

动态规划 (Dynamic Programming)

吉建民

USTC

`jianmin@ustc.edu.cn`

2022 年 9 月 12 日

Used Materials

Disclaimer: 本课件大量采用了 Rich Sutton's RL class, David Silver's Deep RL tutorial 和其他网络课程课件, 也采用了 GitHub 中开源代码, 以及部分网络博客内容

Table of Contents

背景

策略评估 (Policy Evaluation)

策略迭代 (Policy Iteration)

值迭代 (Value Iteration)

异步 (Asynchronous) 动态规划

收敛性证明

动态规划

- ▶ 动态规划 (Dynamic Programming): 通过把原问题分解为子问题的方式求解的一类方法。
 - ▶ 通过记住求解过的子问题来节省时间
- ▶ 动态规划适用的问题需要满足两个性质:
 - ▶ 最优子结构 (Optimal Substructure): 整个问题的最优解可以通过求解子问题得到 (通过子问题的最优解构造出全局最优解)
 - ▶ Bellman's Principle of Optimality: An optimal policy has the property that whatever the initial state and initial decision are, the remaining decisions must constitute an optimal policy with regard to the state resulting from the first decision.
 - ▶ 重叠子问题 (Overlapping Subproblems):
 - ▶ 子问题多次重复出现
 - ▶ 子问题的求解结果可以储存下来并再次使用
- ▶ MDPs 满足上述性质:
 - ▶ 贝尔曼方程给出递归分解方法
 - ▶ 值函数可以作为子问题的求解结果

动态规划与强化学习

- ▶ 动态规划是强化学习算法的基础，但一般不能直接应用
 - ▶ 动态规划需要完整的 MDP 模型 (Perfect Model)，同时有较高的计算量
 - ▶ 强化学习可以看作在没有完整模型的情况下，用少量的计算尽量达到动态规划的效果
- ▶ 动态规划求解 MDPs
 - ▶ 预测 (Prediction): 计算特定策略 π 的值函数 V_π
 - ▶ 输入: MDP (S, A, T, R) 和策略 π
 - ▶ 输出: 值函数 V_π
 - ▶ 控制 (Control):
 - ▶ 输入: MDP (S, A, T, R)
 - ▶ 输出: 最优值函数 V^* 或最优策略 π^*

贝尔曼方程

► 三种写法:

$$V^*(s) = \max_a \left(R(s, a) + \gamma \sum_{s'} T(s, a, s') V^*(s') \right),$$

$$Q^*(s, a) = R(s, a) + \gamma \sum_{s'} T(s, a, s') \max_{a'} Q^*(s', a'),$$

$$C^*(s, a) = \gamma \sum_{s'} T(s, a, s') \max_{a'} (R(s', a') + C^*(s', a')).$$

► 三种表达的关系:

$$V^*(s) = \max_a Q^*(s, a)$$

$$V^*(s) = \max_a (R(s, a) + C^*(s, a))$$

$$Q^*(s, a) = R(s, a) + C^*(s, a)$$

$$Q^*(s, a) = R(s, a) + \gamma \sum_{s'} T(s, a, s') V^*(s')$$

$$C^*(s, a) = \gamma \sum_{s'} T(s, a, s') V^*(s')$$

$$C^*(s, a) = \gamma \sum_{s'} T(s, a, s') \max_{a'} Q^*(s', a')$$

Table of Contents

背景

策略评估 (Policy Evaluation)

策略迭代 (Policy Iteration)

值迭代 (Value Iteration)

异步 (Asynchronous) 动态规划

收敛性证明

迭代策略评估 (Iterative Policy Evaluation)

- ▶ 策略评估：给定策略 π ，计算值函数 V_π
 - ▶ $V_\pi(s) = E_\pi (R_{t+1} + \gamma R_{t+2} + \gamma^2 R_{t+3} + \dots \mid S_t = s)$
 - ▶ $V_\pi(s) = \sum_a \pi(a \mid s) (R(s, a) + \gamma \sum_{s'} T(s, a, s') V_\pi(s'))$
- ▶ 迭代策略评估：给定策略 π ，采用迭代的方式基于贝尔曼方程计算出其值函数 V_π
 - ▶ $V_1 \rightarrow V_2 \rightarrow \dots \rightarrow V_\pi$
- ▶ 同步迭代算法：
 - ▶ 在第 $k+1$ 次迭代中，对所有状态 $s \in S$ ，基于 $V_k(s')$ 更新 $V_{k+1}(s)$ ，直到收敛
 - ▶ $V_\pi^{k+1}(s) = \sum_a \pi(a \mid s) (R(s, a) + \gamma \sum_{s'} T(s, a, s') V_\pi^k(s'))$

迭代策略评估算法

Iterative policy evaluation

Input π , the policy to be evaluated

Initialize an array $V(s) = 0$, for all $s \in \mathcal{S}^+$

Repeat

$\Delta \leftarrow 0$

 For each $s \in \mathcal{S}$:

$v \leftarrow V(s)$

$V(s) \leftarrow \sum_a \pi(a|s) \sum_{s',r} p(s',r|s,a) [r + \gamma V(s')]$

$\Delta \leftarrow \max(\Delta, |v - V(s)|)$

until $\Delta < \theta$ (a small positive number)

Output $V \approx v_\pi$

策略评估直接求解

- ▶ 给定特定策略 π , MDP 变成 Markov Reward Process (MRP)

$$\begin{aligned} V_{\pi}(s) &= \sum_a \pi(a | s) \left(R(s, a) + \gamma \sum_{s'} T(s, a, s') V_{\pi}(s') \right) \\ &= \sum_a \pi(a | s) R(s, a) + \gamma \sum_a \pi(a | s) \sum_{s'} T(s, a, s') V_{\pi}(s') \\ &= R_s^{\pi} + \gamma \sum_{s'} T_{s's}^{\pi} V_{\pi}(s') \end{aligned}$$

其中 $R_s^{\pi} = \sum_a \pi(a | s) R(s, a)$, $T_{s's}^{\pi} = \sum_a \pi(a | s) T(s, a, s')$.

- ▶ 将所有状态 $s \in S$ 表达为向量, 上式的矩阵形式为:

$$V_{\pi} = R^{\pi} + \gamma T^{\pi} V_{\pi}$$

直接求解结果为:

$$V_{\pi} = (I - \gamma T^{\pi})^{-1} R^{\pi}$$

计算复杂性为 $O(N^3)$

Table of Contents

背景

策略评估 (Policy Evaluation)

策略迭代 (Policy Iteration)

值迭代 (Value Iteration)

异步 (Asynchronous) 动态规划

收敛性证明

策略迭代

- ▶ 给定策略 π
 - ▶ 评估 (Evaluate) 策略 π , 计算 V_π
 - ▶ 迭代策略评估

$$V_\pi^{k+1}(s) = \sum_a \pi(a | s) \left(R(s, a) + \gamma \sum_{s'} T(s, a, s') V_\pi^k(s') \right)$$

- ▶ 直接策略评估

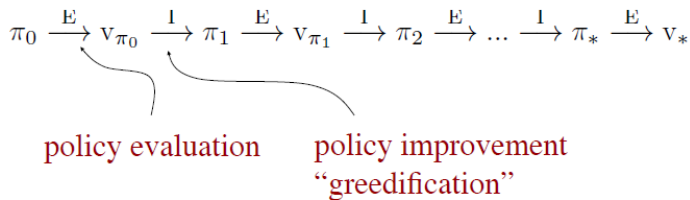
$$V_\pi = (I - \gamma T^\pi)^{-1} R^\pi$$

- ▶ 改进 (Improve) 策略, 基于 V_π 采用贪婪策略
 - ▶ $Q_\pi(s, a) = R(s, a) + \gamma \sum_{s'} T(s, a, s') V_\pi(s')$
 - ▶ 基于 $Q_\pi(s, a)$ 采用贪婪策略得到 π'

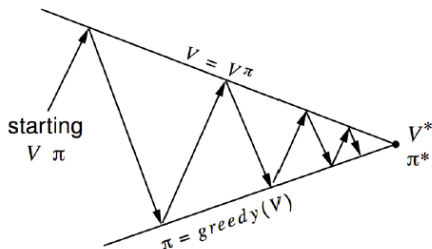
$$\begin{aligned} \pi'(s) &= \operatorname{argmax}_a Q_\pi(s, a) \\ &= \operatorname{argmax}_a \left(R(s, a) + \gamma \sum_{s'} T(s, a, s') V_\pi(s') \right) \end{aligned}$$

- ▶ 当策略不再改进, 则已经是最优的。

策略迭代 (con't)



策略迭代 (con't)

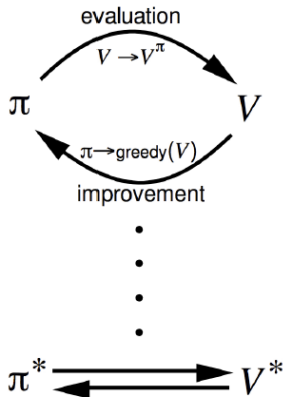


Policy evaluation Estimate v_π

Iterative policy evaluation

Policy improvement Generate $\pi' \geq \pi$

Greedy policy improvement



策略迭代算法

1. Initialization

$V(s) \in \mathbb{R}$ and $\pi(s) \in \mathcal{A}(s)$ arbitrarily for all $s \in \mathcal{S}$

2. Policy Evaluation

Repeat

$\Delta \leftarrow 0$

For each $s \in \mathcal{S}$:

$v \leftarrow V(s)$

$V(s) \leftarrow \sum_{a \in \mathcal{A}} \pi(a|s) (r(s, a) + \gamma \sum_{s' \in \mathcal{S}} T(s'|s, a) V(s'))$

$\Delta \leftarrow \max(\Delta, |v - V(s)|)$

until $\Delta < \theta$ (a small positive number)

3. Policy Improvement

policy-stable $\leftarrow true$

For each $s \in \mathcal{S}$:

$a \leftarrow \pi(s)$

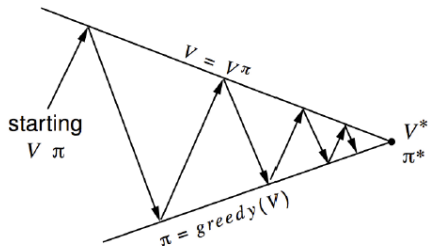
$\pi(s) \leftarrow \arg \max_a r(s, a) + \gamma \sum_{s' \in \mathcal{S}} T(s'|s, a) v_{\pi(s')}$

If $a \neq \pi(s)$, then *policy-stable* $\leftarrow false$

If *policy-stable*, then stop and return V and π ; else go to 2

广义策略迭代 (Generalized Policy Iteration)

- Generalized Policy Iteration (GPI): any interleaving of policy evaluation and policy improvement, independent of their granularity.



Policy evaluation Estimate v_π

Any policy evaluation algorithm

Policy improvement Generate $\pi' \geq \pi$

Any policy improvement algorithm

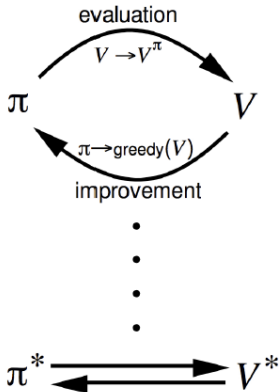


Table of Contents

背景

策略评估 (Policy Evaluation)

策略迭代 (Policy Iteration)

值迭代 (Value Iteration)

异步 (Asynchronous) 动态规划

收敛性证明

值迭代 (Value Iteration)

- ▶ 当已知所有子问题的解, $V^*(s')$, 则最终结果 $V^*(s)$ 可以通过下面公式计算

$$V^*(s) \leftarrow \max_a \left(R(s, a) + \gamma \sum_{s'} T(s, a, s') V^*(s') \right)$$

- ▶ 值迭代算法就是将上述公式转换为迭代形式:

$$V_{k+1}^*(s) \leftarrow \max_a \left(R(s, a) + \gamma \sum_{s'} T(s, a, s') V_k^*(s') \right)$$

- ▶ 基于 $V^*(s)$ 可以计算出最优策略 π^*

$$\pi^*(s) = \operatorname{argmax}_a \left(R(s, a) + \gamma \sum_{s'} T(s, a, s') V^*(s') \right)$$

值迭代算法

Value iteration

Initialize array V arbitrarily (e.g., $V(s) = 0$ for all $s \in \mathcal{S}^+$)

Repeat

$\Delta \leftarrow 0$

For each $s \in \mathcal{S}$:

$v \leftarrow V(s)$

$V(s) \leftarrow \max_a \sum_{s',r} p(s',r|s,a)[r + \gamma V(s')]$

$\Delta \leftarrow \max(\Delta, |v - V(s)|)$

until $\Delta < \theta$ (a small positive number)

Output a deterministic policy, $\pi \approx \pi_*$, such that

$\pi(s) = \operatorname{argmax}_a \sum_{s',r} p(s',r|s,a)[r + \gamma V(s')]$

策略迭代与值迭代的区别

```
1. Initialization
    $v(s) \in \mathbb{R}$  and  $\pi(s) \in \mathcal{A}(s)$  arbitrarily for all  $s \in \mathcal{S}$ 

2. Policy Evaluation
   Repeat
      $\Delta \leftarrow 0$ 
     For each  $s \in \mathcal{S}$ :
        $temp \leftarrow v(s)$ 
        $v(s) \leftarrow \sum_{s'} p(s'|s, \pi(s)) [r(s, \pi(s), s') + \gamma v(s')]$ 
        $\Delta \leftarrow \max(\Delta, |temp - v(s)|)$ 
   until  $\Delta < \theta$  (a small positive number)

3. Policy Improvement
    $policy\_stable \leftarrow true$ 
   For each  $s \in \mathcal{S}$ :
      $temp \leftarrow \pi(s)$ 
      $\pi(s) \leftarrow \arg \max_a \sum_{s'} p(s'|s, a) [r(s, a, s') + \gamma v(s')]$ 
     If  $temp \neq \pi(s)$ , then  $policy\_stable \leftarrow false$ 
   If  $policy\_stable$ , then stop and return  $v$  and  $\pi$ ; else go to 2
```

Figure 4.3: Policy iteration (using iterative policy evaluation) for v_* . This algorithm has a subtle bug, in that it may never terminate if the policy continually switches between two or more policies that are equally good. The bug can be fixed by adding additional flags, but it makes the pseudocode so ugly that it is not worth it. :-)

finding optimal
value function

```
Initialize array  $v$  arbitrarily (e.g.,  $v(s) = 0$  for all  $s \in \mathcal{S}^+$ )

Repeat
   $\Delta \leftarrow 0$ 
  For each  $s \in \mathcal{S}$ :
     $temp \leftarrow v(s)$ 
     $v(s) \leftarrow \max_a \sum_{s'} p(s'|s, a) [r(s, a, s') + \gamma v(s')]$ 
     $\Delta \leftarrow \max(\Delta, |temp - v(s)|)$ 
  until  $\Delta < \theta$  (a small positive number)

Output a deterministic policy,  $\pi$ , such that
   $\pi(s) = \arg \max_a \sum_{s'} p(s'|s, a) [r(s, a, s') + \gamma v(s')]$ 
```

one policy
update (extract
policy from the
optimal value
function

Figure 4.5: Value iteration.

策略迭代与值迭代的区别 (con't)

- ▶ 策略迭代的第二步 policy evaluation 与值迭代的第二步 finding optimal value function 十分相似，除了后者用了 max 操作，前者没有 max。因此后者可以得出 optimal value function，而前者不能得到 optimal function。
- ▶ 策略迭代分为两步：策略估值和策略提升，策略估值时必须使得状态值函数迭代到收敛才能进入下一步进行策略提升，而值迭代时，其实也是策略估值和策略提升，只是由于策略估值简化了才使得两步合并起来了，具体就是在策略估值时状态值函数只迭代一次，不需要收敛，因此在当前某个状态下，只需要对每个可能的动作都计算一下采取这个动作后到达的下一个状态的期望价值，然后直接选择其中一个最大的值作为当前状态的估值。
- ▶ 策略迭代的收敛速度更快一些，在状态空间较小时，最好选用策略迭代方法（先让状态值函数收敛，再更新策略）。当状态空间较大时，值迭代的计算量更小一些（状态空间太大时，如果采用策略迭代，策略估值时每次都要迭代到收敛太慢了）。

同步 (Synchronous) 动态规划算法

Problem	Bellman Equation	Algorithm
Prediction	Bellman Expectation Equation	Iterative Policy Evaluation
Control	Bellman Expectation Equation + Greedy Policy Improvement	Policy Iteration
Control	Bellman Optimality Equation	Value Iteration

- ▶ 迭代策略评估每步迭代的时间复杂性为 $O(mn^2)$, 其中 m actions, n states
- ▶ 策略迭代每步迭代的时间复杂性为 $O(mn^2 + n^3)$, 其中 policy improvement 用 $O(mn^2)$, policy evaluation 用 $O(n^3)$ (求解线性方程的时间复杂度)
- ▶ 值迭代每步迭代的时间复杂性为 $O(mn^2)$
- ▶ 上述算法采用 state-value function $V_\pi(s)$ 或 $V^*(s)$, 也可以采用 action-value function $Q_\pi(s, a)$ 或 $Q^*(s, a)$, 值迭代每步的时间复杂性提高为 $O(m^2n^2)$

Table of Contents

背景

策略评估 (Policy Evaluation)

策略迭代 (Policy Iteration)

值迭代 (Value Iteration)

异步 (Asynchronous) 动态规划

收敛性证明

异步 (Asynchronous) 动态规划

- ▶ 同步动态规划：将所有状态同时更新，一次更新整个状态空间
 - ▶ 当状态空间很大时，计算代价较大
 - ▶ 同步动态规划，目前能处理百万级别状态空间的问题
- ▶ 异步动态规划：选择状态，分别更新
 - ▶ 可以避免更新一些与最优策略无关的状态（比如，利用 Agent 的历史经验）
 - ▶ 当所有状态被持续选择，则保证收敛
- ▶ 异步动态规划算法：
 - ▶ 原位动态规划 (In-place dynamic programming)
 - ▶ 优先扫描 (Prioritized sweeping)
 - ▶ 实时动态规划 (Real-time dynamic programming)

原位动态规划 (In-place dynamic programming)

- ▶ 同步值迭代存储两份值函数

$$V_{new}(s) \leftarrow \max_a \left(R(s, a) + \gamma \sum_{s'} T(s, a, s') V_{old}(s') \right)$$
$$V_{old}(s) \leftarrow V_{new}$$

- ▶ 原位动态规划只存储一份，随时更新

$$V(s) \leftarrow \max_a \left(R(s, a) + \gamma \sum_{s'} T(s, a, s') V(s') \right)$$

- ▶ 原位动态规划即时计算即时更新。这样可以减少保存的状态价值的数量，节约内存。代价是收敛速度可能稍慢。

优先扫描 (Prioritized sweeping)

- ▶ 优先级动态规划 (prioritised sweeping): 该算法对每一个状态进行优先级分级, 优先级越高的状态其状态价值优先得到更新。通常使用贝尔曼误差来评估状态的优先级, 贝尔曼误差即新状态价值与前次计算得到的状态价值差的绝对值。这样可以加快收敛速度, 代价是需要维护一个优先级队列。
- ▶ 贝尔曼误差 (Bellman Error)

$$\left| \max_a \left(R(s, a) + \gamma \sum_{s'} T(s, a, s') V(s') \right) - V(s) \right|$$

实时动态规划 (Real-time dynamic programming)

- ▶ 实时动态规划 (real-time dynamic programming): 实时动态规划直接使用个体与环境交互产生的实际经历来更新状态价值, 对于那些个体实际经历过的状态进行价值更新。这样个体经常访问过的状态将得到较高频次的价值更新, 而与个体关系不密切、个体较少访问到的状态其价值得到更新的机会就较少。收敛速度可能稍慢。
- ▶ 基本思想: 只更新与 Agent 有关的状态 (可能经历到的状态)
- ▶ 每步实际状态为 S_t , 行动为 A_t , 效用为 r_{t+1} , 每步状态更新为:

$$V(S_t) \leftarrow \max_a \left(R(S_t, a) + \gamma \sum_{s'} T(S_t, a, s') V(s') \right)$$

近似动态规划 (Approximate Dynamic Programming)

- ▶ 计算近似的值函数 $\hat{V}(s, w)$
 - ▶ 通过机器学习方法，根据采样状态集 $\tilde{S}$ ，估计值函数

$$\tilde{V}_k(s) = \max_a \left(R(s, a) + \gamma \sum_{s'} T(s, a, s') \hat{V}(s', w_k) \right)$$

- ▶ 再利用数据集 $\{\langle s, \tilde{V}_k(s) \rangle \mid s \in \tilde{S}\}$ ，学习得到新的值函数 $\hat{V}(\cdot, w_{k+1})$
- ▶ 利用近似的值函数 $\hat{V}(\cdot, w)$ 进行动态规划求解

Table of Contents

背景

策略评估 (Policy Evaluation)

策略迭代 (Policy Iteration)

值迭代 (Value Iteration)

异步 (Asynchronous) 动态规划

收敛性证明

算法收敛性

- ▶ 迭代策略评估算法的结果是否收敛到 V_π ?
- ▶ 策略迭代算法的结果是否收敛到 π^* ?
- ▶ 值迭代算法的结果是否收敛到 V^* ?
- ▶ 算法是否一定收敛, 结果是否唯一, 收敛速度多快 ?
- ▶ 收缩映射定理 (Contraction Mapping Theorem) 可以回答这些问题

收缩映射定理 (Contraction Mapping Theorem)

- ▶ 范数 (norm): $\|\vec{x}\|$ 表示向量在向量空间中的长度
 - ▶ $\|\vec{x}\|_1 := \sum_{i=1}^n |x_i|$
 - ▶ $\|\vec{x}\|_2 := \sqrt{\sum_{i=1}^n x_i^2}$
 - ▶ $\|\vec{x}\|_p := (\sum_{i=1}^n |x_i|^p)^{1/p}$
 - ▶ $\|\vec{x}\|_\infty := \max_i |x_i|$
- ▶ An operator F on a normed vector space $\mathcal{X}$ is a γ -**contraction**, for $0 < \gamma < 1$, provided for all $x, y \in \mathcal{X}$

$$\|F(x) - F(y)\| \leq \gamma \|x - y\|$$

Theorem (Contraction Mapping Theorem)

For any metric space $\mathcal{V}$ that is complete (i.e. closed) under an operator $T(v)$, where T is a γ -contraction,

- *T converges to a unique fixed point*
- *At a linear convergence rate of γ*

值函数空间

- ▶ 考虑由所有可能值函数构成的向量空间 $\mathcal{V}$
- ▶ 其中每个向量 V 都是 $|S|$ 维
- ▶ 向量空间中每个点对应一个值函数 $V(\cdot)$
- ▶ 贝尔曼操作作为向量空间 $\mathcal{V}$ 中的操作
 - ▶ 贝尔曼期望操作 (Bellman expectation backup operator) F^π

$$F^\pi(V) = R^\pi + \gamma T^\pi V$$

- ▶ 贝尔曼最优操作 (Bellman optimality backup operator) F^*

$$F^*(V) = \max_a (R(a) + \gamma T(a)V)$$

- ▶ 两个值函数 U 和 V 的距离由 ∞ -norm 定义

$$\|U - V\|_\infty = \max_s |U(s) - V(s)|$$

- ▶ 这些贝尔曼操作都使得值函数更靠近 (closer), 从而收敛到唯一解

贝尔曼操作

- ▶ 贝尔曼期望操作是 γ -contraction

$$\begin{aligned}\|F^\pi(U) - F^\pi(V)\|_\infty &= \|(R^\pi + \gamma T^\pi U) - (R^\pi + \gamma T^\pi V)\|_\infty \\ &= \|\gamma T^\pi(U - V)\|_\infty \\ &\leq \|\gamma T^\pi\| \|U - V\|_\infty \\ &\leq \gamma \|U - V\|_\infty\end{aligned}$$

- ▶ 贝尔曼最优操作是 γ -contraction

$$\|F^*(U) - F^*(V)\|_\infty \leq \gamma \|U - V\|_\infty$$

收敛性结果

- ▶ 贝尔曼期望操作 F^π 和贝尔曼最优操作 F^* 各自有唯一的不动点
- ▶ V_π 是贝尔曼期望操作 F^π 的不动点
- ▶ 迭代策略评估算法收敛到 V_π ，策略迭代算法收敛到 V^*
- ▶ V^* 是贝尔曼最优操作 F^* 的不动点
- ▶ 值迭代算法收敛到 V^*

小结

- ▶ 动态规划算法求解给定模型的 MDP 问题，是强化学习算法的基础
- ▶ 动态规划算法的核心是贝尔曼方程
- ▶ 策略评估算法，给定策略计算其相应的值函数
- ▶ 策略迭代算法，先评估策略，再改进策略，直到收敛
- ▶ 值迭代算法，采用贝尔曼最优操作不断迭代，直到收敛