

多智能体遗传算法用于超高维函数优化*

钟伟才 薛明志 刘 静 焦李成

西安电子科技大学雷达信号处理国家重点实验室, 西安 710071

摘要 基于智能体对环境的感知与反作用的能力提出了一种新的函数优化方法——多智能体遗传算法。该方法将智能体固定在网格上, 而每个智能体为了增加自身能量将与其邻域展开竞争或合作, 同样, 智能体也可利用自身的知识进行自学习来增加能量。理论分析证明算法具有全局收敛性。实验结果表明, 多智能体遗传算法对维数高达甚至 10000 的函数, 都能以较少的计算量获得高质量的解, 充分说明算法具有很快的收敛速度。

关键词 智能体 遗传算法 函数优化 进化计算

遗传算法是通过模拟生物在自然界中的进化过程而形成的一种全局优化算法, 目前已得到广泛的应用^[1]。对于函数优化问题, 遗传算法是一种很有潜力的方法, 主要在于它对函数的类型、搜索空间的形状等都没有任何限制。但当维数增高时, 搜索空间变大, 变量间耦合度增强, 从而导致遗传算法陷入局部极值的概率增大。为提高算法性能, 提出了一些新方法, 如非单调适应值标度变换方法^[2], 基于函数分解的可伸缩宏进化算法^[3], 自适应演化算法^[4](AEA)等, 均获得了良好的效果。

文献[5]利用智能体解决了高达 7000 个皇后的问题, 显示出智能体的优越性能。受其启发, 本文针对函数优化问题, 将智能体对环境的感知和反作用的能力与遗传算法的搜索方式相结合, 提出了多智能体遗传算法 (multi-agent genetic algorithm, MAGA), 设计了 4 个进化算子来实现智能体间的竞争、合作、自学习等行为, 其收敛速度远远高于传统算法。而收敛速度慢, 尤其当搜索空间较大时, 一直是影响遗传算法实用性的关键因素之一, 但本文提出的方法对上千、上万维的函数, 都能快速找到高质量的解, 这充分显示了方法的实用价值, 同时也显示出智能体与遗传算法解函数优化问题的巨大潜力。

1 多智能体遗传算法

根据文献[6]的定义, 任何一个能够感知环境并反作用于环境的实体都可看作为智能体, 而所有由多个自治成分构成的系统都可表达为一个智能体系统, 其特性主要表现为: (1) 每个智能体只有局部感知能力; (2) 没有系统的全局控制; (3) 数据是分散的; (4) 计算是异步的。由此可见, 对于不同问题, 智能体的具体含义不同, 我们要根据问题灵活设计智能体的含义与行为。

1.1 用于函数优化的智能体

函数优化问题的数学模型为

$$\min f(x) \quad x = (x_1, \dots, x_n) \in S, \quad (1)$$

其中 $f(x)$ 为目标函数, $S \subseteq R^n$ 表示边界为 $x_i \leq x_i \leq \bar{x}_i, i = 1, \dots, n$ 的 n 维搜索空间。

在本文中, 一个智能体就表示目标函数的一个解, 即一个实值向量 $(x_1, \dots, x_n)$ 。所有智能体都被放在一个规模为 $L_{\text{size}} \times L_{\text{size}}$ 的网格 L 上, 每个智能体占一个格点位置且不能移动。由于智能体只有局部感知能力, 因此它只能与周围的智能体发生相互作用, 这就构成了如图 1 所示的智能体网格。根据图 1, 智能体邻域的概念定义如下:

2003-01-17 收稿, 2003-04-28 收修改稿

* 国家自然科学基金重点资助项目(批准号: 60133010)

E-mail: neouma@163.com

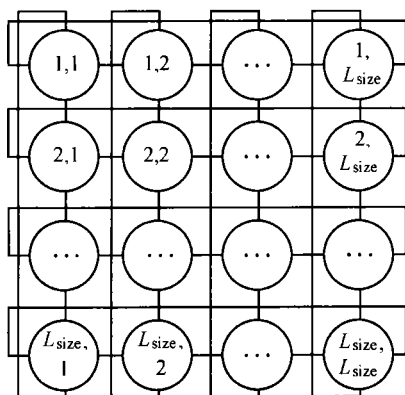


图 1 智能体网格

圆圈代表智能体，圆中的数字代表该智能体的位置，有连线的两个智能体才能发生相互作用

定义 1 将位置 (i, j) 的智能体表示成 $L_{i,j}$, $i, j = 1, 2, \dots, L_{size}$, $L_{i,j}$ 的邻域为 $Nbs_{i,j}$, $Nbs_{i,j} = \{L_{i-1,j}, L_{i,j-1}, L_{i+1,j}, L_{i,j+1}\}$, 其中

$$i_1 = \begin{cases} i-1 & i \neq 1 \\ L_{size} & i = 1 \end{cases}, j_1 = \begin{cases} j-1 & j \neq 1 \\ L_{size} & j = 1 \end{cases},$$

$$i_2 = \begin{cases} i+1 & i \neq L_{size} \\ 1 & i = L_{size} \end{cases}, j_2 = \begin{cases} j+1 & j \neq L_{size} \\ 1 & j = L_{size} \end{cases}.$$

每个智能体都具有一定的能量，其值等于目标函数

$$e_k = \begin{cases} \underline{x}_k & (m_k + U(-1, 1) \times (m_k - l_k)) < \underline{x}_k \\ \bar{x}_k & (m_k + U(-1, 1) \times (m_k - l_k)) > \bar{x}_k, k = 1, \dots, n, \\ m_k + U(-1, 1) \times (m_k - l_k) & \text{其他} \end{cases} \quad (2)$$

$$m'_k = (m_k - \underline{x}_k) / (\bar{x}_k - \underline{x}_k), k = 1, \dots, n, \quad (3)$$

$$New_{i,j} = (m'_1, \dots, m'_{i_1-1}, m'_{i_2}, m'_{i_2-1}, \dots, m'_{i_1+1}, m'_{i_1}, m'_{i_2+1}, \dots, m'_n), 1 < i_1 < n, 1 < i_2 < n, i_1 < i_2, \quad (4)$$

$$e_k = \underline{x}_k + e'_k \cdot (\bar{x}_k - \underline{x}_k), k = 1, \dots, n. \quad (5)$$

邻域正交交叉算子：正交交叉算子是文献[7]提出的一种新交叉算子，篇幅所限，这里不再详述。而邻域正交交叉算子将正交交叉算子^[7]作用在 $L_{i,j}$ 和 $Max_{i,j}$ 上以产生一个新的智能体来代替 $L_{i,j}$ 。

变异算子：根据(6)式产生新智能体 $New_{i,j} = (e_1, e_2, \dots, e_n)$ ：

$$e_k = \begin{cases} l_k & U(0, 1) < \frac{1}{n}, k = 1, \dots, n, \\ l_k + G(0, 1/t) & \text{其他} \end{cases} \quad (6)$$

值的相反数。智能体进化的目标就是尽可能的增大自身的能量，这就导致了智能体间激烈的竞争，但这种竞争只能存在于智能体与其邻域间。最终，能量低的智能体将死亡，而它的邻域将占领它的位置。当然，合作也可能发生在智能体与其邻域间。另一方面，由于智能体具有智能，它也可利用自己的知识来进行自学习。

1.2 智能体进化算子

邻域竞争算子和邻域正交交叉算子利用邻域的信息来分别实现智能体间的竞争与合作，变异算子和智能体自学习算子利用智能体自身的知识来增加能量。设进化算子作用在智能体 $L_{i,j} = (l_1, l_2, \dots, l_n)$ 上， $Max_{i,j} = (m_1, m_2, \dots, m_n)$ 为 $L_{i,j}$ 邻域内能量最大的智能体：

邻域竞争算子：若 $L_{i,j}$ 比 $Max_{i,j}$ 的能量大，则可继续存活在网格上；否则必须死亡，空出的格点将被 $Max_{i,j}$ 占据——如果 $U(0, 1) < P_o$, $Max_{i,j}$ 则按(2)式产生新智能体 $New_{i,j} = (e_1, e_2, \dots, e_n)$ 放在空出的格点上，否则先按(3)式将 $Max_{i,j}$ 映射到区间 $[0, 1]$ 上，然后根据(4)式确定 $New'_{i,j} = (e'_1, e'_2, \dots, e'_n)$ ，最后根据(5)式将 $New'_{i,j}$ 映射回区间 $[\underline{x}_k, \bar{x}_k]$ 上以得到 $New_{i,j}$ 。上述 $U(\cdot, \cdot)$ 为均匀分布的随机数， P_o 为占据概率。

其中 $G(0, 1/t)$ 为 Gauss 分布的随机数， t 为进化代数。然后用 $New_{i,j}$ 代替 $L_{i,j}$ 。

智能体自学习算子：该算子类似于一个小规模的多智能体遗传算法，其详细步骤在算法 1 中给出。

算法 1 智能体自学习算子

$sGen$ 表示最大代数， $sR \in [0, 1]$ 表示相对搜索半径， $sL_{size} < L_{size}$, sP_m 为预先设定的参数， $Eng(\cdot)$ 表示智能体的能量。 $sBest^l$ 和 $sCBest^l$ 分别表示到第 l 代为止产生的最优智能体和第 l 代中的最优

智能体.

步骤1: 根据(7), (8)式产生 $sL_{size} \times sL_{size}$ 的智能体网格 $sL_{i,j}^0$, $i', j' = 1, 2, \dots, sL_{size}$, 更新 $sBest^0$, $l \leftarrow 0$;

$$e_{i',j',k} = \begin{cases} \underline{x}_k & l_k \cdot U(1-sR, 1+sR) < \bar{x}_k \\ \bar{x}_k & l_k \cdot U(1-sR, 1+sR) > \bar{x}_k, k = 1, \dots, n \\ l_k \cdot U(1-sR, 1+sR) & \text{其他} \end{cases} \quad (8)$$

步骤2: 对 sL^l 上的每个智能体执行邻域竞争算子, 最终得到 $sL^{l+1/2}$;

步骤3: 对 $sL^{l+1/2}$ 上的每个智能体, 如果 $U(0, 1) < sP_m$, 则执行变异算子, 最终得到 sL^{l+1} ;

步骤4: 从 sL^{l+1} 中找出 $sCBest^{l+1}$, 如果 $\text{Eng}(sCBest^{l+1}) > \text{Eng}(sBest^l)$, 则令 $sBest^{l+1} \leftarrow sCBest^{l+1}$, 否则令 $sBest^{l+1} \leftarrow sBest^l$, $sCBest^{l+1} \leftarrow sBest^l$;

步骤5: 如果 $l < sGen$, 则令 $l \leftarrow l+1$, 并转步骤2; 否则转步骤6;

步骤6: $L_{i,j} \leftarrow sBest^l$.

1.3 多智能体遗传算法

多智能体遗传算法将邻域竞争算子作用在每个智能体上, 因此能量较低的智能体将被从网格上清除出去, 以使得有较大的发展空间留给能量较高的智能体. 为了降低计算代价, 智能体自学习算子只作用在当代中最优的智能体上. 每个智能体的进化只利用其邻域或自身的信息, 因此整个算法没有全局性控制, 从而是完全异步的. 算法2详细描述了多智能体遗传算法的整体流程.

算法2 多智能体遗传算法

$Best^t$ 和 $CBest^t$ 分别表示到第 t 代为止产生的最优智能体和第 t 代中的最优智能体.

步骤1: 初始化智能体网格 L^0 , 更新 $Best^0$, $t \leftarrow 0$;

步骤2: 对 L^t 上的每个智能体执行邻域竞争算子, 最终得到 $L^{t+1/3}$;

步骤3: 对 $L^{t+1/3}$ 中的每个智能体, 如果 $U(0, 1) < P_c$, 则执行邻域正交叉算子, 最终得到 $L^{t+2/3}$;

步骤4: 对 $L^{t+2/3}$ 中的每个智能体, 如果 $U(0, 1) < P_m$, 则执行变异算子, 最终得到 L^{t+1} ;

$$sL_{i',j'} = \begin{cases} L_{i,j} & i' = 1 \text{ and } j' = 1 \\ \text{New}_{i',j'} & \text{其他} \end{cases}, \quad (7)$$

其中 $\text{New}_{i',j'} = (e_{i',j',1}, e_{i',j',2}, \dots, e_{i',j',n})$ 由(8)式确定;

步骤5: 从 L^{t+1} 中找出 $CBest^{t+1}$, 将智能体自学习算子作用在 $CBest^{t+1}$ 上;

步骤6: 如果 $\text{Eng}(CBest^{t+1}) > \text{Eng}(Best^t)$, 则令 $Best^{t+1} \leftarrow CBest^{t+1}$, 否则令 $Best^{t+1} \leftarrow Best^t$, $CBest^{t+1} \leftarrow Best^t$;

步骤7: 如果停止准则满足, 则输出 $Best^t$, 停止; 否则令 $t \leftarrow t+1$, 并转步骤2.

2 多智能体遗传算法的收敛性分析

如果 $(x_1, \dots, x_n)$ 中的分量 x_i 用 M 位二进制串表示, 则相当于将区间 $[\underline{x}_i, \bar{x}_i]$ 量化为 2^M 个离散值, 其精度为 $\epsilon = (\bar{x}_i - \underline{x}_i)/2^M$, 因此本文直接通过实数编码来分析 MAGA 的收敛性. 设所需精度为 ϵ , 则搜索空间 S 可看成是离散空间, 其大小为 $|S| = \prod_{i=1}^n \lceil (\bar{x}_i - \underline{x}_i)/\epsilon \rceil$, 每个元素 $x \in S$ 就是一个智能体. 令 $E = \{\text{Eng}(x) | x \in S\}$, 显然 $|E| \leq |S|$, 于是 E 可以写成 $E = \{E^1, E^2, \dots, E^{|E|}\}$, 这里 $E^1 > E^2 > \dots > E^{|E|}$. 根据能量的不同又可把集合 S 划分为若干非空子集 $\{S^i\}$, 这里

$$S^i = \{x | x \in S \text{ and } \text{Eng}(x) = E^i\}, \quad i = 1, 2, \dots, |E|, \quad (9)$$

$$\sum_{i=1}^{|E|} |S^i| = |S|; \forall i \in \{1, 2, \dots, |E|\},$$

$$S^i \neq \emptyset; \forall i \neq j, S^i \cap S^j = \emptyset; \bigcup_{i=1}^{|E|} S^i = S, \quad (10)$$

显然, E^1 即为全局最优解, 且集合 S^1 包含了所有能量等于 E^1 的智能体.

在 MAGA 的整个进化过程中, 智能体的个数是保持不变的. 令 $\mathcal{S}$ 表示所有智能体网格的集合, 由于一个智能体网格中的智能体是允许相同的, 因此有(11)式:

$$|\mathcal{L}| = \left(|S| + L_{\text{size}} \times L_{\text{size}} - 1 \right) / (L_{\text{size}} \times L_{\text{size}}). \quad (11)$$

为了衡量智能体网络的优劣, 定义智能体网络 L 的能量为

$$\text{Eng}(L) = \max\{\text{Eng}(L_{i,j}) \mid i, j = 1, \dots, L_{\text{size}}\}, \quad (12)$$

这样对 $\forall L \in \mathcal{L}$, 均有 $E^{|\mathcal{L}|} \leq \text{Eng}(L) \leq E^1$. 因此可以把集合 $\mathcal{L}$ 划分为非空子集 $\{\mathcal{L}^i\}$,

$$\mathcal{L}^i = \{L \mid L \in \mathcal{L} \text{ and } \text{Eng}(L) = E^i\}, \quad i = 1, 2, \dots, |\mathcal{E}|, \quad (13)$$

$$\sum_{i=1}^{|\mathcal{E}|} |\mathcal{L}^i| = |\mathcal{L}|; \forall i \in \{1, 2, \dots, |\mathcal{L}|\}, \mathcal{L}^i \neq \emptyset; \forall i \neq j, \mathcal{L}^i \cap \mathcal{L}^j = \emptyset; \bigcup_{i=1}^{|\mathcal{E}|} \mathcal{L}^i = \mathcal{L}, \quad (14)$$

集合 $\mathcal{L}^1$ 包含所有能量为 E^1 的智能体网络.

令 $L^j, i = 1, 2, \dots, |\mathcal{E}|, j = 1, 2, \dots, |\mathcal{L}^i|$ 表示 $\mathcal{L}^i$ 中的第 j 个智能体网络. 在智能体进化算子的作用下, 从 L^j 到 L^{kl} 的状态转移可表示为 $L^j \rightarrow L^{kl}$. 令 $p_{ij,kl}$ 表示从 L^j 到 L^{kl} 的转移概率, $p_{ij,k}$ 表示从 L^j 到 $\mathcal{L}^k$ 中任一网络的转移概率, 而 $p_{i,k}$ 表示从 $\mathcal{L}^i$ 中任一网络到 $\mathcal{L}^k$ 中任一网络的转移概率. 显然有:

$$p_{ij,k} = \sum_{l=1}^{|\mathcal{L}^k|} p_{ij,kl}; \sum_{k=1}^{|\mathcal{E}|} p_{ij,k} = 1; p_{i,k} \geq p_{ij,k}, \quad (15)$$

明确了上述概念后, 在文献[8]的基础上, 我们将给出 MAGA 全局收敛性的理论证明.

定理 1^[9] 令 P' 是一个 n 阶可归约随机矩阵, 即通过相同的行变换和列变换后可以得到 $P' = \begin{pmatrix} C & 0 \\ R & T \end{pmatrix}$, 其中 C 是一个 m 阶本原随机矩阵且 $R, T \neq 0$. 则

$$P'^\infty = \lim_{k \rightarrow \infty} P'^k = \lim_{k \rightarrow \infty} \begin{pmatrix} C^k & 0 \\ \sum_{i=0}^{k-1} T^i R C^{k-i} & T^k \end{pmatrix} = \begin{pmatrix} C^\infty & 0 \\ R^\infty & 0 \end{pmatrix} \quad (16)$$

是一个稳定的随机矩阵且 $P'^\infty = 1' p'^\infty$, 这里 p'^∞

$= p'^0 P'^\infty$ 惟一确定且与初始分布无关, p'^∞ 满足 $p'^\infty_i > 0, 1 \leq i \leq m$ 且 $p'^\infty_i = 0, m < i \leq n$.

定理 2 在多智能体遗传算法中, 对 $\forall i, k \in \{1, 2, \dots, |\mathcal{E}|\}$, 有 $p_{i,k} = \begin{cases} > 0 & k \leq i \\ = 0 & k > i \end{cases}$.

证明 对 $\forall L^j \in \mathcal{L}^i, i = 1, 2, \dots, |\mathcal{E}|, j = 1, 2, \dots, |\mathcal{L}^i|$, 均 $\exists x^* = (x_1, \dots, x_n) \in L^j$, $\text{Eng}(x^*) = E^i$. 设在智能体进化算子的作用下, L^j 转变为 L^{kl} ; 若 L^j 为第 t 代, 则 L^{kl} 为第 $(t+1)$ 代, 为了方便分别记为 L^t 和 L^{t+1} .

首先, 由算法 2 Step 6 可得(17)式.

$$\begin{aligned} \text{Eng}(\text{CBest}^{t+1}) &= \text{Eng}(\text{Best}^{t+1}) \geq \\ \text{Eng}(\text{Best}^t) &= \text{Eng}(x^*), \end{aligned} \quad (17)$$

因此有: $\text{Eng}(L^{t+1}) \geq \text{Eng}(L^t) \Rightarrow k \leq i \Rightarrow \forall k > i, p_{ij,kl} = 0 \Rightarrow \forall k > i, p_{ij,k} = \sum_{l=1}^{|\mathcal{L}^k|} p_{ij,kl} = 0 \Rightarrow \forall k > i, p_{i,k} = 0$. (18)

其次, 在邻域竞争算子中 x^* 必然胜过其所有邻域, 因此 $x^* \in L^{t+1/3}$; 对 x^* 执行邻域正交叉算子的概率为 P_c , 因此 x^* 属于 $L^{t+2/3}$ 的概率为 $(1 - P_c)$, 而对 x^* 执行变异算子的概率 Pr_1 为 $Pr_1 = (1 - P_c) \cdot P_m > 0$.

$\exists x' \in L^{t+1}(L^{kl}), k \leq i, \text{Eng}(x') = E^k$, 不妨设 x' 有 n_1 个分量, $x'_1, \dots, x'_{n_1}$, 与 x^* 对应位置上的分量不同. 通过变异算子由 x^* 产生 x' 的概率 Pr_2 为 $Pr_2 = \left(1 - \frac{1}{n}\right)^{n-n_1} \cdot \prod_{i=1}^{n_1} \frac{\sqrt{t}}{\sqrt{2\pi}} e^{-\frac{(x_i - x'_i)^2}{2}} > 0$.

由以上分析可知, 在智能体进化算子作用后, 从 L^j 转到 $\mathcal{L}^k$ 中任一智能体网络的概率 $p_{ij,k}$ 为 $p_{ij,k} > Pr_1 \times Pr_2 > 0$. 因此, $\forall k \leq i, p_{i,k} \geq p_{ij,k} > 0$. 证毕.

该定理说明, 能量低的网格可以转移到能量相同或更高的网格, 但能量高的网格不能转移到能量低的网格, 因此一旦 MAGA 进入到集合 $\mathcal{L}^1$ 中就不可能逃逸出来.

定理 3 多智能体遗传算法具有全局收敛性.

证明 每个 $\mathcal{L}^i, i = 1, 2, \dots, |\mathcal{E}|$ 可以看作时齐有限 Markov 链上的一个状态, 根据定理 2, 该 Markov 链的转移矩阵 P' 可表示为(19)式:

$$P' = \begin{pmatrix} p_{1.1} & 0 & \cdots & 0 \\ p_{2.1} & p_{2.2} & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ p_{|E|.1} & p_{|E|.2} & \cdots & p_{|E|.|E|} \end{pmatrix} = \begin{pmatrix} C & 0 \\ R & T \end{pmatrix}, \quad (19)$$

其中 $R = (p_{2.1}, p_{3.1}, \dots, p_{|E|.1})^T > 0$, $T \neq 0$, $C = (p_{1.1}) = (1) \neq 0$.

根据定理1可得(20)式.

$$P'^{\infty} = \lim_{k \rightarrow \infty} P'^k = \lim_{k \rightarrow \infty} \begin{pmatrix} C^k & 0 \\ \sum_{i=0}^{k-1} T^i R C^{k-i} & T^k \end{pmatrix} = \begin{pmatrix} C^{\infty} & 0 \\ R^{\infty} & 0 \end{pmatrix}, \quad (20)$$

其中 $C^{\infty} = 1$, $R^{\infty} = (1, 1, \dots, 1)^T$, 这样 P'^{∞} 就是一个稳定的随机矩阵, 且 $P'^{\infty} =$

$$\begin{pmatrix} 1 & 0 & \cdots & 0 \\ 1 & 0 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 1 & 0 & \cdots & 0 \end{pmatrix}.$$

由此可得 $\lim_{l \rightarrow \infty} Pr \{Eng(L^l) = E^1\} = 1$, 其中 Pr 为概率, 所以多智能体遗传算法具有全局收敛性. 证毕.

3 数值实验与分析

在使用 MAGA 前, 需要预先设定一些参数: $L_{size} \times L_{size}$ 实际对应于遗传算法中群体的大小, 因此 L_{size} 不易选得过大或过小, 一般选 5~10 即可. P_o 实际控制了算法在利用已有信息与探索新空间之间的平衡——当 $P_o < 0.5$ 时, 侧重于探索新空间, 而当 $P_o > 0.5$ 时, 侧重于利用已有信息. P_c 一般选小于 0.5 即可, 否则计算量增加太大. P_m 相应于一般的变异概率. 智能体自学习算子实际是一个小规模 MAGA, 所以其 4 个参数也非常易于选取——从计算量方面考虑, sL_{size} 一般选小于 5, 而 $sGen$ 一般选 5~10. sR 控制搜索半径, 起到局部搜索的作用, 不易选得过大, 一般取小于 0.3. sP_m 与 P_m 类似. 本实验中 MAGA 参数设置为: $L_{size} = 5$, $P_o = 0.2$, $P_c = 0.1$, $P_m = 0.1$, $sL_{size} = 3$, $sR = 0.2$, $sP_m = 0.05$, $sGen = 10$; 停止准则为: 如果 $f_{min} \neq 0$, 则考察式子 $|f_{best} - f_{min}| < \epsilon \cdot |f_{min}|$ 是否

成立, 否则考察式子 $|f_{best}| < \epsilon$, 此处 f_{min} 和 f_{best} 分别表示全局最优解和到当前代为止找到的最优解. 实验所用的 4 个多峰标准函数如(21)~(24)式所示.

Schwefel 函数: $f_1(x) = \sum_{i=1}^n (-x_i \sin \sqrt{|x_i|})$, $S = [-500, 500]^n$, $f_{min} = -418.983 \times n$, (21)

Rastrigin 函数: $f_2(x) = \sum_{i=1}^n [x_i^2 - 10 \cos(2\pi x_i) + 10]$, $S = [-5.12, 5.12]^n$, $f_{min} = 0$, (22)

Ackley 函数: $f_3(x) = -20 \exp \left(-0.2 \sqrt{\frac{1}{n} \sum_{i=1}^n x_i^2} \right) - \exp \left(\frac{1}{n} \sum_{i=1}^n \cos(2\pi x_i) \right) + 20 + e$, $S = [-32, 32]^n$, $f_{min} = 0$, (23)

Griewank 函数: $f_4(x) = \frac{1}{4000} \sum_{i=1}^n x_i^2 - \prod_{i=1}^n \cos(x_i/\sqrt{i}) + 1$, $S = [-600, 600]^n$, $f_{min} = 0$. (24)

3.1 MAGA 与 AEA^[4]对 20~1000 维函数优化结果比较

在 AEA 的停止准则中, 对 f_1 取 $\epsilon = 10^{-4}$, 对 f_2 取 $\epsilon = 10^{-1}$, 对 f_3 和 f_4 均取 $\epsilon = 10^{-3}$. 为了一致, 在 MAGA 中, 4 个函数均取 $\epsilon = 10^{-4}$. 表 1 给出了平均函数评价次数的比较结果, 其中 MAGA 的结果为 50 次随机运行的平均. 由表 1 可见, 对于 f_1 , 当 $n \leq 100$ 时, MAGA 所用评价次数略大于 AEA, 而当 $200 \leq n \leq 1000$ 时, 略小于 AEA. 对于 f_2 , 当 $n \leq 200$ 时, MAGA 所用评价次数大于 AEA, 而当 $400 \leq n \leq 1000$ 时, 小于 AEA. 对于 f_3 和 f_4 , MAGA 所用评价次数均远远小于 AEA, 而且 MAGA 仅用几千次评价就使解的精度达到了 10^{-4} . 总的来说, 虽然 MAGA 的评价次数随着维数的增长有所增长, 但增长速度明显慢于 AEA, 显示出了其在解高维问题时的优越性能.

表 1 MAGA 与 AEA 对 20~1000 维函数优化平均函数评价次数的比较

n	f_1		f_2		f_3		f_4	
	MAGA	AEA	MAGA	AEA	MAGA	AEA	MAGA	AEA
20	2,483	1,603	4,301	1,247	3,583	7,040	2,566	3,581
100	5,443	5,106	10,265	4,789	5,410	22,710	4,447	17,228
200	7,284	8,158	14,867	10,370	6,051	43,527	5,483	36,760
400	12,368	13,822	17,939	23,588	6,615	78,216	6,249	61,975
1000	22,827	23,867	20,083	46,024	7,288	160,940	7,358	97,660

3.2 MAGA对1000~10000维函数优化的性能分析

本实验对 n 在 1000~10000 中采样, 间隔为 500. 将 MAGA 在每个采样的维数上随机运行 50 次, 图 2~5 给出了平均评价次数随维数的变化. 由图可知, 对于 f_1 , 虽然评价次数随维数的增长较明显, 但即使当维数增到 10000 时, 评价次数也少于 200000. 对于 f_2 , 维数对评价次数没有明显的影响, 尤其当维数大于 1000 时, 评价次数则分布在 10000~20000 间. 对于 f_3 , 随着维数的增长, 评价次数增长速度降低, 当维数增到 10000 时, 也只用了约 8000 次评价. 对于 f_4 , 评价次数随维数的增长速度也比较缓慢, 当维数大于 2000 时, 评价次数则分布在 10000~30000 间.

将进化算法用于上万维的函数优化问题在国内外还均未见报道, 而本实验的结果充分证明了 MAGA 对超高维函数优化的能力, 这主要得益于智能体领域间的作用和智能体的自学习能力. 在进化算法中, 评价次数是反映算法时间复杂度的主要因素, 而 MAGA 与传统遗传算法相比, 除了评价次数也没有更复杂的操作, 由此可见 MAGA 具有很快的收敛速度, 这对于实际应用具有很大价值.

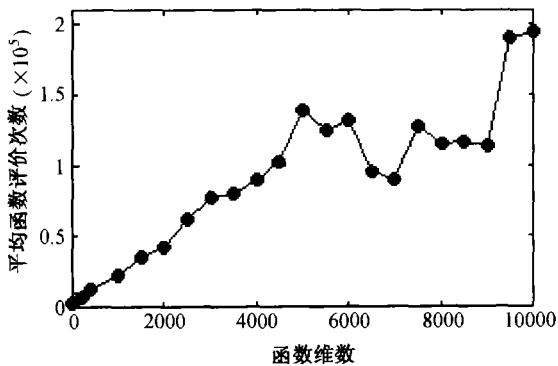


图 2 f_1 平均评价次数随维数的变化

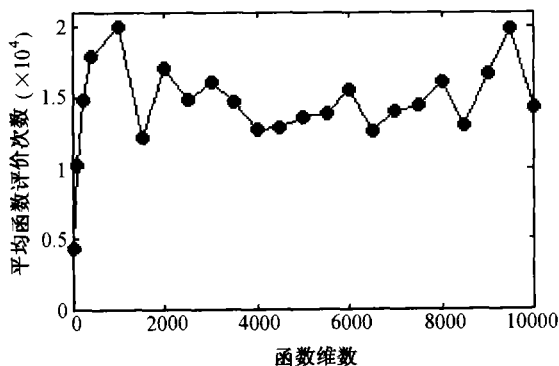


图 3 f_2 平均评价次数随维数的变化

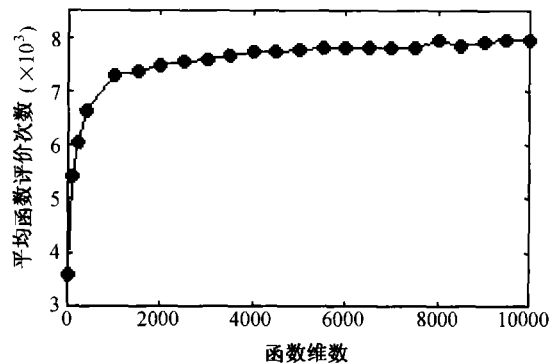


图 4 f_3 平均评价次数随维数的变化

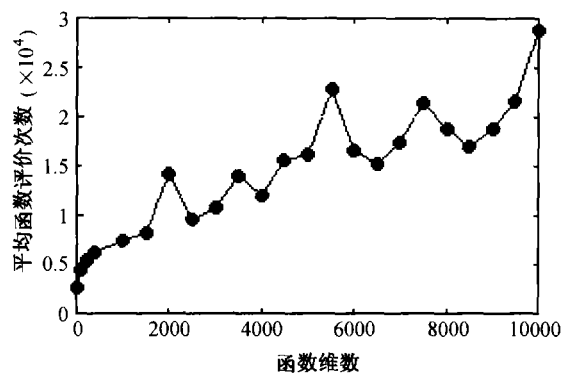


图 5 f_4 平均评价次数随维数的变化

参 考 文 献

- 1 潘 进, 等. 用遗传算法确定短序列正交多子波. 自然科学进展, 2000, 10(3): 266
- 2 李敏强, 等. 遗传算法的一种非单调适应值标度变换方法. 自然科学进展, 2001, 11(5): 530
- 3 谢 涛, 等. 基于函数分解的可伸缩宏进化算法. 自然科学进展, 2001, 11(6): 662
- 4 潘正君, 等. 演化计算. 北京: 清华大学出版社, 1998. 72~76
- 5 韩 靖, 等. AER 模型中的智能涌现. 模式识别与人工智能, 2002, 15(2): 134
- 6 Russell S, et al. Artificial Intelligence: A modern Approach. New York: Prentice-Hall, 1995
- 7 Leung Y W, et al. An orthogonal genetic algorithm with quantization for global numerical optimization. IEEE Trans on Evolutionary Computation, 2001, 5(1): 41
- 8 Rudolph G. Convergence analysis of canonical genetic algorithms. IEEE Trans on Neural Networks, Special Issue on Evolutional Computing, 1994, 5(1): 96
- 9 Iosifescu M. Finite Markov Processes and Their Applications. Chichester: Wiley, 1980