infty}^{+\infty} e^T(t) e(t) dt; \\ e = y_r - y, y = Pu\}. \quad (7)$$

且该解满足边界条件

$$u(\pm\infty) = 0, x(\pm\infty) = 0, y(\pm\infty) = 0. \quad (8)$$

3 基于稳定逆的迭代学习控制 (Stable-inversion based iterative learning control)

3.1 基于稳定逆的迭代学习控制设计 (Stable-inversion based iterative learning controller design)

采用最大值原理和梯度法来求解上述最优跟踪问题.引入 Lagrange 乘子向量 $\lambda(t)$, 构成 Hamilton 函数

$$H(x, u, \lambda, t) = \frac{1}{2} (y_r - Cx)^T (y_r - Cx) + \lambda^T (Ax + Bu).$$

给定 $u_k(t)$, 则欲使性能指标 $J(u_k)$ 值的下降, 下一控制输入 u_{k+1} 的搜索方向应为性能指标 $J(u_k)$ 关于 u_k 的负梯度方向, 即为

$$\Delta u_k = -\nabla J(u_k) = -\frac{\partial H}{\partial u_k}(x_k, u_k, \lambda, t) = -B^T \lambda(t).$$

由协态方程, 有

$$\dot{\lambda}(t) = -\frac{\partial H}{\partial x_k}[x_k, u_k, \lambda, t] = -A^T \lambda(t) + C^T e_k(t).$$

其中, $e_k(t) = y_r(t) - y_k(t)$. 将梯度法的求解过程理解为迭代运行过程, 则 e_k 即为第 k 次运行的误差, 由此可得到采用迭代学习控制过程逼近系统稳定逆的算法为

$$u_{k+1}(t) = u_k(t) + \alpha z_k(t). \quad (9)$$

其中, α 为学习步长, $z_k(t)$ 通过以下方程求解

$$\begin{cases} \dot{\lambda}_k(t) = -A^T \lambda_k(t) + C^T e_k(t), \\ z_k(t) = -B^T \lambda_k(t), \end{cases} \quad (10)$$

边界条件为 $\lambda_k(\pm\infty) = 0$.

系统(10)恰为系统(1)的伴随系统 P^* , 且

$$P^*(s) = -B^T(sI + A^T)C^T = P^T(-s). \quad (11)$$

理论上, 采用式(9)和(10)可以迭代计算出系统的稳定逆, 但还存在计算上的问题. 注意到方程(10)的边界条件比较特殊. 要求解该方程, 方法相当繁琐, 必

须首先根据矩阵 A 的特征值对系统进行相似变换, 以使新的矩阵为 Jordan 标准型, 且使 A 的具有正实部的特征值包含在左上角的 Jordan 块中, 使 A 的具有负实部的特征值包含在右下角的 Jordan 块中, 然后将状态相应分为两部分, 与正实部的特征值相关的状态由 $+\infty$ 到 $-\infty$ 反向积分, 与负实部的特征值相关的状态由 $-\infty$ 到 $+\infty$ 正向积分, 最后, 合并状态, 再通过逆相似变换求出原方程的解. 这样做的目的是保证计算的稳定性^[6]. 但如果 A 的特征值均具有负实部, 则只需简单将原方程由 $+\infty$ 到 $-\infty$ 反向积分即可得到解. 本文利用这一特点, 通过引入线性二次型最优输出调节器来简化计算过程.

本文取线性二次型性能指标为^[10]

$$J(u) = \frac{1}{2} \int_{-\infty}^{+\infty} (y^T(t) y(t) + \rho u^T(t) u(t)) dt. \quad (12)$$

其中, $\rho > 0$. 针对系统(1), 结合最优输出调节器的迭代学习控制的具体算法如下

$$\begin{cases} u_k(t) = v_k(t) - \frac{1}{\rho} B^T K x_k(t), \\ v_{k+1}(t) = v_k(t) + \frac{1}{\rho} z_k(t). \end{cases} \quad (13)$$

其中, 常值对称正定矩阵 K 满足代数 Riccati 方程

$$KA + A^T K - \frac{1}{\rho} K B B^T K + C^T C = 0. \quad (14)$$

$z_k(t)$ 通过以下微分方程求解

$$\begin{cases} \dot{\xi}_k(t) = -\left(A - \frac{1}{\rho} B B^T K\right)^T \xi_k(t) + C^T e_k(t), \\ z_k(t) = -B^T \xi_k(t). \end{cases} \quad (15)$$

边界条件为 $\xi_k(\pm\infty) = 0$.

可以看到, 式(9)和(10)相当于仅具有开环作用的迭代学习控制, 而结合最优输出调节器的算法(13)~(15)则是同时具有开环和闭环作用的迭代学习控制, 实际上是基于无限时间 LQR 的最优开闭环综合迭代学习律的非因果表达.

迭代学习算法(13)~(15)与 Amann 等人^[3]提出的算法比较, 有两点大的不同. 一是上述迭代学习律的各参数都是定常的, 而 Amann 算法的各参数都是时变的, 因此, 上述算法比 Amann 算法简单, 计算量大大降低; 二是边界条件不同, 上述迭代学习算法的边界条件使得算法最终将得到系统的稳定逆, 所以, 对于非最小相位系统, Amann 算法存在系统输出极限误差, 而上述算法则可以使系统输出精确跟踪期望轨迹.

3.2 算法收敛性分析 (Convergence analysis of algorithm)

先叙述一个估计系统的 H_∞ 范数的引理.

引理 1^[11] 设系统 P 由传递函数表示为 $P(s) = C(sI - A)^{-1}B$, A, B 和 C 为相应维数的常值矩阵, 若 A 为稳定阵, 则 $\|P\|_\infty < 1$ 的充分必要条件为存在正定阵 X 满足 Riccati 不等式

$$XA + A^T X + XBB^T X + C^T C < 0. \quad (16)$$

定理 1 设系统 P 如式(1)表示, 期望轨迹满足假设 1, 迭代学习控制律如式(13) ~ (15) 所示, 则该迭代学习控制系统的轨迹跟踪误差序列 $\{e_k\}$ 收敛于零, 其中 $e_k = y_r - y_k$.

证 期望轨迹满足假设 1, 故系统针对该期望轨迹的稳定逆存在. 式(13)表明由 v_k 到 y_k 的闭环反馈系统 G 可表示为

$$\begin{cases} \dot{x}_k(t) = \left[A - \frac{1}{\rho} BB^T K \right] x_k(t) + B v_k(t), \\ y_k(t) = C x_k(t). \end{cases} \quad (17)$$

记 $\bar{A} = A - \frac{1}{\rho} BB^T K$, 则系统 G 的传递函数可表示为 $G(s) = C(sI - \bar{A})^{-1}B$, 再由式(15)知, 从 e_k 到 z_k 的系统为 G 的伴随系统 G^* , 有 $G^*(s) = G^T(-s)$.

因为闭环系统为最优无限时间输出调节器系统, 因此 $\bar{A}$ 是稳定的. 考虑系统

$$\bar{G}(s) = C(sI - \bar{A})^{-1} \frac{1}{\sqrt{\eta}} B.$$

其中, $\eta > 0$, 则由式(14)可得到

$$\begin{aligned} K\bar{A} + \bar{A}^T K + K \left(\frac{1}{\sqrt{\eta}} B \right) \left(\frac{1}{\sqrt{\eta}} B \right)^T K + C^T C = \\ \left(\frac{1}{\eta} - \frac{1}{\rho} \right) K B B^T K. \end{aligned}$$

由于 K 是对称正定阵, 由引理 1 知, 当 $\eta > \rho$ 时, 有 $\|\bar{G}\|_\infty = \frac{1}{\sqrt{\eta}} \|G\|_\infty < 1$. 由 η 的任意性知, 此时必有 $\|G\|_\infty < \sqrt{\rho}$. 再由伴随算子的性质^[12], 有

$$\frac{1}{\rho} \|G(s)G^T(-s)\|_\infty = \frac{1}{\rho} \|G(s)\|_\infty^2 < 1. \quad (18)$$

对式(13)取 Laplace 变换, 并代入误差表达式, 可得

$$\begin{aligned} E_{k+1}(s) &= Y_r(s) - G(s)V_{k+1}(s) = \\ Y_r(s) - G(s)V_k(s) - \frac{1}{\rho} G(s)G^T(-s)E_k(s) &= \\ \left(1 - \frac{1}{\rho} G(s)G^T(-s) \right) E_k(s). \end{aligned} \quad (19)$$

由式(18)及 $G(s)G^T(-s)$ 为共轭对称正定阵, 有

$$\left\| 1 - \frac{1}{\rho} G(s)G^T(-s) \right\|_\infty < 1. \quad (20)$$

由式(19), (20)和压缩映射原理, 可知轨迹跟踪误差序列 $\{e_k\}$ 收敛于零. 证毕.

定理 1 保证了迭代学习算法(13) ~ (15)的收敛性.

下面进一步分析算法鲁棒收敛的条件.

按照标称系统的模型参数 $\{A, B, C\}$ 设计迭代学习算法(13) ~ (15). 若真实被控系统的参数为 $\{\bar{A}, \bar{B}, \bar{C}\}$, 记标称系统(1)和真实系统的由输入到状态的传递函数分别为 $H(s) = (sI - A)^{-1}B$ 和 $\tilde{H}(s) = (sI - \bar{A})^{-1}\bar{B}$, 则真实的被控系统为

$$y(s) = \tilde{P}(s)u(s) = \tilde{C}\tilde{H}(s)u(s). \quad (21)$$

假设 $\tilde{H}(s)$ 与标称系统的 $H(s)$ 的不同在于 $\tilde{H}(s)$ 存在乘性不确定性 $\Delta(s)$,

$$\tilde{H}(s) = (1 + \Delta(s))H(s). \quad (22)$$

记

$$S(s) = \left(1 + \frac{1}{\rho} B^T K H(s) \right)^{-1},$$

$$T(s) = 1 - S(s),$$

$$v_k(t) = u_k(t) + \frac{1}{\rho} B^T K x_k(t).$$

由式(13)和(17), 可得到标称系统由 v_k 到 y_k 的传递函数为

$$G(s) = CH(s)S(s) = P(s)S(s). \quad (23)$$

不确定系统由 v_k 到 y_k 的传递函数为

$$\begin{aligned} \tilde{G}(s) &= C \left(sI - \bar{A} + \frac{1}{\rho} \bar{B} \bar{B}^T K \right)^{-1} \bar{B} = \\ C \left(I + \frac{1}{\rho} \tilde{H}(s) B^T K \right)^{-1} \tilde{H}(s) &= \\ \tilde{P}(s) \left(1 + \frac{1}{\rho} B^T K \tilde{H}(s) \right)^{-1}. \end{aligned} \quad (24)$$

将式(22)和(23)代入式(24), 可得

$$\begin{aligned} \tilde{G} &= C(1 + \Delta)H \left(1 + \frac{1}{\rho} B^T K(1 + \Delta)H \right)^{-1} = \\ C(1 + \Delta)HS(1 + \Delta(1 - S))^{-1} &= \\ (1 + \Delta)(1 + \Delta T)^{-1}G &= \\ (1 + \Delta S(1 + \Delta T)^{-1})G. \end{aligned} \quad (25)$$

另一方面, 由迭代学习律(13) ~ (15)及系统方程(21)可知, 误差传递方程为

$$E_{k+1}(s) = \left(1 - \frac{1}{\rho} \tilde{G}(s)G^T(-s) \right) E_k(s). \quad (26)$$

由式(25), 可得

$$\begin{aligned} 1 - \frac{1}{\rho} \tilde{G}(s)G^T(-s) &= \\ 1 - \frac{1}{\rho} G(s)G^T(-s) - \Delta S(1 + \Delta T)^{-1} \frac{1}{\rho} G(s)G^T(-s). \end{aligned} \quad (27)$$

因此,保证算法收敛的一个充分条件为

$$\left\| 1 - \frac{1}{\rho} G(s) G^T(-s) - \Delta S(1 + \Delta T)^{-1} \frac{1}{\rho} G(s) G^T(-s) \right\|_{\infty} < 1. \quad (28)$$

定理 2 设不确定系统 $\bar{P}$ 如式(21)和(22)所示,若系统的不确定性满足

$$\| \Delta(s) S(s) (1 + \Delta(s) T(s))^{-1} \|_{\infty} < 1, \quad (29)$$

则按标称系统参数设计的迭代学习控制律(13)~(15)收敛.

证 记 $b(s) = \Delta(s) S(s) (1 + \Delta(s) T(s))^{-1}$, 由条件(29)可知,对于任意的 $\omega \in \mathbb{R}$,

$$|b(j\omega)| < 1. \quad (30)$$

同时,由定理1的证明知 $\|G\|_{\infty} < \sqrt{\rho}$,因此,对于任意的 $\omega \in \mathbb{R}$,

$$0 \leq \frac{1}{\rho} G(j\omega) G^T(-j\omega) < 1. \quad (31)$$

即 $\frac{1}{\rho} G(j\omega) G^T(-j\omega)$ 为小于1的非负实数.记 $a(s) = \frac{1}{\rho} G(s) G^T(-s)$,由上式可知,对于任意的 $\omega \in \mathbb{R}$, $a(j\omega)$ 也为小于1的非负实数,且

$$|1 - a(j\omega)| = 1 - a(j\omega). \quad (32)$$

由式(29)和(31)可以得到,对于任意的 $\omega \in \mathbb{R}$,

$$|b(j\omega)| + a(j\omega) < a(j\omega).$$

将上式两边同时加1并移项,得到

$$1 - a(j\omega) + |b(j\omega)| + a(j\omega) < 1. \quad (33)$$

因此,由式(32)和(33),对于任意的 $\omega \in \mathbb{R}$,

$$|1 - \frac{1}{\rho} G(j\omega) G^T(-j\omega) - \Delta(j\omega) S(j\omega) (1 + \Delta(j\omega) T(j\omega))^{-1} \frac{1}{\rho} G(j\omega) G^T(-j\omega)| =$$

$$|1 - a(j\omega) - b(j\omega) a(j\omega)| <$$

$$|1 - a(j\omega)| + |b(j\omega)| + a(j\omega) =$$

$$1 - a(j\omega) + |b(j\omega)| + a(j\omega) < 1.$$

所以,式(28)成立. 证毕.

定理2表明本文提出的迭代学习控制设计允许一定的模型失配.

3.3 迭代学习律的实际控制方式 (Practical mode in application of iterative learning law)

由迭代学习控制律(13)~(15)的边界条件可以看到,它要求系统的运行时间从 $-\infty$ 到 $+\infty$,这在实际过程中是无法实现的,因此,必须考虑迭代学习律(13)~(15)在应用中的实际控制方式.

本文采用将期望轨迹平移的方法来近似实现迭代学习控制律(13)~(15).若系统的期望轨迹 y_d 在 $(-\infty, +\infty)$ 上充分光滑地延拓为 y_r . 取充分大的

T_a ,考虑新的期望轨迹为 $y_d^{(new)}$

$$y_d^{(new)}(t) = y_r(t - T_a), \quad \forall t \in (-\infty, +\infty). \quad (34)$$

每次系统的运行时间也由原来的 $[0, T]$ 改为 $[0, 2T_a + T]$. 注意到闭环系统为最优无限时间输出调节器系统,保证了 $\bar{A}$ 的特征值均具有负实部,因此,针对方程(15),双边界条件之一 $\xi_k(-\infty) = 0$ 恒成立.所以,迭代学习律(13)~(15)的边界条件相应地修改为

$$\xi_k(2T_a + T) = 0. \quad (35)$$

在上述修改的条件下,迭代学习律(13)~(15)是可以实现的.此时,迭代学习控制所感兴趣的 $[0, T]$ 时间段的期望轨迹被平移到了 $[T_a, T_a + T]$ 段,因此,在迭代学习控制作用下,系统的输出将在 $[T_a, T_a + T]$ 时间段逼近期望轨迹.

上述方法是将原来在无限长时间段上求稳定逆的问题在一个有限段时间上进行近似,由稳定逆定义的条件3)可以知道,当时间朝正、负两个方向增大时,稳定逆 $u_c(t)$ 的值逐渐逼近于零,因此,这样的近似造成的损失并不大.

上述运行过程从零时刻开始,但所感兴趣时段的初始时刻为 T_a ,在 T_a 时刻之前的控制作用的目的是为了调整系统运行到 T_a 时刻的内部状态,以保证系统输出能够在 $[T_a, T_a + T]$ 时段精确跟踪期望轨迹.

4 仿真研究 (Simulation illustrations)

首先将采用稳定逆的方案与不采用稳定逆的方案进行对比.

例 1 假设非最小相位的不稳定线性系统为

$$P(s) = \frac{s - 5}{(s - 2)(s - 3)},$$

系统在 $t \in [0, 6]$ 上的期望轨迹为 $y_d(t) = \sin(0.5\pi t)$. 本例分别采用不使用和使用稳定逆的迭代学习控制方案控制上述系统,进行仿真对比研究.两种方案均采用式(13)~(15)进行迭代学习控制算法设计,且均取 ρ 为 0.01,初始控制均为零输入.不同之处在于对期望轨迹的处理和算法边界条件的选取.

不使用稳定逆的方案采用原期望轨迹,算法的边界条件为 $\xi_k(6) = 0$.

使用稳定逆的方案将期望轨迹进行延拓和平移,平移时间为 $T_a = 10$,因此,期望轨迹经延拓和平移后成为 $t \in [0, 26]$ 上新的期望轨迹

$$y_r^{(new)} = \begin{cases} \sin(0.5\pi(t-10)), & t \in [10, 16], \\ 0, & t \in [0, 10] \cup [16, 26]. \end{cases}$$

此时,算法的边界条件为 $\xi_k(26) = 0$.

图1表示系统输出的误差范数随迭代学习次数

变化的情况,其中,“*”代表不采用稳定逆方案的结果,“o”代表采用稳定逆方案的结果.图1表明,经过5次左右迭代,两种算法的输出误差都收敛了,这说明它们的收敛速度都很快,但不采用稳定逆的方案有较大的固定极限误差,极限误差范数达到0.19左右,而与之相反,基于稳定逆的方案将趋于零.

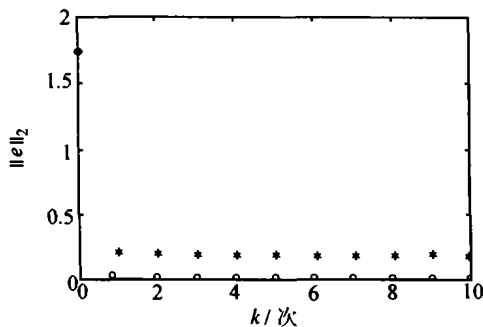


图1 输出误差随迭代次数变化情况

Fig. 1 Output errors vs iterations

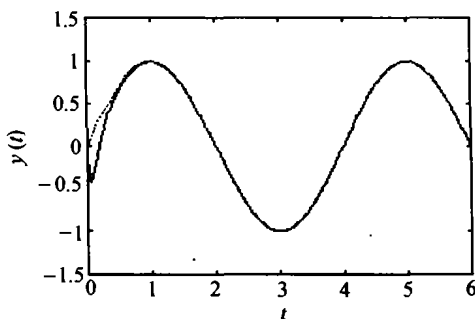


图2 不采用稳定逆时迭代运行输出
Fig. 2 Output not using stable inversion

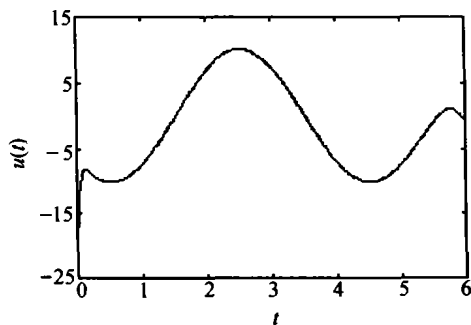


图4 不采用稳定逆时迭代运行输入
Fig. 4 Control input not using stable inversion

下一个例子以车载倒立摆为例,说明基于稳定逆的迭代学习控制的效果.

例2 考虑车载倒立摆在平衡点附近的线性化系统^[13]

$$M\ddot{\theta} = (M + m)g\theta - u,$$

$$M\ddot{x} = u - mg\theta.$$

其中, M 和 m 分别为小车与摆的质量, l 为摆长,

图2和图3分别表示不采用稳定逆方案时和采用稳定逆方案时系统第10次运行的输出轨迹与期望轨迹的比较,其中,虚线代表期望轨迹,实线代表系统的实际输出轨迹.图4和图5分别表示不采用稳定逆方案时和采用稳定逆方案时系统第10次运行的系统输入.图2表明不采用稳定逆的方案在初始的一段时间内无法精确跟踪期望轨迹,而图3表明采用稳定逆的方案则从一开始即准确跟踪了期望轨迹.产生这样的结果的原因可从图4和图5看出,不采用稳定逆的方案,控制是从零时刻开始作用,尽管从一开始作用时,控制的作用就很大,但仍无法使系统跟踪上期望轨迹,这是由系统非最小相位的延迟作用的本质决定的.而由图5可以看到,采用稳定逆的方案是从“零”时刻之前就开始对系统进行控制作用的,注意到期望轨迹经平移后,“零”时刻实际上是 $t = T_a$ 处,这样,就“克服”了非最小相位的延迟.

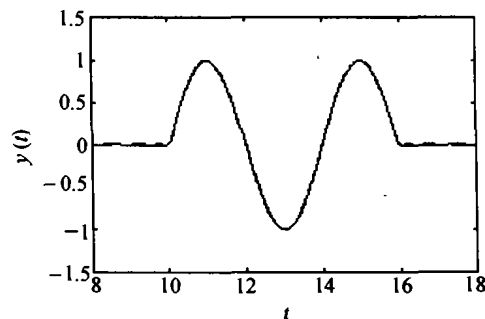


图3 采用稳定逆时迭代运行输出
Fig. 3 Output using stable inversion

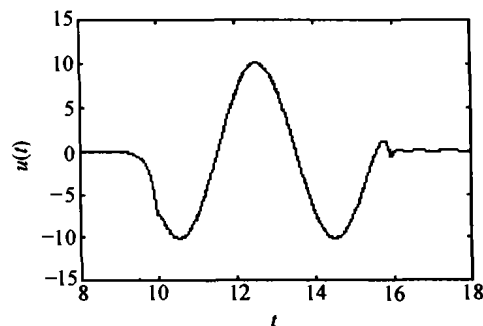


图5 采用稳定逆时迭代运行输入
Fig. 5 Control input using stable inversion

x 代表小车的位置, θ 表示摆偏离竖直位置的角度.期望小车的位置在 $t \in [0, 6]$ 时间段按期望轨迹

$$y_d(t) = 0.1t^2(6-t)\sin(0.5\pi t)$$

变化.本例假设 $M = 2 \text{ kg}$, $m = 0.1 \text{ kg}$, $l = 0.5 \text{ m}$, $g = 9.81 \text{ m/s}^2$. 则上述系统的极点为 $0, 0, 4.5388$ 和 -4.5388 , 零点为 4.4294 和 -4.4294 , 因此, 该系统是具有不稳定零极点的非最小相位系统. 按系统的

参数用迭代学习律(13)~(15)设计控制,取 ρ 为0.01,初始控制为零输入.将期望轨迹进行延拓和平移,平移时间为 $T_a = 10$,此时,算法的边界条件为 $\xi_k(26) = 0$.

图6表示系统的输出误差范数随迭代学习次数变化的情况,可以看到,算法很快就收敛了.图7代表系统第20次运行时的输出轨迹与期望轨迹的比较,其中,虚线代表期望轨迹,实线代表系统的实际输出轨迹.图8表示系统第20次运行时的输入.可以看到系统的输出轨迹精确地跟踪了期望轨迹,而控制也是从“零”时刻之前就开始对系统进行作用了.

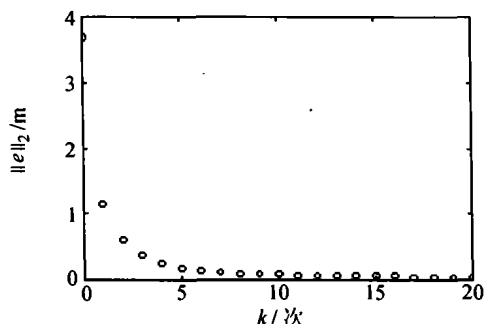


图6 输出误差随迭代次数变化情况

Fig. 6 The output errors vs. iterations

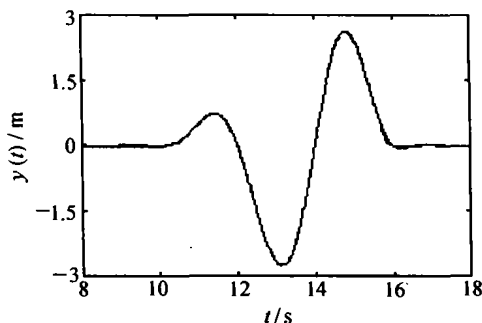


图7 第20次迭代运行的系统输出

Fig. 7 System output in 20th iteration

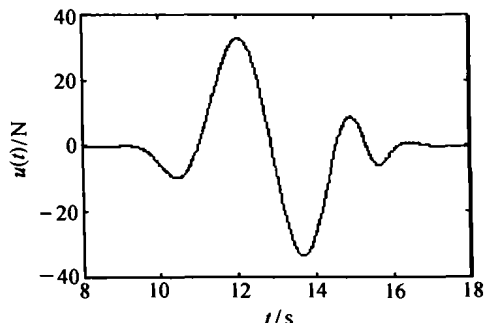


图8 第20次迭代运行的系统输入

Fig. 8 Control input in 20th iteration

5 结论(Conclusion)

本文根据非因果的稳定逆理论,针对非最小相

位系统上提出了一种基于稳定逆的最优开闭环综合迭代学习控制,并通过将期望轨迹延拓和平移的方式给出了此种非因果的学习律在实际应用中的运用方式.在非最小相位对象上的仿真结果验证了算法的可行性和有效性.

参考文献(References):

- [1] MOORE K L. *Iterative Learning Control for Deterministic Systems* [M]. London: Springer, 1993.
- [2] ROH C L, LEE M N, CHUNG M J. ILC for non-minimum phase system [J]. *Int J of Systems Science*, 1996, 27(4):419-424.
- [3] AMANN N, OWENS D H. Non-minimum phase plants in iterative learning control [A]. *Proc of the 2nd IEE Int Conf on Intelligent Systems Engineering* [C]. Hamburg-Harburg: [s. n.], 1994: 107-112.
- [4] GHOSH J, PADEN B. Iterative learning control for nonlinear non-minimum phase plants [J]. *J of Dynamic Systems, Measurement, and Control*, 2001, 123(1):21-30.
- [5] DEVASIA S, CHEN D, PADEN B. Nonlinear inversion-based output tracking [J]. *IEEE Trans on Automatic Control*, 1996, 41(7): 930-942.
- [6] HUNT L R, MEYER G, SU R. Noncausal inverses for linear systems [J]. *IEEE Trans on Automatic Control*, 1996, 41(4):608-611.
- [7] GHOSH J, PADEN B. Pseudo-inverse based iterative learning control for plants with unmodelled dynamics [A]. *Proc of the American Control Conf* [C]. Chicago, Illinois: [s. n.], 2000:472-476.
- [8] SOGO T, KINOSHITA K, ADACHI N. Iterative learning control using adjoint systems for nonlinear non-minimum phase systems [A]. *Proc of the 39th IEEE Conf on Decision and Control* [C]. Sydney: [s. n.], 2000:3445-3446.
- [9] JEONG G M, CHOI C H. Iterative learning control for linear discrete time nonminimum phase systems [J]. *Automatica*, 2002, 38(2):287-291.
- [10] ANDERSON B O D, MOORE J B. *Optimal Control-Linear Optimal Control* [M]. Englewood Cliffs: Prentice Hall, 1989.
- [11] 申铁龙. H_∞ 控制理论及应用 [M]. 北京:清华大学出版社, 1996.
(SHEN Tielong. *H_∞ Control Theory and Applications* [M]. Beijing: Tsinghua University Press, 1996.)
- [12] 柳重堪. 应用泛函分析 [M]. 北京:国防工业出版社, 1986.
(LIU Chongkan. *Applied Functional Analysis* [M]. Beijing: National Defence Industry Press, 1986.)
- [13] OGATA K. *Designing Linear Control Systems with Matlab* [M]. Englewood Cliffs: Prentice Hall, 1993.

作者简介:

刘山 (1970—),男,博士.现为浙江大学智能系统与决策研究所助理研究员.研究方向为智能控制,学习控制及其在工业中的应用. E-mail: slui@iipc.zju.edu.cn;

吴铁军 (1950—),男,现为浙江大学智能系统与决策研究所所长,教授,博士生导师.研究方向为大系统智能控制,非线性控制,混杂动态系统控制与决策及其在复杂系统中的应用等.