

doi:10.11918/j.issn.0367-6234.2016.05.029

可拓聚类适应度共享小生境遗传算法研究

李中华, 张泰山

(中南大学 信息科学与工程学院, 410083 长沙)

摘 要: 针对遗传算法易陷入早熟收敛和全局搜索能力差等缺点,提出一种基于可拓理论的小生境遗传算法. 算法首先构造了遗传编码物元和可拓遗传算子,然后通过可拓聚类方法实现小生境群体的划分,结合适应度共享技术和聚类代表个体保存策略,维持稳定多样的小生境. 仿真实验表明,该算法能可靠、快速地收敛到全局最优解,有效避免早熟收敛,其收敛速度和求解精度均优于简单遗传算法和常规小生境算法.

关键词: 遗传算法;小生境;可拓聚类;适应度共享;代表个体;早熟收敛

中图分类号: TP319.4 **文献标志码:** A **文章编号:** 0367-6234(2016)05-0178-06

Research of fitness sharing niche genetic algorithms based on extension clustering

LI Zhonghua, ZHANG Taishan

(School of Information Science and Engineering, Central South University, 410083 Changsha, China)

Abstract: To solve the problems of premature convergence and weak ability in global search of the genetic algorithm, a fitness sharing niche genetic algorithm based on extenics is proposed. The algorithm build the matter-element code and extension genetic operator, create niche groups by extension clustering, and preserve the stability of niche groups by combining fitness sharing mechanism and elitist retention strategy. Experiments show that the algorithm can solve the optimal performance with global search ability and fast convergence rate. It is proved to be more effective and accurate than standard geneic algorithm and normal niche genetic algorithm.

Keywords: genetic algorithm; niche; extension clustering; fitness sharing; elitist retention; premature convergence

基本遗传算法(SGA)是一类模拟生物进化的智能优化算法,已广泛应用于工程领域的优化求解. 但大量实践和研究表明,基本遗传算法存在全局搜索能力差和早熟收敛等问题,被称为基因漂移^[1]. 针对这一问题,学者提出了各种小生境遗传算法(NGA),但这些经典的小生境技术均存在参数设置问题^[2-5]. 这些参数的预设与优化需要一定的先验知识,处置不当将难以维持稳定的小生境,从而限制了算法的有效运用.

本文基于可拓理论^[6]和全局性优化考虑,提出可拓聚类适应度共享小生境遗传算法(ECNGA). 算法基本原理为:构造可拓遗传算法模型,提出遗传信息物元编码和可拓遗传算子;利用可拓聚类方法对每代种群划分形成小生境群体;根据适应度共享机制调整

个体适应度,再种群进化,并结合聚类小生境的代表个体保存策略,维持小生境的稳定性和多样性,从而有效地防止算法早熟收敛,增强全局搜索能力.

1 可拓聚类分析

定义 1 设聚类 $H_p = \{R_1, R_2, \dots, R_{n_p}\}$, 类内样本个体的物元模型^[6]为

$$R = [N, C, V] = \begin{bmatrix} N & C_1 & v_1 \\ & C_2 & v_2 \\ & \vdots & \vdots \\ & C_m & v_m \end{bmatrix}.$$

对 $i = 1, 2, \dots, m$, 各物元特征 C_i 的属性值分别为 $v_{1i}v_{2i}\cdots v_{n_p i}$. 令

$$a_{pi} = \min(v_{1i}, v_{2i}, \dots, v_{n_p i}),$$

$$b_{pi} = \max(v_{1i}, v_{2i}, \dots, v_{n_p i}),$$

$$M_{pi} = \frac{(v_{1i} + v_{2i} + \dots + v_{n_p i})}{n_p},$$

收稿日期: 2015-03-05.

作者简介: 李中华(1968—), 男, 博士, 副教授;

张泰山(1940—), 男, 教授, 博士生导师.

通信作者: 李中华, chinali@csu.edu.cn.

称类 H_p 各特征属性取值区间为 $V_{pi}(C_{pi}) = [a_{pi}, b_{pi}]$.

定义 R_p 为类 H_p 的中心物元, 记为

$$R_p = \begin{bmatrix} N & C_1 & M_{p1} \\ & C_2 & M_{p2} \\ & \vdots & \vdots \\ & C_m & M_{pm} \end{bmatrix}.$$

若 $n_p = 1$ 时, 令 $V_{pi}(C_{pi}) = [a_i, b_i]$, 即各特征属性取值为节域区间, 同时令中心物元为 $R_p = R_1$.

定义 2 关联度计算准则^[6-8]:

1) 已知类 $H_p = \{R_1, R_2, \dots, R_{n_p}\}$, 取待聚类样本 $R_x = [N, C, V]$, 样本的 m 维特征观测值为 $V = (x_1, x_2, \dots, x_m)^T$.

2) 对类 H_p , 计算 R_x 各特征属性 x_i 的基本关联函数为

$$k_p(x_i) = \begin{cases} \frac{x_i - a_{pi}}{M_{pi} - a_{pi}}, & x_i \leq M_{pi}; \\ \frac{b_{pi} - x_i}{b_{pi} - M_{pi}}, & x_i > M_{pi}. \end{cases}$$

3) 计算综合关联度 $K_x(H_p)$. 若规定各项聚类特征的权重系数为 $\lambda_p = [\lambda_{p1}, \lambda_{p2}, \dots, \lambda_{pm}]$, 则

$$K_x(H_p) = \sum_{i=1}^m \lambda_{pi} k_p(x_i).$$

权重系数表示各项特征在聚类评价体系中的相对重要程度, 可采用专家评价法、层次分析法或比重权数法确定^[7].

关联度 $K_x(H_p)$ 表示样本 R_x 与类 $H_p = \{R_1, R_2, \dots, R_{n_p}\}$ 的关联程度, 其取值大小是判定样本是否归属该类的直接依据.

可拓聚类方法的基本流程如下:

Step 1 按照聚类样本的指标或维数建立物元模型为

$$R_i = [N, C, V] = \begin{bmatrix} N & C_1 & v_{i1} \\ & C_2 & v_{i2} \\ & \vdots & \vdots \\ & C_m & v_{im} \end{bmatrix}.$$

($i = 1, 2, \dots, n$, n 是样本总数)

令 $a_j = \min_{i=1}^n v_{ij}, b_j = \max_{i=1}^n v_{ij}$, 则 $V_j = [a_j, b_j]$ 是特征 $C_j(j = 1, 2, \dots, m)$ 的节域区间.

Step 2 初始化聚类. 将待聚类物元按某种规则排序(例如遗传优化的适应度排序), 并产生一随机数 $k(1 < k < n)$, 取前 k 个物元形成 k 个初始类.

Step 3 依次取剩余样本 R_x , 计算 R_x 对类 $H_p(p = 1, 2, \dots, k)$ 的关联度 $K_x(H_1), K_x(H_2), \dots, K_x(H_k)$.

Step 4 判定待聚类样本 R_x 的类属. 令 $K_x(H_q) =$

$\max\{K_x(H_1), K_x(H_2), \dots, K_x(H_k)\}$. 若 $K_x(H_q) > 0$, 则 R_x 属于类 $H_q; K_x(H_q) = 0$, 可认为 R_x 属于或不属于类 $H_q; K_x(H_q) < 0, R_x$ 不属于已知类, 需增加新的聚类, 即设样本 R_x 为新类别, 聚类数目 $k = k + 1$.

Step 5 调整新增样本的类 H_q 或新增类的属性取值区间、中心物元, 并求平均关联度. H_q 类内平均关联度记为

$$\bar{K}_q = \frac{1}{n_q} \sum_{i=1}^{n_q} K_i(H_q).$$

式中, n_q 是 H_q 类内样本数目, $K_i(H_q)$ 是类内第 i 个样本的关联度.

Step 6 类属约简. 取类 H_q 的中心物元 R_q , 计算 R_q 对所有类 H_p 的关联度 $K_q(H_p)(p = 1, 2, \dots, k, p \neq q)$. $K_q(H_p)$ 表示类 H_q, H_p 之间的相似性. 若 $K_q(H_p) \geq \bar{K}_q$, 表示两个类属相似度较强, 可合并为一个类, 聚类数 $k = k - 1$, 并计算约简后的取值区间、中心物元及平均关联度.

Step 7 跳转 Step3, 处理所有待聚类的 $n - k$ 个样本物元.

Step 8 聚类检验.

1) 对类 $H_p(p = 1, 2, \dots, k)$, 计算并固定中心物元 R_p ;

2) 计算所有物元 R_x 对所有类的关联 $K_x(H_p)(x = 1, 2, \dots, n, p = 1, 2, \dots, k)$;

3) 令 $K_x(H_q) = \max\{K_x(H_1), \dots, K_x(H_k)\}$, 聚类样本 R_x 归属至类 H_q ;

4) 循环以上 3 步, 直至所有样本个体的聚类归属不变.

Step 9 输出聚类结果, 包括聚类数目 k , 每一个聚类 H_p 下的样本数 n_p , 及类内样本关于该聚类的关联度 $K_x(H_p)$.

定义 3 若样本集合 $X(t)$ 有 k 个聚类, $X(t) = \{H_1, H_2, \dots, H_k\}$, 各聚类样本数目为 $(n_1, n_2, \dots, n_k)$, 集合总样本数目 $n = \sum_{i=1}^k n_i$, 则集合聚类的熵^[5]为

$$E_t = - \sum_{i=1}^k \frac{n_i}{n} \ln\left(\frac{n_i}{n}\right).$$

熵是反映集合内样本聚类多样性程度的重要指标.

2 可拓聚类适应度共享小生境遗传算法

2.1 DNA 编码物元

定义 4 设遗传信息为物元的集合, 记为 $R_i = [N, C, V]$. 基因特征属性 $C_j(j = 1, 2, \dots, m)$ 的量值为 v_{ij} , 其定义域 $V_j = [a_j, b_j]$.

将属性量值用遗传算法编码 $(l_1, l_2, \dots, l_i)$ 表

示,则每项特征所对应的编码为

$$D_{ij}(V_j) = \begin{cases} l_1, v_{ij} \in X_1; \\ l_2, v_{ij} \in X_2; \\ \vdots \\ l_t, v_{ij} \in X_t. \end{cases}$$

式中, $X_1, X_2, \dots, X_t$ 为在 $[-\infty, +\infty]$ 的 t 个区间.

经过量值编码后,将遗传信息物元扩展定义为 DNA 编码物元^[7],记为

$$R_i = \begin{bmatrix} N & C_1 & D_{i1} \\ & C_2 & D_{i2} \\ & \vdots & \vdots \\ & C_m & D_{im} \end{bmatrix}.$$

2.2 可拓遗传算子

定义 5 设 DNA 编码物元 $R_0 \in R$, 将 R_0 改变为另一个同类 DNA 编码物元或多个同类编码物元的可拓变换,称为 R_0 的可拓遗传算子,记作

$$T\{R_0\} = \{R_1\} \text{ 或 } T\{R_0\} = \{R_1, R_2, \dots\}.$$

可拓遗传算子可用事元^[6]的形式化语言表示为

$$T = \begin{bmatrix} O_T & ct1 & vt1 \\ & ct2 & vt2 \\ & ct3 & vt3 \\ & ct4 & vt4 \\ & \vdots & \vdots \end{bmatrix} = \begin{bmatrix} \text{遗传} & \text{支配编码} & vt1 \\ & \text{变换量} & vt2 \\ & \text{遗传结果} & vt3 \\ & \text{施动编码} & vt4 \\ & \vdots & \vdots \end{bmatrix}.$$

式中, O_T 是动作的名称,表示实施遗传操作的类型,即

$$O_T \in \{\text{复制, 交叉, 变异, 置换, 增删, 扩缩, } \dots\}.$$

可拓遗传可解析为: $vt4$ 对 $vt1$ 实施变换 O_T , 变换量为 $vt2$, 变换结果为 $vt3$. 结合物元可拓变换的多种形式, 可得一些常用或特有的可拓遗传算子.

1) 复制算子(或称选择算子) $T\{R\} = \{R\}$, 即

$$\begin{bmatrix} \text{复制} & ct1 & R \\ & ct2 & R \\ & ct3 & R \end{bmatrix} \{R\} = \left\{ \begin{bmatrix} N & C_1 & D_1 \\ & \vdots & \vdots \\ & C_m & D_m \end{bmatrix} \right\}.$$

常用的复制选择方法有适应度比例选择法、锦标赛选择法和排序选择法.

2) 交叉算子 $T\{R_1, R_2\} = \{R'_1, R'_2\}$. 基本的有单点交叉、两点交叉、多点交叉等.

例如, 父代个体 $\{R_1, R_2\}$, 运用两点交叉算子, 操作结果为

$$\begin{bmatrix} ct1 & R_1 \\ ct2 & (i, j) \\ ct3 & \{R'_1, R'_2\} \\ ct4 & R_2 \end{bmatrix} \{R_1, R_2\} =$$

$$\left\{ \begin{bmatrix} N & C_1 & D_{11} \\ & \vdots & \vdots \\ & C_{i-1} & D_{i-1} \\ & C_i & D_{2i} \\ & \vdots & \vdots \\ & C_j & D_{2j} \\ & C_{j+1} & D_{1j+1} \\ & \vdots & \vdots \\ & C_m & D_{1m} \end{bmatrix}, \begin{bmatrix} N & C_1 & D_{21} \\ & \vdots & \vdots \\ & C_{i-1} & D_{2i-1} \\ & C_i & D_{1i} \\ & \vdots & \vdots \\ & C_j & D_{1j} \\ & C_{j+1} & D_{2j+1} \\ & \vdots & \vdots \\ & C_m & D_{2m} \end{bmatrix} \right\}.$$

式中, (i, j) 表示交叉的基因座位置.

若为实数编码的算术交叉算子, 取随机数 $\alpha \in (0, 1)$, 则有

$$\begin{bmatrix} \text{交叉} & ct1 & R_1 \\ & ct2 & \alpha \\ & ct3 & \{R'_1, R'_2\} \\ & ct4 & R_2 \end{bmatrix} \{R_1, R_2\} = \left\{ \begin{bmatrix} N & C_1 & \alpha D_{11} + (1 - \alpha) D_{21} \\ & \vdots & \vdots \\ & C_m & \alpha D_{1m} + (1 - \alpha) D_{2m} \end{bmatrix}, \begin{bmatrix} N & C_1 & \alpha D_{21} + (1 - \alpha) D_{12} \\ & \vdots & \vdots \\ & C_m & \alpha D_{2m} + (1 - \alpha) D_{1m} \end{bmatrix} \right\}.$$

3) 置换变异算子 $T\{R\} = \{R'\}$, 即

$$\begin{bmatrix} \text{置换} & ct1 & R \\ & ct2 & (i, l_i) \\ & ct3 & R' \end{bmatrix} \{R\} = \left\{ \begin{bmatrix} N & C_1 & D_1 \\ & \vdots & \vdots \\ & C_{j-1} & D_{i-1} \\ & C_j & l_i \\ & C_{j+1} & D_{i+1} \\ & \vdots & \vdots \\ & C_{1m} & D_m \end{bmatrix} \right\}.$$

表示随机取基因座 $i (1 \leq i \leq m)$, 在 i 处以编码 l_i 替换.

若为实数编码的变异算子, 则可取 $l_i = U_{\min}^i + r(U_{\max}^i - U_{\min}^i)$, 其中随机数 $r \in [0, 1]$, $[U_{\min}^i, U_{\max}^i]$ 表示变异点处基因值的取值范围.

4) 扩缩算子 $T\{R\} = \alpha R$, 是指扩大或缩小随机选择的基因编码的量值, 即

$$\begin{bmatrix} \text{扩缩} & ct1 & R \\ & ct2 & (i, \alpha) \\ & ct3 & \alpha R \\ & \vdots & \vdots \end{bmatrix} \{R\} = \left\{ \begin{bmatrix} N & C_1 & D_1 \\ & \vdots & \vdots \\ & C_{i-1} & D_{i-1} \\ & C_i & \alpha D_i \\ & C_{i+1} & D_{i+1} \\ & \vdots & \vdots \\ & C_{1m} & D_m \end{bmatrix} \right\}.$$

当 $\alpha > 1$ 时为扩大算子, 当 $0 < \alpha < 1$ 时为缩小算子。

5) 增删算子 $T\{R\} = \{R'\}$, 即

$$\left[\begin{array}{ccc} \text{增删} & ct1 & R \\ & ct2 & i \\ & ct3 & R' \\ & \vdots & \vdots \end{array} \right] \{R\} = \left\{ \left[\begin{array}{ccc} N & C_1 & D_1 \\ & \vdots & \vdots \\ & C_{j-1} & D_{i-1} \\ & C_j & D'_i \\ & C_{j+1} & D_{i+1} \\ & \vdots & \vdots \\ & C_{lm} & D_m \end{array} \right] \right\}.$$

当 $D'_i = U_{\max}^i$ 时为增加算子, 当 $D'_i = U_{\min}^i$ 时为删除算子。

2.3 小生境构造及适应度共享

基于可拓聚类方法构造小生境, 对个体适应度共享。其基本过程为:

1) 对每代种群 $X(t)$ 进行可拓聚类。聚类数目 k , 每一个类 H_p 下的样本数 n_p , 类内样本关联度 $K_x(H_p)$ 。

2) 定义每个聚类中个体的小生境数^[8-9]为

$$m_i = n_p \times K_i(H_p), i = 1, 2, \dots, n_p.$$

式中, m_i 为个体 i 的小生境数, 其值越大, 表示其相似个体越多, 反之越稀疏。

3) 适应度共享, 调整小生境中个体的适应度为

$$f'_i = f_i / m_i = f_i / (n_p \times K_i(I_p)).$$

式中 f_i, f'_i 分别为共享前、后的适应度。

ECNGA 算法对种群聚类, 并选出每类的代表个体^[10]保存至下一代种群。代表个体保存策略如下:

Step 1 标记种群 $X(t)$ 的聚类代表个体为 $\{a_1(t), a_2(t), \dots, a_k(t)\}$;

Step 2 遗传生成种群 $X(t+1)$, 并可拓聚类;

Step 3 取 $a_i(t) (i = 1, 2, \dots, k)$, 判断其聚类归属 $H_p (p \in [1, 2, \dots, k'])$;

Step 4 以代表个体替换聚类 H_p 中适应度最低的个体;

Step 5 比较代表个体 $a_i(t)$ 与目前聚类中适应度最高个体的 $b_p(t+1)$, 若 $f(a_i(t)) < f(b_p(t+1))$, 则重新标记 $b_p(t+1)$ 为代表个体。

采用适应度共享机制和代表个体保存策略进行遗传操作, 可使种群维持稳定且多样的小生境, 提高算法的全局搜索能力。

2.4 交叉概率 P_c 和变异概率 P_m

为促进种群的多样性, P_c, P_m 应与小生境熵 E_i 呈负相关关系。另外, 对适应度高的个体, 可增大 P_c 以加快其进化速度, 减小 P_m 以避免破坏其结构模式。因此, 设计自适应调整 P_c 和 P_m 为

$$P_c = \begin{cases} \eta_c \cdot \exp(-E_i) \cdot \frac{f_{\max} - \bar{f}}{f_{\max} - f_c}, & f_c \geq \bar{f}; \\ \eta_c \cdot \exp(-E_i), & f_c < \bar{f}. \end{cases}$$

$$P_m = \begin{cases} \eta_m \cdot \exp(-E_i) \cdot \frac{f_{\max} - f_m}{f_{\max} - \bar{f}}, & f_m \geq \bar{f}; \\ \eta_m \cdot \exp(-E_i), & f_m < \bar{f}. \end{cases}$$

式中: $0 < \eta_c < 1$ 且 $0 < \eta_m \leq 1$; $f_{\max}, \bar{f}$ 为每代种群中最大的适应度及平均适应度; f_c 为交叉两个体中较大的适应度; f_m 为变异个体的适应度。

2.5 ECNGA 可拓小生境遗传算法流程

Step 1 算法初始化, 即设定种群规模 n, P_c, P_m 等参数和置终止准则, 随机生成初始种群 $X(t) = \{R_i, i = 1, 2, \dots, n\}$, 置进化代数 $t \leftarrow 0$ 。

Step 2 计算种群 $X(t)$ 中个体适应度 f_i , 按适应度降序排列。

Step 3 对种群 $X(t) = \{R_i, i = 1, 2, \dots, n\}$ 进行可拓聚类, 构造 k 个小生境。

Step 4 代表个体保存并更新标记。

Step 5 计算种群中所有个体的共享适应度 f'_i 。

Step 6 种群进化, 即可拓遗传形成种群 $X(t+1)$ 。

1) 可拓选择从父代 $X(t)$ 中运用选择算子选择出 $n/2$ 对个体。

2) 可拓交叉 $n/2$ 对个体依概率 P_c 进行交叉, 形成 n 个中间个体。

3) 可拓变异依概率 P_m 对 n 个中间个体变异, 形成子代 $X(t+1)$ 。

Step 7 终止检验。若不满足终止准则, 置 $t \leftarrow t+1$ 并跳转 step2; 满足时则输出种群 $X(t+1)$ 中最优解。

2.6 算法收敛性分析

引理 1^[10-11] 基本遗传算法 SGA 收敛到全局最优解的概率小于 1。

引理 2^[10-11] 保留最优个体的 GA 算法以概率 1 收敛到全局最优解。

定理 1 采用代表个体保存策略的 ECNGA 算法以概率 1 收敛到全局最优解。

证明 标记代表个体 $\{a_1(t), a_2(t), \dots, a_k(t)\} \in X(t)$, 令

$$a_p(t) = \max\{a_1(t), a_2(t), \dots, a_k(t)\}.$$

即 $a_p(t)$ 为 $X(t)$ 的最优个体。

同理, 令 $X(t+1)$ 的最优个体 $b_q(t+1) = \max\{b_1(t+1), b_2(t+1), \dots, b_k'(t+1)\}$ 。

1) 设 $p = q$, 若 $a_p(t) > b_q(t+1)$, 则最优个体

$a_p(t)$ 被保存;反之即存在超过 $a_p(t)$ 的最优个体 $b_q(t+1)$,且 $b_q(t+1)$ 被保存.

2)设 $p \neq q, a_p(t) > b_q(t+1)$,则 $a_p(t) > b_p(t+1)$,最优个体 $a_p(t)$ 在类 H_p 中保存; $a_p(t) < b_q(t+1)$,则 $a_q(t) < b_q(t+1)$,最优个体 $b_q(t+1)$ 在类 H_q 中保存.

综上,采用代表个体保存策略的 ECNGA 算法在每代选择操作前保留最优个体,依据引理 2, ECNGA 算法以概率 1 收敛到全局最优解.

3 仿真实验及分析

实验 1 经典多峰 Shubert 函数寻优.

$$\begin{aligned} \min \{f(x_1, x_2)\} &= \sum_{i=1}^5 \text{icos}[(i+1)x_1 + 1] \cdot \\ &\sum_{i=1}^5 \text{icos}[(i+1)x_2 + 1], \\ \text{s.t. } -10 &\leq x_1, x_2 \leq 10. \end{aligned}$$

Shubert 函数有 760 个局部最优点和 18 个全局最优点,全局最优目标函数值 $-186.730\ 9$. 取适应度函数为

$$\text{Fit}(x_1, x_2) = \begin{cases} -0.05f(x_1, x_2), & f(x_1, x_2) < 0; \\ 0, & f(x_1, x_2) \geq 0. \end{cases}$$

同时比较 SGA、模糊聚类 FCNGA^[4] 及 ECNGA 三种算法. 参考文献[4]设置算法参数:每次随机产生初始种群,规模 $n=100$, 进化代数 $G_{\max}=500$, 交叉算子 $\eta_c=0.6$, 变异算子 $\eta_m=0.01$, 终止条件为目标函数值与当前最优解差 0.1% 时或进化代数达 $G_{\max}$, 优化重复次数 50. 仿真结果如表 1 所示,进化曲线图 1 所示.

表 1 多峰函数寻优仿真结果对比

算法类型	平均进化代数	最优解精度	寻优正确率/%
SGA	456.2	-186.088 0	56
FCNGA	177.4	-186.718 9	72
ECNGA	126.3	-186.730 9	94

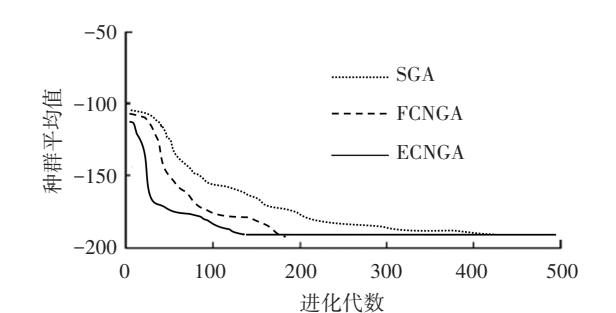


图 1 种群平均值进化曲线

从表 1 数据及图 1 进化曲线可知,ECNGA 算法明显优于其他两种算法,搜索能力较强,有更高的求解质量和更快的收敛速度,并避免了出现早熟而陷

入局部极值的情形.

实验 2 自调整模糊控制器优化设计.

图 2 所示为基于可拓遗传优化的自调整模糊控制系统结构图. 图 2 中各物理量分别为系统给定值 $r(k)$ 、输出响应 $y(k)$ 、误差 $e(k)$ 、误差变化率 $e_c(k)$ 及控制输出量 $u(k)$.

模糊控制规则解析式为

$$U = -[\alpha E + (1 - \alpha)E_c],$$

$$\alpha = \frac{1}{n} \cdot (\alpha_H - \alpha_L) \cdot |E| + \alpha_L.$$

式中: E 、 E_c 为误差、误差变化率量化取整, U 为控制量量化取整,修正因子 $\alpha \in [\alpha_L, \alpha_H]$ 且 $0 < \alpha_L \leq \alpha_H < 1$.

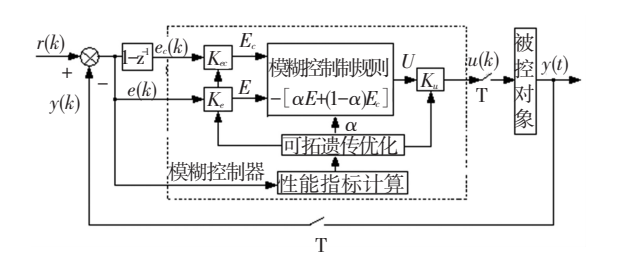


图 2 可拓遗传优化的自调整模糊控制系统结构图

图 2 中可拓遗传优化机构的作用是根据所选取的寻优目标函数,搜索修正因子 α 、量化因子 K_e 、 K_{ec} 及比例因子 K_u 等复杂多维参数的全局最优解,构成最优模糊控制器.

选取 ITAE 作为寻优目标函数,即

$$J(\text{ITAE}) = \min \left[\sum_{k=1}^n kT^2 |e(k)| \right].$$

式中, T 为采样周期, n 为决定 ITAE 指标的总时间.

在 MATLAB/SIMULINK 中构建寻优模糊控制系统. 设对象为多数工控过程的二阶惯性加延迟的模型:

$$G(s) = \frac{e^{-0.2s}}{(2s+1)(5s+1)}.$$

设置寻优约束条件^[12]为

$$0 < K_e < \frac{1}{2\delta}, 0 < K_{ec} < \frac{1}{\delta}, 0 < K_u < \frac{u_{\max}}{l},$$

$$0 < \alpha_L \leq \alpha_H < 1.$$

式中: δ 为控制要求的稳态误差,取 $\delta=0.01$; $u_{\max}$ 为被控对象允许输入的最大值,仿真模块内取 $u_{\max}=12\text{ V}$.

设优化参数编码物元 $R = [N \ C \ D]$,物元特征 $C = (K_e, K_{ec}, K_u, \alpha_L, \alpha_H)^T$, D 为对应浮点数编码. 为对比观察,采用参考文献[12]的 SGA 方法同时参数寻优. 随机生成规模 $n=200$ 的初始种群,取遗传参数 $\eta_c=0.6$ 、 $\eta_m=0.02$ 和进化终止代数 $G_{\max}=100$.

ECNGA 算法寻优后获得最优控制参数为
 $(K_e, K_{ec}, K_u, \alpha_L, \alpha_H) = (12.7, 36.1, 1.57, 0.23, 0.84)$,
同等条件下,SGA 算法优化参数为
 $(K_e, K_{ec}, K_u, \alpha_L, \alpha_H) = (13.6, 41.4, 1.69, 0.26, 0.78)$.

将两组优化参数分别代入自调整模糊控制仿真模块,即可得控制结果曲线,如图 3、4 所示. 采样周期 $T = 0.1\text{ s}$.

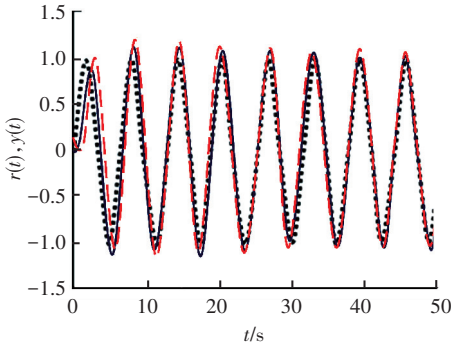


图 3 正弦信号跟踪

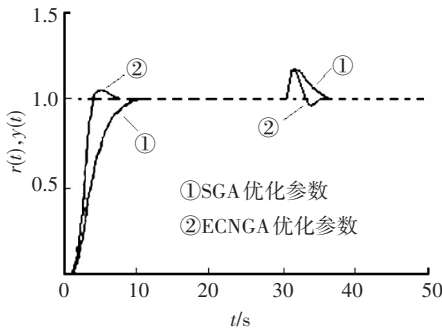


图 4 阶跃信号跟踪

1) 正弦信号跟踪实验. 如图 3 曲线所示,点形虚线为给定信号 $r(t) = \sin t$, 取 ITAE 控制指标总时间为 50. 线段虚线为 SGA 优化的模糊控制正弦响应曲线,其 $ITAE = 273.1$, 实线为 ECNGA 优化后的模糊控制正弦响应曲线,其 $ITAE = 211.5$.

2) 阶跃信号跟踪实验. 如图 4 所示,给定阶跃信号 $r(t) = \varepsilon(t)$, 10 s 前为模糊控制阶跃响应上升阶段. 曲线①为 SGA 优化参数, $ITAE = 11.2$, 稳态 $e(\infty) = 0.004\ 7$; 曲线②为 ECNGA 优化参数, $ITAE = 5.9$, 稳态 $e(\infty) = -0.001\ 4$. 在 30 s 处施加负载扰动时,可观察优化模糊控制器能更快恢复稳定.

对比两组优化参数可知,ECNGA 优化方法可获得全局最优解,其模糊控制器在信号跟踪的动静态性能方面优于 SGA 优化模糊控制,其响应曲线虽然稍有超调,但调节时间大大缩短,控制精度和稳态性能也有了比较明显的改进,而且对扰动的鲁棒稳定性更好.

4 结 论

1)使用可拓聚类分析对种群划分小生境,按照共享机制对个体的适应度进行调整,可以增大种群多样性和有效防止早熟收敛;结合聚类代表个体保存策略,维持稳定的小生境,以增强全局搜索能力.

2)引入度量种群多样性的小生境熵自动调整进化参数,可有效平衡算法快速性和多样性的矛盾.

3)该算法优于标准遗传算法和常规小生境遗传算法,具有全局寻优能力和较高的搜索效率,是一种适用于复杂寻优问题的有效优化算法.

参考文献

[1] 玄光男,程润伟. 遗传算法与工程优化[M].北京:清华大学出版社,2004:21-30.
[2] CHEN C H, LIU T K, CHOU Y H. A novel crowding genetic algorithm and its applications to manufacturing robots[J]. IEEE Transactions on Industrial Informatics, 2014, 10(3): 1705-1706.
[3] KIM Hyoungjin, LIOU Mengsing. New fitness sharing approach for multi-objective genetic algorithm[J]. Journal of Global Optimization, 2013, 55(3): 579-595.
[4] 谭艳艳,许峰.自适应模糊聚类小生境遗传算法[J].计算机工程与应用, 2009, 45(4): 52-55.
[5] 陆青,梁昌勇,杨善林,等.面向多模态函数优化的自适应小生境遗传算法[J].模式识别与人工智能, 2009, 22(1): 91-99.
[6] 杨春燕,蔡文.可拓学[M].北京:科学出版社,2014: 18-96.
[7] 王科俊,徐晶,王磊,等.基于可拓遗传算法的机器人路径规划[J].哈尔滨工业大学学报,2006, 38(7): 1135-1138.
[8] 赵敏,林道荣,瞿波,等.一种新的基于小生境模拟退火的遗传算法[J].辽宁工程技术大学学报, 2013, 32(3): 367-372.
[9] SANTRA A K., CHRISTY C. J, NAGARAJAN B. Cluster based hybrid niche mimetic and genetic algorithm for text document categorization [J]. International Journal of Computer Science Issues, 2011, 8(2): 450-456.
[10] 何琳,王科俊,李国斌,等.最优保留遗传算法及其收敛性分析[J].控制与决策, 2000, 15(1): 63-66.
[11] LI Junhua, LI Ming. An analysis on convergence and convergence rate estimate of elitist genetic algorithms in noisy environments [J]. Optik, 2013, 124(24): 6780-6785.
[12] 葛平淑,郭烈.参数自调整的月球车路径跟踪模糊控制器设计[J].计算机工程与应用, 2012, 48(11): 55-59.

(编辑 王小唯 苗秀芝)