

TSP における SDPSO の最適化

日大生産工 ○紅林亜佳 日大生産工 山内 ゆかり

1. まえがき

粒子群最適化（PSO）は、鳥や魚などの群れに見られる社会行動を基にモデル化された最適解探索アルゴリズムであり、シンプルかつ高性能であり多くの分野で使われている。しかし、通常のPSOアルゴリズムでは、局所解に陥りやすいという問題がある。Zhenhua Yuらはその問題を解決すべく、SDPSO(1)を提案し、従来のGWOやPSOよりも効率の良い経路を迅速に計画することができるようになったと報告がされている。しかし、今回の実験環境である巡回セールスマン問題で探索を考えた場合、探索過程においてより良い都市ルートの可能性がある粒子同士を最適解として認識できない可能性がある。

そのため本研究では、SDPSOのさらなる精度向上を目指して、包括的な手法であるCL法（Comprehensive learning）(2)を導入し、従来よりも局所会に陥りづらく、大域最適解に収まりやすい最適化探索アルゴリズムを目指す。

2. 従来研究

2-1. 巡回セールスマン問題

巡回セールスマン問題（TSP）は都市の集合と各都市間のコスト（距離等）が与えられ、すべての都市をちょうど1度ずつ巡り出発地に戻る。その時の最小コスト（最短距離）の経路を求める、組み合わせ最適化問題である。

2-2. PSOアルゴリズム

まず、PSOアルゴリズムとは、上記で述べたように、鳥や魚などの群れに見られる社会行動を基にモデル化された最適解探索アルゴリズムである。解空間内に複数の粒子が存在し、その粒子が i で表され、速度 v_i^t 、位置 x_i^t で表され、以下の式で定義される。

$$\begin{aligned} v_i^{t+1} = & wv_i^t + c_1r_1(p_i^t - x_i^t) + \\ & c_2r_2(g^t - x_i^t) \end{aligned} \quad (1)$$

$$i=1,2,\dots,S, t=1,2,\dots,K$$

$$x_i^{t+1} = x_i^t + v_i^{t+1} \quad (2)$$

$$i=1,2,\dots,S, t=1,2,\dots,K$$

ここで、 S は粒子数、 K は最大反復回数、 w は完成重みパラメータ、 c_1, c_2 は学習係数と呼ばれる。 r_1, r_2 は0か1の乱数。 p^t, g^t は個体の集団の過去の最適位置を表す。式(1)によって速度が更新される。また、PSOの特徴として、すべての粒子群で過去の探索含め最も適合度が高い位置をグローバルベスト、1粒子における探索の中で最も適合度が高い位置をパーソナルベストとし、これを繰り返し探索の中で更新することによって最適解を探索する。

2-3. SDPSOアルゴリズム

SDPSOアルゴリズムは、PSO、SAアルゴリズム、DLS(3)の3つを組み合わせたものである。まず、PSOとは、2-2で述べたものである。一方で、SA（Simulated Annealing）アルゴリズムとは、「焼きなまし法」と呼ばれており、金属を高温の状態にして、徐々に温度を下げることで秩序がある構造を作り出す焼きなましの技術をコンピュータ上で再現したアルゴリズムである。今回の最適化問題では、局所最適解に陥ることを避けて、大域最適解に探索するように、現在の解よりも悪い解を確率 p で受け入れる。また、この確率 p は（1）式のように計算をし、 T は温度関数と呼ばれ、（2）式で計算される。

また、下記の式の t は時間（探索回数）、 K は最大探索回数、 g^t は時刻 t のグローバルベストである。これを組み合わせることにより、PSOアルゴリズムの反復中の多様性が悪化し、アルゴリズムが局所最適解に陥りやすいという問題点が解決できるのである。

$$p = \exp\left(\frac{-(Fit(g^{t+1}) - Fit(g^t))}{T^t}\right) \quad (3)$$

$$T^{t+1} = r \times T^t \quad (4)$$

Comprehensive SDPSO

Asuka KUREBAYASHI and Yukari YAMAUCHI

$$r = \frac{K-t}{K} \quad (5)$$

また、DLS (dimensional learning strategy) 方式に関しては、次元学習戦略ということで、パーソナルベストにグローバルベストを1次元ずつ当てはめていき、今のパーソナルベストより精度が良かった場合入れ替えるというものである。そして、今回のTSPでは、次元数を都市数と考えて計算する。

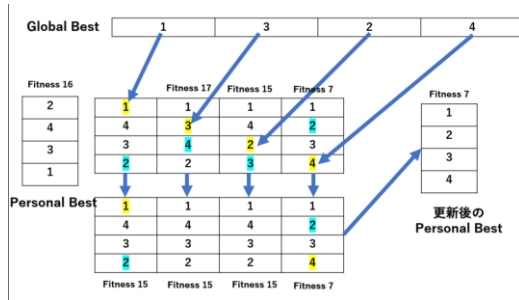


Figure1 DLS 方式イメージ図

2-4. SDPSO収束の加速

収束速度を落とさないため、DLS方式を採用する。これにより、各粒子の連続非更新数がカウントが記録されてカウントがある一定の閾値 m を超えると、大域最適解に従って各次元の現在の粒子の位置が更新される。また閾値 m の値が大きすぎると多様性がなくなってしまう為、以下の式で定義する。

$$m = \frac{t}{K} (M_{max} - M_{min}) \quad (6)$$

3. 提案手法

SDPSOでは、通常のPSOと比べて多様な探索ができ局所最適解に陥りにくくなった。しかし、より大域最適解を探索できるように包括的な探索を行う、CL法 (Comprehensive learning) を導入する。この包括的学習では、近隣の粒子の最良位置も考慮して探索を行う。これにより、探索空間における情報の幅が広がり、より多様な解を見つける能力が向上する。アルゴリズムの流れとしては、1.初期化。2.適応度の計算。3.近隣粒子の特定。各粒子に対して近隣粒子を特定する。4.位置の更新。中心粒子は、パーソナルベスト、グローバルベストに加えて近隣粒子を参考にしながら位置を更新する。5.1~4のアル

ゴリズムを反復する。これによりさらに多様な探索が可能になり、精度向上が見込める。

4. 実験および検討

実験は、以下の図のように都市数48のアメリカの都市間の距離の最短ルートを求める巡回セールスマン問題での精度 (適応度) を通常のPSO、SDPSOと比較する。

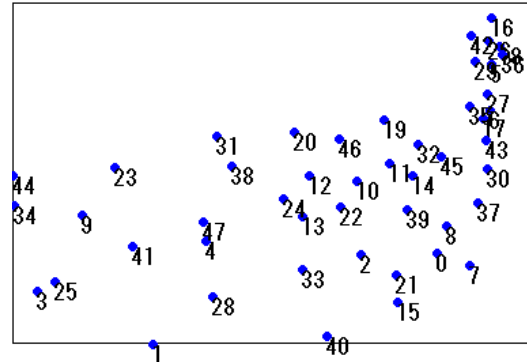


Figure 2 巡回セールスマン問題環境

その探索の中で1番適合度 (Fitness) の値が小さいものを最適解 (BestFitness) とし、その最適解の小さいほうが精度が高いと考えて、比較する。

5. まとめ

本研究では、PSOの欠点である、局所解に陥りやすいという欠点を解決するために、SAアルゴリズムとDLSアルゴリズムを加えて、多様な探索をし、大域最適解に陥りやすくさせる。さらにここに、CL方式を加えて、さらなる精度向上を目指す。

参考文献

- 1) Zhenhua Yu, Zhijie Si, Xiaobo Li, Dan Wang, Houbing Song, "Senior A Novel Hybrid Particle Swarm Optimization Algorithm for Path Planning of UAVs", IEEE,(2022)
- 2) J. J. Liang, A. K. Qin, Ponnuthurai Nagaratnam, Suganthan, Senior, and S. Baskar, "Comprehensive Learning Particle Swarm Optimizer for Global Optimization of Multimodal Functions", IEEE ,(2006)
- 3) Guiping Xu, Quanlong Cui, Xiaohu Shi, Hongwei Ge, Zhi-Hui Zhan, Heow Pueh Lee, Yanchun Liang, Ran Tai, Chunguo Wu, "Particle swarm optimization based on dimensional learning strategy", Science Direct ,(2019)