

CS181 期末复习

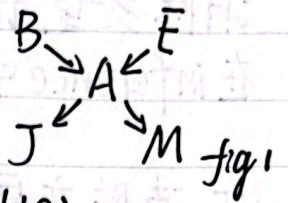
BN Exact Inference

BN is directed acyclic graph. It implicitly encode 联合分布.

$$P(x_1, \dots, x_n) = \prod_{i=1}^n P(x_i | \text{parents}(x_i))$$

Eg. $P(+b, +e, +a, +j, +m)$ (fig1)

$$= P(+b) P(+e) P(+a | +b, +e) P(+j | +a) P(+m | +a)$$



Outline: Enumeration. Variable Elimination. NPC. Sampling

① By enumeration: evidence: $E_1, \dots, E_k = e_1, \dots, e_k$

Query: Q Hidden variables: $H_1, \dots, H_r$

欲求 $P(Q | e_1, \dots, e_k)$, 实际求: $\text{sum out } H \text{ to get joint of Query}$ and Evidence

$$= \sum_{h_1, \dots, h_r} P(Q, h_1, \dots, h_r, e_1, \dots, e_k) \times \frac{1}{Z} \text{ normalize}$$

如 fig1 中: $P(B | +j, +m) = \frac{1}{Z} \sum_{e, a} P(B) P(e) P(a | B, e) P(+j | a) P(+m | a)$

这个式子完全正确, 但计算 Σ 中大量重复项, 不高效

② Variable Elimination: 两步:

a. 把求和操作尽可能 move inward

b. 之后, 从内至外作计算

$$P(B | j, m) = \alpha \sum_{e, a} P(B) P(e) P(a | B, e) P(j | a) P(m | a)$$

$$= \alpha P(B) \sum_e P(e) \sum_a P(a | B, e) P(j | a) P(m | a)$$

思路: 固定其余的 (e), 让一个 sum (a)

之后计算流程: $\text{sum over } a, \text{ then sum over } e$

注意到 $P(a | B, e)$, 启示: 不应以数为单位, 而应以 Factor

Def: Factor: multi-dimensional array to represent $P(Y_1 \dots Y_n | X_1 \dots X_n)$



$R \rightarrow T$

Operation 1: Join Factors: $P(R) \times P(T|R) \Rightarrow P(R, T)$

2: Eliminate: $P(R, T) \xrightarrow{\text{sum } R} P(T)$

消除变量法依然在 inference 问题上 NP-hard, 但比 enum 快多了
在 inference 中, 把除了 q_i 与 e 的其它 Hidden variables

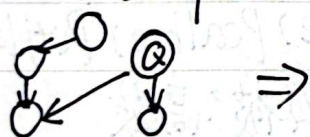
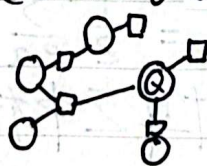
VE 消除变量顺序也十分重要, *can greatly affect the size of the largest factor*

△: 是否有算法总能返回产生最小 factors 的 order? 不存在!

这是个 NP-Hard 问题, 可从 3SAT 归约至此问题

但是 对于 poly-tree* BN 这种特殊 BN, VE 复杂度 is linear in the BN size if 采取如下策略: *leaves towards root*

① convert to factor graph ② take Q as root ③ Eliminate from


 $\Rightarrow$


Factor Graph

* Poly-tree: A directed graph with no undirected cycles.

Approximate Inference

Sampling: Goal: P Basic idea: 在采样分布 S 下抽 N 个 sample, 进而计算 quantity, 其收敛至 P
此方法快而简单, 且内存开销少 ($O(n)$)

Prior Sampling: For $i=1, 2, \dots, n$ (topological order!)

Sample X_i from $P(X_i | \text{parents}(X_i))$

最后样本的频率代表 BN's 联合分布概率!

i.e., the sampling procedure is consistent

即: 与 evidence 不一致的采样

Reject Sampling: 在 prior 基础上微改:

e.g., 欲近似 $P(C=a, b)$, 则在 prior 采样过程中, 若 $[A \neq a \text{ or } B \neq b]$, 立即拒绝该样本

It is consistent for conditional probabilities !!

Likelihood Weighting: reject 采样的问题在于: 若 evidence 发生概率很低, 那么便会拒掉很多 sample.

Algorithm: Input: evidence $e_1, \dots, e_k$

$w = 1.0$

for $i=1, \dots, n$

if X_i is an evidence variable

$x_i = e_i$ (observed)

$w = w * P(x_i | \text{parents}(X_i))$

else: Sample x_i from $P(X_i | \text{parents}(X_i))$

return $(X_1, X_2, \dots, X_n), w$

下游 Var 受上游影响, 但上游不受下游影响!

It is consistent for conditional probabilities !!



Gibbs Sampling: Generate each sample by making a random change to preceding sample

固定 evidence, 对剩下的 non-evidence var, sample its value conditioned on all other var, i.e., var in Markov Blanket.

随机游

流程: ↓ fix evidence → initialize (Randomly)

→ 选一个非 evidence var, 在 Markov Blanket 分布下采样它

→ 不断重复上一个操作

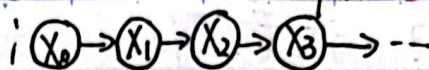
△ Theorem: Gibbs sampling is consistent

综上: prior sampling 针对任意 probability, 而 Reject, Likelihood 与 Gibbs sampling 是针对固定的一个条件概率的



Probabilistic Temporal Models

Markov Models:



Value of X : **state**

Stationary Assumption: $P(X_t | X_{t-1})$ 不变! $P(X_0)$

$P(X_t | X_{t-1}) \rightarrow$ **Transition Model**

$$P(X_0, \dots, X_T) = P(X_0) \prod_t P(X_t | X_{t-1})$$

Δ : '可'也'不可'视之为BN! '可'是因为它是DAG, '不可'是因为它无限多变量在Model中, X_{t+1} 与 $X_0, \dots, X_{t-1}$ 是条件独立的 given X_t ! 即只与上一时间步有关, 故称 **first-order Markov Model**

(k -order allows dependency on k earlier steps)

Forward Algorithm: 推理 t 时 X_t 分布 (given $P(X_0)$)

$$P(X_t) = \sum_{x_{t-1}} P(X_t, X_{t-1} = x_{t-1}) = \sum_{x_{t-1}} \underbrace{P(x_{t-1})}_{\text{Probability of previous iteration}} \underbrace{P(X_t | X_{t-1} = x_{t-1})}_{\text{Transition model}}$$

考场上灵机应变作 vectorization

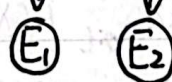
观察到: 对于大多数 chains: initial 分布影响 $\rightarrow$ end up with a distribution independent of initial distribution $\Rightarrow$ stationary 分布

$$P_{\infty}(X) = P_{\infty+1}(X) = \sum_x P(X|x) P_{\infty}(x)$$

Hidden Markov Models: $X \rightarrow E$ (Evidence)

New! Emission Model: $P(E_t | X_t)$ $X_0 \rightarrow X_1 \rightarrow X_2 \rightarrow \dots$

$P(X_0, X_1, E_1, \dots, X_T, E_T)$



$$= P(X_0) \prod_{t=1:T} \underbrace{P(X_t | X_{t-1})}_{\text{先 transit 再 emit}} \underbrace{P(E_t | X_t)}$$

Inductive Bias: E_t 与其它所有条件独立 given X_t

Notation: $X_{a:b} = X_a, X_{a+1}, \dots, X_b$



Task: Filter: $P(x_t | e_{1:t})$: 在给定 $e_{1:t}$, 推 x_t 后验分布
如何求它? 注意到:

$$P(x_{t+1} | e_{1:t+1}) = P(x_{t+1} | e_{1:t}, e_{t+1}) \quad \text{Bayes' Rule}$$

$$= \alpha P(e_{t+1} | x_{t+1}, e_{1:t}) P(x_{t+1} | e_{1:t})$$

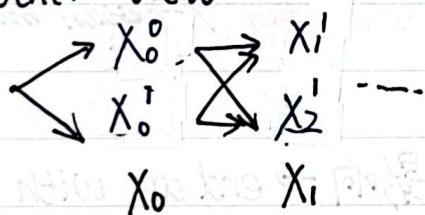
$$\text{normalize} = \alpha P(e_{t+1} | x_{t+1}) [P(x_{t+1} | e_{1:t})] \rightarrow \text{predict.}$$

$$= \underbrace{\alpha P(e_{t+1} | x_{t+1})}_{\text{emit}} \underbrace{\sum_{x_t} P(x_t | e_{1:t})}_{\hookrightarrow P(x_t | e_{1:t})} \underbrace{P(x_{t+1} | x_t)}_{\text{transit}}$$

Cost Per Time Step: $O(|X|^2)$ where $|X|$ is the number of states

Trick: α_s are constant! 只需最后一步时 normalize 即可。

Another view:



每个 arc 代表 $x_{t+1} \rightarrow x_t$, 权重视为:

$$P(x_t | x_{t+1}) \cdot P(e_t | x_t)$$

Each path $\rightarrow$ a sequence of states

$\Rightarrow$ Path 上 weight 的乘积 is proportional to that state sequence's probability. 故 all path 便是在 filter:

$$P(x_{t+1} | e_{1:t+1}) = \sum_{x_{1:t}} P(x_{1:t+1} | e_{1:t+1})$$

可用 dynamic programming. 除了 Filter, 有上述 view 还可:

Task: Most Likely Explanation: $\arg\max_{x_{0:t}} P(x_{0:t} | e_{1:t})$

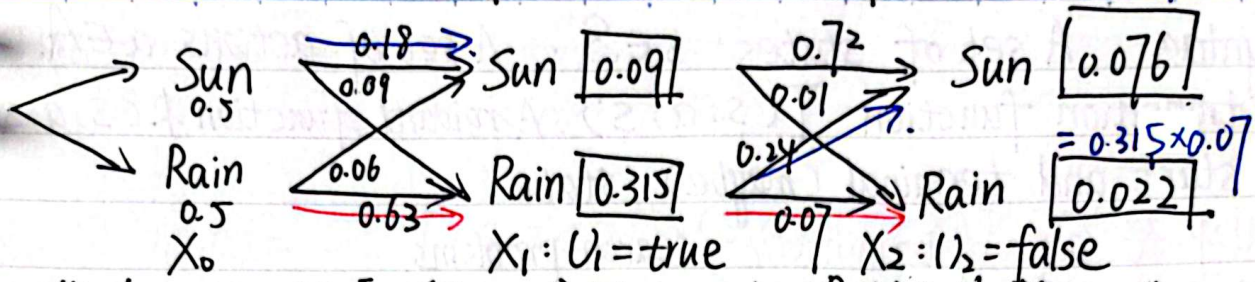
Viterbi 算法: $\arg\max_{x_{0:t}} P(x_{0:t} | e_{1:t}) = \arg\max_{x_{0:t-1}} P(x_{0:t-1} | e_{1:t-1}) \cdot P(x_t | x_{t-1}) \cdot P(e_t | x_t)$ $\hookrightarrow$ viterbi 中项

$\Rightarrow$ 知 Π_{t-1} 情况, 则在 $\{P(x_t | x_{t-1}) P(e_t | x_t)\}$ 找最大项即其对应 x_t . 注意过程中, 推 x_{t+1} 时, 要用到:

Π_i 情况 & $x_i \rightarrow x_{i+1}$ arc's weight (而此 weight 要用 e_i), 并记录路径.



Eg.	W_{t-1}	$P(W_t W_{t-1})$		W_t	$P(W_t W_t)$	
		Sun	Rain		true	false
	Sun	0.9	0.1	Sun	0.2	0.8
	Rain	0.3	0.7	Rain	0.9	0.1

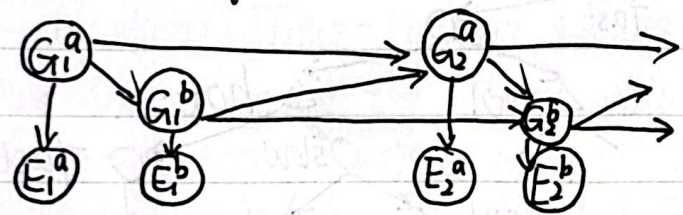


蓝红 path 都需留! 推至 t 时, 谁代表的数大, 选哪个 x_t , 然后按照 path 倒推回 $x_{0:t-1}$

Time & Space Complexity: $O(|X|^2 T)$ & $O(|X| T)$

Dynamic Bayes Net (DBN)

Track Multiple Hidden Variables via Multiple Evidence



Observation: 算每个 particle 的权重从 1 变为多少; (归一化)
Resample: 从 N 个 sample 中的 weight 为概率, 放回地抽 N 次

每个 HMM 都是 DBN, 且每个离散 DBN 也能用一个 HMM 表达

Inference in DBN: Particle Filtering* 近似解

(Track samples of X , not all values; Samples are called particles)

则 $P(X)$ 实际上是用 a list of N particles 表示的

一般, $N \ll |X|$. $N \uparrow$, 越精确

X_0 时, particles 在 $P(X_0)$ 下分布。之后, 每个 sample 在 $P(X_{t+1} | x_t)$ 下采样, 并予以 $Weight = P(e_t | x_t)$. t 步后, 有 x_t 样本以及 weight, 加和便是 X_t 分布 (normalize), 便可用于 $t+1$ 时 inference. 很像 likelihood weighting!

每个 sample 在 $P(X_{t+1} | x_t)$ 下 propagate 采样, $Weight = P(e_t | x_t)$. 用此时的加权样本 ~~在~~ $P(X_{t+1})$. Resample N 个样本且 weight 为 1, 然后 propagate...

很像 likelihood weighting, 但它对于 HMM 不适合, 因为 sample 权重掉的太快; 第 t 步 resample 后, 一个 state 中样本数量占比代表 t 时该 state 的概率分布

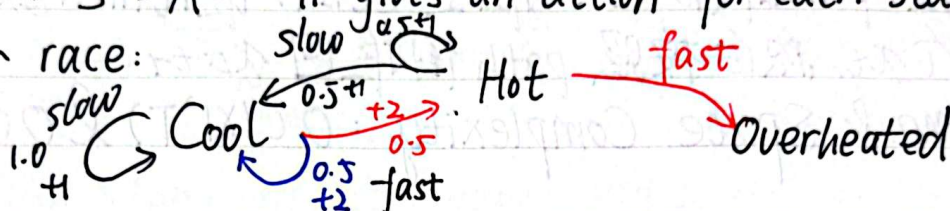
Markov Decision Process

Syntax: A set of states $s \in S$. A set of actions $a \in A$.
 A transition function $T(s, a, s')$. A reward function $R(s, a, s')$.
 A start and terminal (maybe) state

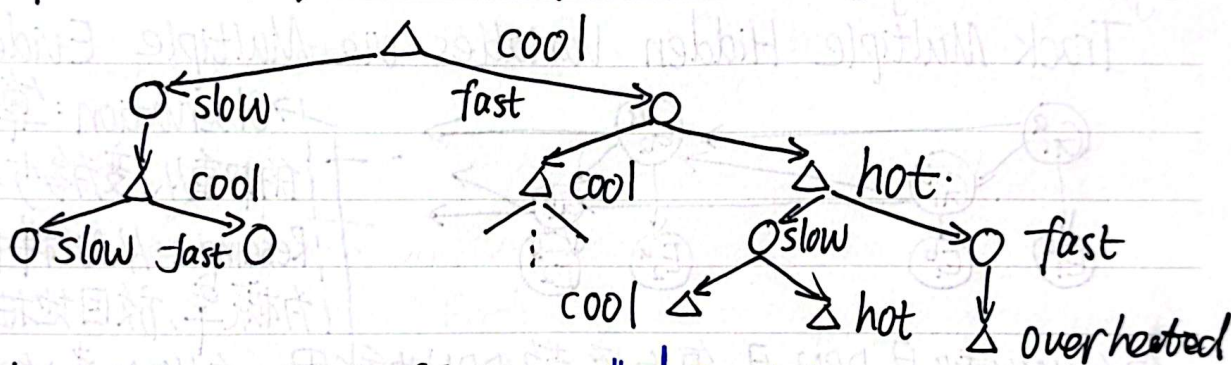
MDP 是 non-deterministic search problems.

Policy: $\pi^*: S \rightarrow A$: π gives an action for each state

Example: car race:



故: MDP state projects an expectimax-like search tree



Δ : state, s $\circ$: action, (s, a) , called q_h -state

$(s, a) \rightarrow s'$ is called a transition
 $T(s, a, s') = p(s' | s, a)$, 且有 $R(s, a, s')$

对于每个 s 时, 看眼前与未来的收益都很重要! - 一种权衡法是 Discounting, How? Level $\uparrow$, 就给那个收益乘上 γ^i 因子。本 level, $i=0$, 再下一层, $i=1$, --

Solving DMP:

value of s : $V^*(s) =$ expected utility starting from s and act optimally.
 value of q_h -state (s, a) : $Q^*(s, a) =$ expected utility starting out having taken action a from s and thereafter acting optimally

Optimal Policy: $\pi^*(s)$ = optimal action from state s

How to compute? **Recursive Definition:**

$$V^*(s) = \max_a Q^*(s, a), \text{ 而:}$$

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

$$\text{故: } V^*(s) = \max_a \sum_{s'} T(s, a, s') [R(s, a, s') + \gamma V^*(s')]$$

这被称为 Bellman Equation

但这种计算方式有2个问题: ① States are repeated

② Tree goes on forever (一种解法: depth-limited)

Value-Iteration: Time-limited Values

Define $V_k(s)$ to be the optimal value of s if the game ends in k more steps. 等价地: **depth-k expectimax**

从 $V_0(s) = 0$ 开始, 递归时: 若有 $V_k(s)$, 推 $V_{k+1}(s)$:

$$V_{k+1}(s) \leftarrow \max_a \sum_{s'} T(s, a, s') [R(s, a, s') + \gamma V_k(s')]$$

重复直至收敛。每一次递归: complexity 为 $O(S^2A)$

Δ : Theorem: will converge to unique optimal values

$$(\{V_{k+1}(s)\} \Rightarrow \overline{V_{k+1}(s)})$$

折扣因子 γ

Policy Extraction: 假设 $V^*(s)$ 已有。那么, 如何 act 呢?

$$\pi^*(s) = \arg\max_a \sum_{s'} T(s, a, s') [R(s, a, s') + \gamma V^*(s')]$$

来以此行动。上述便是 value-iteration.

Q-value Iteration: (流程与原理上与 value iteration 近似)

从 $Q_0(s, a) = 0$ 开始, 若知 $Q_k(s, a)$ (即全部信息)

$$Q_{k+1}(s, a) \leftarrow \sum_{s'} T(s, a, s') [R(s, a, s') + \gamma \max_{a'} Q_k(s', a')]$$

Repeat till converge. (通常比 Value Iteration 更快收敛)



Policy Iteration: 一种 value iter 的 alternative approach

Policy Evaluation $\rightarrow$ Policy Improvement $\rightarrow$ Repeat till converge!

Step1: 评估: Define $V^\pi(s)$ = expected utility starting in s and following fixed π :

$$V^\pi(s) = \sum_{s'} T(s, \pi(s), s') [R(s, \pi(s), s') + \gamma V^\pi(s')]$$

* 有 fixed policy, 算 utility 巨快!

显然此式要 iterate 才 converge, $\exists \lambda V_k^\pi(s)$

$$V_{k+1}^\pi(s) \leftarrow \sum_{s'} T(s, \pi(s), s') [R(s, \pi(s), s') + \gamma V_k^\pi(s')]$$

Or: $\vec{V}^\pi(s) = \sum_{s'} T(s, \pi(s), s') [R(s, \pi(s), s') + \gamma \vec{V}^\pi(s')]$

Step2: 改进: $\pi_{i+1}(s) = \arg\max_a \sum_{s'} T(s, a, s') [R(s, a, s') + \gamma V^\pi(s')]$

Repeat the two steps until policy converges

Δ : 要做 policy extraction

Summary: Value Iteration: update state value, don't track policy. Finally use state value to extract policy

Policy Iteration: May converge faster. After the policy is evaluated, a new policy is chosen

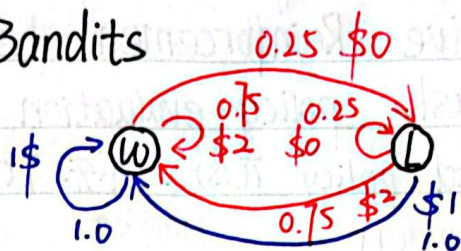


Reinforcement Learning

Intro Example: Double Bandits

Action: Blue, Red

States: Win, Lose



解 MDP 是 offline planning. 但这意味着: 你需要知道所有的信息, 包括 s take a 后 s' 抵达概率分布

若概率不知, 只知 reward 呢? 甚至 reward 也不知呢?

那么在 play 的时候, 实则在 'learning' (RL)。这的确是 MDP 问题, 但仅 compute 是学不了的。Need to act, to figure it out.

Key ideas:

Exploration vs. Exploitation. $\Rightarrow$ try unknown actions to gain info & eventually have to use what we know

强化学习中, 假设是 MDP: 有 states $s \in S$, 有 actions,

有 model $T(s, a, s')$, 有 reward function $R(s, a, s')$

我们 seek for $\pi(s)$, 但不知 T or R

Basic idea: 行动, 然后观察结果 (new states, rewards)。由此学习, 然后学着最大化期望收益

Model-based Learning

先用 experience 学近似 model, 然后运用它 as if the learned model is correct

Step 1: Learn Empirical MDP Model

数 对于 s, a 下的 s' 分布, 归一化后以代表 $\hat{T}(s, a, s')$; 并且发现每个 $\hat{R}(s, a, s')$ (当从 s 采取 a 抵达 s')

step 2: Solve the MDP



Model-Free Learning

这属于 Passive Reinforcement Learning.

Simplified task: policy evaluation

learn the state values

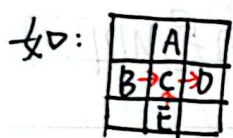
Input: a fixed policy $\pi(s)$. 不知 $T(s,a,s')$ & $R(s,a,s') \Rightarrow$ Goal:

Direct Evaluation

在这种范式下, 按照 policy 行动。但这不是 offline learning!

Act according to π . 每次到一个 s , 记录下:

s 之后的 the sum of discounted rewards; 然后求和



Episode 1 B, $\rightarrow$, C, -1 2. B, $\rightarrow$, C, -1

C, $\rightarrow$, D, -1 C, $\rightarrow$, A, -1

D, exit, x, 10 A, exit, x, -10

Direct Evaluation Example

则此时 State Value 计算: Epi 1 中: B 之后: 8, C 之后 9, D 10

Ep 2 中: B 之后 -12, C 之后 -11, A 之后 -10

则 B: $\frac{-12+8}{2} = -2$, C: $\frac{9-11}{2} = -1$, A: -10, D: 10, E: ?

Direct Evaluation 有一些缺点: 浪费了 state 之间的联系. 每个 state 单独学

新方法: Policy Evaluation

Simplified Bellman updates calculate V for a fixed policy.

$$V_0^{\pi}(s) = 0, \quad V_{k+1}^{\pi}(s) \leftarrow \sum_{s'} T(s, \pi(s), s') [R(s, \pi(s), s') + \gamma V_k^{\pi}(s')]$$

But don't know T & R ! How can we update V without T & R ?

Idea: Take samples of outcome s' (by doing the action!)* and 平均

$$\text{sample}_i = R(s, \pi(s), s'_i) + \gamma V_k^{\pi}(s'_i), \quad V_{k+1}^{\pi}(s) \leftarrow \frac{1}{n} \sum_i \text{sample}_i$$

*: 注意: can't rewind 故此 idea X!

New idea: Update $V(s)$ each time 经历一个 (s, a, s', r)

$$\textcircled{1} \text{ Sample} = R(s, \pi(s), s') + \gamma V^{\pi}(s')$$

$$\textcircled{2} \quad V^{\pi}(s) = (1-\alpha)V^{\pi}(s) + \alpha \text{ sample} = V^{\pi}(s) + \alpha (\text{sample} - V^{\pi}(s))$$

这个更新下: make recent samples more important

此法称为 Temporal Difference Learning



TD learning 是 model-free way to do policy evaluation. 但若
想利用评估的 values 转化为 new policy.

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

↑ 仍未知 ↑ $Q(s, a)$

Idea: Learn Q-values, not values!

⇒ Q-learning: $Q_{k+1}(s, a)$

Q-value iteration: $\sum_{s'} T(s, a, s') [R(s, a, s') + \gamma \max_{a'} Q_k(s', a')]$

每一次 receive sample (s, a, s', r) , 考虑到旧的 $Q(s, a)$:

$$\text{sample} = R(s, a, s') + \gamma \max_{a'} Q(s', a')$$

更新 $Q(s, a)$: $Q(s, a) \leftarrow (1-\alpha)Q(s, a) + \alpha [\text{sample}]$

Amazingly: 即使 act suboptimally, policy converge to optimal one.
这被称为 off-policy learning.

上述介绍完 learning paradigm, 现介绍探索与利用:

How to explore? Simplest: random (ϵ -greedy)

即小 prob ϵ 会 act randomly, 大 prob ϵ 会 act on current policy
但有问题: 学好了之后, 依然会 act randomly! 一种解法:

lower ϵ over time.

另一种: Exploration Function:

考虑 Q-value estimate 为 u , $s' \rightarrow a'$ count of visit 为 n :

Optimistic Utility: $f(u, n) = u + k/n$

则 Q update: original $Q(s, a) \leftarrow \alpha [R(s, a, s') + \gamma \max_{a'} Q(s', a')]$

后来: $Q(s, a) \leftarrow \alpha [R(s, a, s') + \gamma \max_{a'} f(Q(s', a'), N(s', a'))]$

↓ 这趋向于: 给 under-explored states 加 bonus!



Approximate Q-learning: 原来 Q learning 有 $Q(s,a)$ table.

但现在考虑 Feature-Based Representation 衡量 $Q(s,a)$

Linear form: $V(s) = \sum w_i f_i(s)$

$Q(s,a) = \sum w_i f_i(s,a)$, 则更新逻辑为:

transition = (s, a, r, s')

difference = $[r + \gamma \max_{a'} Q(s', a')] - Q(s, a)$

$Q(s,a) \leftarrow Q(s,a) + \alpha [\text{difference}]$ Exact Q's

✓ $w_i \leftarrow w_i + \alpha [\text{difference}] f_i(s,a)$ Approximate Q's
调整 feature 权重

Regret 是衡量 your rewards 与最优 rewards 间的差距。

⚠ Minimizing regret **going beyond** learning to be optimal:

It requires optimally learning to be optimal!

⇔ 为了 optimal policy, 不必要最小化 regret !!



* 即均值: 即使方值最小的点 *

Unsupervised ML

Clustering: Group together similar instances

'similar': One option: Euclidean Squared: $\text{dist}(x, y) = \sum_i (x_i - y_i)^2$
(也有其它 option in specific domain)

① K-means Approach: K 指分 K 类 $\leftarrow$ Squared Euclidean

$$\phi(\{x_i\}, \{a_i\}, \{c_k\}) = \sum_i \text{dist}(x_i, c_{a_i})$$

Two stage 循环: 更新 assignment $\xrightarrow{\text{①}}$ 更新 means (中心点) $\xrightarrow{\text{②}}$

① fix means c , change assignment

② 在该 assignment 下改 means c

Phase 1: $a_i = \underset{k}{\text{argmin}} \text{dist}(x_i, c_k)$ 每个点找最近 c_k

此步不会使 ϕ 函数上升

Phase 2: $c_k = \frac{1}{|\{i: a_i = k\}|} \sum_{i: a_i = k} x_i$ *, 亦不使 ϕ 上升

② 对于 K-means, 它 non-deterministic! 选择的 initials 很重要!
且可能 converge to 局部最优

GMM: Gaussian Mixture Model Approach

Recall Gaussian: $P(x|\mu, \sigma) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(x-\mu)^2}{2\sigma^2}}$

Given 观察到的 data:

$$\hat{\mu}_{MLE} = \frac{1}{N} \sum_{i=1}^N x_i \quad \hat{\sigma}_{MLE}^2 = \frac{1}{N} \sum_{i=1}^N (x_i - \hat{\mu})^2$$

多元高斯分布定义: n 维样本空间 $\mathcal{X}$ 中随机向量 x , PDF 为:

$$p(x) = \frac{1}{(2\pi)^{n/2} |\Sigma|^{1/2}} e^{-\frac{1}{2}(x-\mu)^T \Sigma^{-1} (x-\mu)}$$

其中 Σ 是 $n \times n$ 协方差矩阵

$\Rightarrow$ 高斯分布完全由均值向量 μ 与 Σ 确定!

KOKUYO



我们记: $p(x|\mu, \Sigma)$

则 GMM 中, 对于 k 类, 则定义高斯混合分布:

$$p_M(x) = \sum_{i=1}^k \alpha_i \cdot p(x|\mu_i, \Sigma_i)$$

每一类算 μ_i, Σ_i , 且 $\alpha_i > 0$ 为混合系数, $\sum_{i=1}^k \alpha_i = 1$

Inductive Bias: Sample 生成由 GMM 给出: 且有 $\alpha_1, \dots, \alpha_k$ 先验

则依 Bayes: $p_M(z_j=i|x_j) = \frac{P(z_j=i) p_M(x_j|z_j=i)}{\sum_{i=1}^k \alpha_i p(x_j|\mu_i, \Sigma_i)}$

$$= \frac{\alpha_i \cdot p(x_j|\mu_i, \Sigma_i)}{\sum_{i=1}^k \alpha_i p(x_j|\mu_i, \Sigma_i)}$$

$= \alpha_i (\text{prior}) p_M(x_j)$

其中 $z_j \in \{1, \dots, k\}$ 表示生成 x_j 的高斯

上在讨论: 在 Inductive Bias 下, x_j 是由 $z_j=i$ 个 Gaussian 生成的概率。式子显示: $p_M(z_j=i|x_j)$ 给出了样本 x_j 由第 i 个 Gaussian 生成的后验概率, 记为 γ_{ji} , $i \in \{1, \dots, k\}$

根据此逻辑, 每个样本 x_j 可依后验概率取 max 来判定 x_j 所属的簇标记: (λ_j)

$$\lambda_j = \operatorname{argmax}_{i \in \{1, \dots, k\}} \gamma_{ji}$$

一言以蔽之: 高斯混合聚类是采用高斯分布对原型进行刻画
簇划分则由原型对应后验概率确定

欲求解 $\{(\alpha_i, \mu_i, \Sigma_i) | 1 \leq i \leq k\}$, 最大化似然:

$$LL(D) = \sum_{j=1}^m \ln \left(\sum_{i=1}^k \alpha_i \cdot p(x_j|\mu_i, \Sigma_i) \right)$$

(即 $\ln \left(\prod_{j=1}^m p_M(x_j) \right)$)

常用 EM 算法优化

$$\mu_i = \frac{\sum_{j=1}^m \gamma_{ji} x_j}{\sum_{j=1}^m \gamma_{ji}} \quad \Sigma_i = \frac{\sum_{j=1}^m \gamma_{ji} (x_j - \mu_i)(x_j - \mu_i)^T}{\sum_{j=1}^m \gamma_{ji}} \quad \alpha_i = \frac{1}{m} \sum_{j=1}^m \gamma_{ji}$$



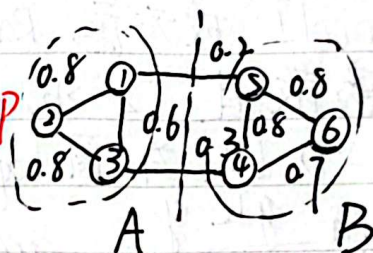
依式可推出 EM 流程:

(E) 根据当前参数计算每个样本属于每个高斯成分的后验 γ_{ji} , 再依三个式子更新参数 $\{(a_i, \mu_i, \Sigma_i) | 1 \leq i \leq k\}$ (M)

③. 谱聚类 Spectral Clustering

用顶点 V 代表数据点, 点间用边 E 相连。E 有 weight W , 权重越大, 说明两点间越 'similar'. 最后 cluster 可转化为 partitioning a similarity graph

则要 minimize weight of between-group connections



即找 groups with the minimal cut:

$$\text{cut}(A, B) = \sum_{i \in A, j \in B} w_{ij}$$

只看 inter-cluster connection 不够, 还有 intra-cluster!

$\Rightarrow$ 也要 maximize the weights of connections within groups

$$\max (\text{assoc}(A, A) + \text{assoc}(B, B)), \text{assoc}(A, A) = \sum_{i, j \in A} w_{ij}$$

命名上述: Criterion: Normalized Cuts

The connectivity between groups:

$$N_{\text{cut}}(A, B) = \frac{\text{cut}(A, B)}{\text{assoc}(A, V)} + \frac{\text{cut}(A, B)}{\text{assoc}(B, V)}$$

$$\text{assoc}(A, V) = \sum_{i \in A, j \in V} w_{ij}$$

$$N_{\text{assoc}}(A, B) = \frac{\text{assoc}(A, A)}{\text{assoc}(A, V)} + \frac{\text{assoc}(B, B)}{\text{assoc}(B, V)}$$

$$\text{则式子上 } N_{\text{cut}}(A, B) = 2 - N_{\text{assoc}}(A, B)$$

既然是个图了, 那么就有 Adjacency Matrix (Symmetric)

定义属于哪一簇: $x \in [-1, 1]^n$, $x_i = \begin{cases} 1, & i \in A \\ -1, & i \in B \end{cases}$ $d_i = \sum_j w_{ij}$

$$\text{则 } N_{\text{cut}}(A, B) = \frac{\sum_{x_i > 0, x_j < 0} -w_{ij} x_i x_j}{\sum_{x_i > 0} d_i} + \frac{\sum_{x_i < 0, x_j > 0} -w_{ij} x_i x_j}{\sum_{x_i < 0} d_i}$$



$$k = \frac{\sum_{i: d_i > 0} d_i}{\sum_i d_i}, \text{ 簇A的度占总度比例}$$

定义矩阵 $W \in \mathbb{R}^{n \times n}$ (adjacency) $D \in \mathbb{R}^{n \times n}$ (度, $d_{ii} = \sum_j w_{ij}$, 其余为0).

$$\text{则 } 4Ncut(A, B) = \frac{(\vec{1} + \vec{x})^T (D - W) (\vec{1} + \vec{x})}{k \vec{1}^T D \vec{1}} + \frac{(\vec{1} - \vec{x})^T (D - W) (\vec{1} - \vec{x})}{(1-k) \vec{1}^T D \vec{1}}$$

$$b = \frac{k}{1-k} \frac{[(\vec{1} + \vec{x}) - b(1-x)]^T (D - W) [(\vec{1} + \vec{x}) - b(1-x)]}{b \vec{1}^T b \vec{1}}$$

$$\min_{\vec{x}} Ncut(\vec{x}) = \min_{\vec{y}} \frac{\vec{y}^T (D - W) \vec{y}}{\vec{y}^T D \vec{y}} \quad \text{s.t. } \vec{y} \in [2-2b, -2b]^n, \vec{y}^T D \vec{1} = 0$$

NP-hard! 故考虑 Rayleigh Quotient 放宽约束, 简化问题

$\vec{y}$ 从离散值 $[2-2b, -2b]^n$ 变为 $\vec{y} \in \mathbb{R}^n$

定义 $L = D - W$, 为图 Laplacian Matrix

$$\text{则: } \min_{\vec{y}} \frac{\vec{y}^T L \vec{y}}{\vec{y}^T D \vec{y}}, \vec{y} \in \mathbb{R}^n, \vec{y}^T D \vec{1} = 0$$

Laplace

广义特征值

$$\Leftrightarrow \max_{\vec{x}} \frac{\vec{x}^T A \vec{x}}{\vec{x}^T B \vec{x}} \quad \text{s.t. } \vec{x}^T B \vec{x} = 1 \Leftrightarrow A \vec{x} = \lambda B \vec{x}, \lambda = -\lambda$$

$$\text{i.e., } (D - W) \vec{y} = \lambda D \vec{y}, \text{ 令 } \vec{z} = D^{1/2} \vec{y}, \text{ 则}$$

第二小的!!!

$$D^{-1/2} (D - W) D^{-1/2} \vec{z} = \lambda \vec{z}, \text{ 求解特征值与其}$$

对应特征方程 λ !

