

机器学习

一. 基本概念与评估指标

留出法 (hold-out): 直接划分为 train-valid (首次划分, 重复实验求平均)

交叉验证法 (cross-validation): k-折子集划分, 每次取 k-1 个子集训练, 1 个验证, 重复 k 次, 取均值. $k=10$ 准确但计算开销高

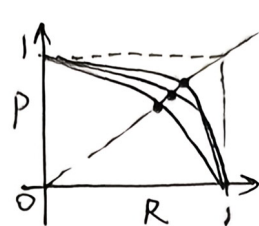
真实 \ 预测	正	反
正	TP	FN
反	FP	TN

(假正例)
(负例)

$$\text{Precision} = P = \frac{TP}{TP+FP}$$

$$\text{Recall} = R = \frac{TP}{TP+FN}$$

P-R 曲线



$P=R$: 平衡点.

平衡点的取值度量分类器的性能

$$F1 = \frac{2}{\frac{1}{P} + \frac{1}{R}} = \frac{2PR}{P+R} \quad F\beta = \frac{(1+\beta^2)PR}{\beta^2 P + R} \quad \begin{cases} \beta > 1 \rightarrow R \uparrow & (\beta = \infty, F\beta = R) \\ \beta < 1 \rightarrow P \uparrow & (\beta = 0, F\beta = P) \end{cases}$$

[Definition] ML: 通过经验 (experience) 提升性能表现的学习系统.

分类: $\begin{cases} \text{监督: label-classification/regression } x \rightarrow y \\ \text{无监督: special info-clustering (interesting patterns)} \\ \text{强化: max reward - 环境交互. label exists, credit assignment.} \end{cases}$

二. 贝叶斯分类:

Naive Bayesian Classifier:

Given $X = (x_1, \dots, x_p)^T$, predict w . ($\max P(w|x) \propto P(x|w)P(w)$).

Eg. $X = (\text{Refund} = \text{No}, \text{Married}, \text{Income} = 120k)$

1. compute $P(\text{Refund} = ? | \text{No/Yes})$, $P(\text{Marital status} = ? | \text{No/Yes})$, sample mean & variance of $\text{Income} | \text{No/Yes}$.

2. compute $P(X|\text{No}) = P(\text{Refund} = \text{No} | \text{No}) \times P(\text{Married} | \text{No}) \times P(\text{Income} = 120k | \text{No})$, $P(X|\text{Yes})$

3. compare $P(X|\text{No})P(\text{No})$ and $P(X|\text{Yes})P(\text{Yes})$

4. the bigger one is the output.

$\hat{=}$ MLE

regression model: $y = f(x, w) + \varepsilon$, $\varepsilon \sim N(0, \sigma^2)$. $p(y|x, w, \sigma) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2\sigma^2}(y-f(x, w))^2}$

$$\mathcal{L}(D, w, \sigma) = \prod_{i=1}^n p(y_i | x_i, w, \sigma), w^* = \arg \max \mathcal{L}$$

$$\mathcal{L}(D, w, \sigma) = \log \mathcal{L} = \sum_{i=1}^n \log p(y_i | x_i, w, \sigma) = -\frac{1}{2\sigma^2} \sum_{i=1}^n (y_i - f(x_i, w))^2 + c(\sigma).$$

$$\mathcal{L}'(D, w, \sigma) = 0 \Leftrightarrow \sum_{i=1}^n (y_i - f(x_i, w)) = \text{RSS}(f) = 0. \quad J_n = \frac{1}{n} \sum_{i=1}^n (y_i - f(x_i))^2 = \text{MSE}$$

Residual Sum of Squares 残差平方和

四. 线性回归与分类

$$(x_i, y_i), f = \sum_{j=1}^m w_j x_{ij} = w^T x, w = [w_1, w_2, \dots, w_m]^T$$

$$\text{MSE} = J_n = \frac{1}{n} \sum_{i=1}^n (y_i - f(x_i))^2 = \frac{1}{n} \sum_{i=1}^n (y_i - w^T x_i)^2 \Rightarrow J_n(w) = (y - X^T w)^T (y - X^T w)$$

$$\nabla J_n = -2X(y - X^T w) = 0 \Rightarrow Xy = XX^T w \Rightarrow w = (XX^T)^{-1} Xy.$$

$$\hat{y} = X^T w = X^T (XX^T)^{-1} Xy.$$

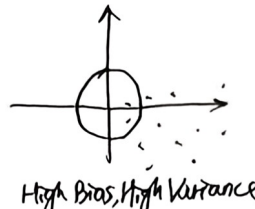
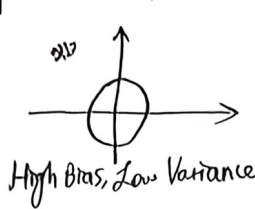
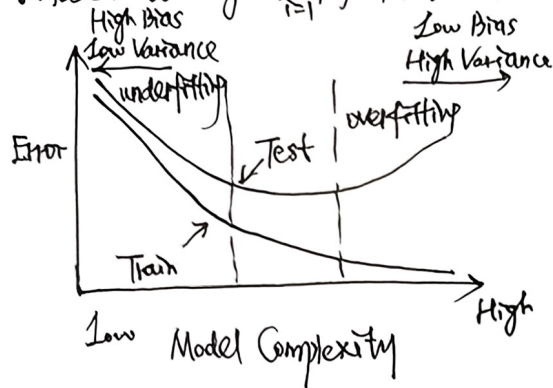


岭回归: $w^* = \arg\min \sum_{i=1}^n (y_i - x_i^T w)^2 + \lambda \sum_{j=1}^p w_j^2 \Leftrightarrow \alpha^* = \arg\min \sum_{i=1}^n (y_i - x_i^T w)^2$ subject to $\sum_{j=1}^p w_j^2 \leq t$.

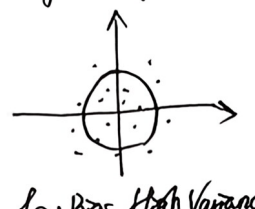
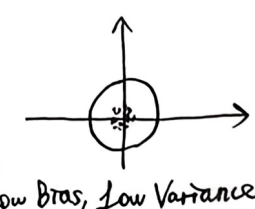
(拉格朗日乘子) $J_n = (y - X^T w)^T (y - X^T w) + \lambda w^T w$.

求导: $\nabla J_n = -2X(y - X^T w) + 2\lambda w = 0 \Rightarrow (XX^T + \lambda I) w = Xy \Rightarrow w^* = (XX^T + \lambda I)^{-1} Xy$.

LASSO: $\hat{w} = \arg\min \sum_{i=1}^n (y_i - x_i^T w)^2 + \lambda \|w\|_1$, Sparse Model



(large λ): over-regularized (underfitting) high bias



Small λ : under-regularized (overfitting) high variance

Classifier:

regression = $f(x) = w^T x \in \mathbb{R}$.

sigmoid / logistic function $\sigma(x) = \frac{1}{1+e^{-x}} \in (0, 1)$

$P(y_i = \pm 1 | x_i, a) = \sigma(y_i a^T x_i) = \frac{1}{1+e^{-y_i a^T x_i}}$

MLE: $\mathcal{L} = \prod_{i=1}^n \sigma(y_i a^T x_i)$

$\ell = -\log \mathcal{L} = -\sum_{i=1}^n \log \sigma(y_i a^T x_i) = -\sum_{i=1}^n \log \frac{1}{1+e^{-y_i a^T x_i}} = \sum_{i=1}^n \log(1+e^{-y_i a^T x_i}) = E(a)$. Gradient Desc

$\min J(w)$. $w_{n+1} = w_n - \gamma \nabla J(w_n)$.

Multi-Class Classifier.

One vs Rest: 多个one分类器

One vs One: 每两个类别作为一个分类器 $\frac{K(K-1)}{2}$ 个分类器

ECOC: 多分类问题转化为多个二分类问题. 海明距离最小. 任意两个类别二元的汉明距离最大

码长为9. 类别为4: $A = \{000000000\}$, $B = \{000111111\}$, $C = \{111000111\}$, $D = \{111111000\}$.

类别	f_1	f_2	...	f_9
A				
B				
C				
D				

训练样本: $\{1, 0, 1, 0, \dots\}$

计算海明距离. 选择距离最小的作为预测结果



SVM: $f(x) = w^T x + b = 0$.

$f(x_p) = w^T(x - \gamma \frac{w}{\|w\|}) + b = 0$. $f(x) = w^T x + b = r \|w\|$. $\Rightarrow r = \frac{f(x)}{\|w\|}$

$\gamma = y \cdot r = y \frac{w^T x + b}{\|w\|}$. geometrical margin

Maximum Margin Classifier: $\max_{w,b} \gamma = \max_{w,b} \frac{y(w^T x + b)}{\|w\|}$, s.t. $y_i \geq \gamma$.

fix $y_i(w^T x_i + b) = 1$. $\max_{w,b} \gamma \Leftrightarrow \min_{w,b} \frac{1}{2} \|w\|^2$ s.t. $y_i = \frac{y_i(w^T x_i + b)}{\|w\|} \geq \gamma$

$y_i = \frac{y_i(w^T x_i + b)}{\|w\|} \geq \gamma \Leftrightarrow y_i(w^T x_i + b) \geq \gamma \|w\|$

MVC: $\min_{w,b} \frac{1}{2} \|w\|^2$, s.t. $y_i(w^T x_i + b) \geq 1$.

支持向量: margin 上的点, 非支持向量的点都不在 margin 边界上.

Slack variables $\min_{w,b,\xi} \frac{1}{2} \|w\|^2 + C \sum_{i=1}^n \xi_i$, s.t. $y_i(w^T x_i + b) \geq 1 - \xi_i$, $\xi_i \geq 0$.

C ↑, $\xi \downarrow$, overfitting.

C ↓, $\xi \uparrow$ underfitting

核函数: 解决线性不可分问题且避免维数灾难 核函数: 直接计算高维内积.

五. 神经网络

1. 感知机: $f(x) = w^T x + b$. predict wrong, set $w_{t+1} = w_t + xy$. ($w^T xy > 0$)

OR: $f(x) = x_1 + x_2 + 0.5$, AND: $f(x) = x_1 + x_2 - 1.5$, XOR: X

$\mathcal{L}(w,b) = \sum_{x \in M} -y_i(w^T x_i + b)$. $\nabla \mathcal{L} = \sum_{x \in M} -x_i y_i$ $w_{k+1} = w_k + \eta \sum_{x \in M} x_i y_i$ (learning rate) (gradient descent)

2. 反向传播:

forward: $z = f(\sum w_{kj} f(\sum u_{ji} x_i + u_{j0}) + u_{k0})$

Activation function 激活函数

$\Delta w = -\eta \Delta J = -\eta \frac{\partial J}{\partial w}$. $w_{k+1} = w_k - \eta \frac{\partial J}{\partial w}$. $\frac{\partial J}{\partial w_{ji}} = \frac{\partial J}{\partial y} \cdot \frac{\partial y}{\partial z^2} \cdot \frac{\partial z^2}{\partial a_1'} \cdot \frac{\partial a_1'}{\partial z_1'} \cdot \frac{\partial z_1'}{\partial w_{ji}}$

例: Input (x_1, x_2) . hidden (h_1, h_2) output $\hat{y}$. $f = \text{sigmoid}$. $\mathcal{L} = \text{MSE} = \frac{1}{2} (y - \hat{y})^2$.

$x = [x_1, x_2] = [0.5, -0.8]$, $y = 1$. $w' = \begin{bmatrix} w'_{11} & w'_{12} \\ w'_{21} & w'_{22} \end{bmatrix} = \begin{bmatrix} 0.1 & 0.3 \\ -0.2 & 0.4 \end{bmatrix}$, $b' = [0.1, -0.1]$, $w^2 = [0.2, -0.5]$, $b^2 = 0.05$.

forward: $z_1' = w'_{11} x_1 + w'_{12} x_2 + b_1' = 0.31$. $a_1' = \sigma(0.31) = 0.58$

$z_2' = w'_{21} x_1 + w'_{22} x_2 + b_2' = -0.27$. $a_2' = \sigma(-0.27) = 0.43$.

$z^2 = w_2^1 a_1' + w_2^2 a_2' + b^2 = -0.05$, $a^2 = \sigma(-0.05) = \hat{y} = 0.49$. ($\hat{y} = \frac{1}{1+e^{-z}}$)

backward: $\frac{\partial \mathcal{L}}{\partial y} = \hat{y} - y = -0.51$. $\frac{\partial \mathcal{L}}{\partial z^2} = \hat{y}(1-\hat{y}) = 0.25$. $\delta^2 = \frac{\partial \mathcal{L}}{\partial y} \cdot \frac{\partial y}{\partial z^2} = -0.13$

$\frac{\partial \mathcal{L}}{\partial w_2^1} = \delta^2 \cdot a_1' = -0.07$, $\frac{\partial \mathcal{L}}{\partial w_2^2} = \delta^2 \cdot a_2' = -0.06$. $\frac{\partial \mathcal{L}}{\partial b^2} = \delta^2 = -0.13$.

$\frac{\partial z^2}{\partial a_1'} = w_2^1 = 0.2$, $\frac{\partial z^2}{\partial a_2'} = w_2^2 = -0.5$, $\frac{\partial a_1'}{\partial z_1'} = a_1'(1-a_1') = 0.24$, $\frac{\partial a_2'}{\partial z_2'} = 0.25$.

$\delta_1' = \delta^2 \times w_2^1 \times \frac{\partial a_1'}{\partial z_1'} = -0.01$, $\delta_2' = \delta^2 \times w_2^2 \times \frac{\partial a_2'}{\partial z_2'} = 0.02$.

$\frac{\partial \mathcal{L}}{\partial w_{ji}'} = \delta_1' \times x_i$, $\frac{\partial \mathcal{L}}{\partial w_{21}'} = \delta_1' \times x_1$, $\frac{\partial \mathcal{L}}{\partial w_{22}'} = \delta_1' \times x_2$, $\frac{\partial \mathcal{L}}{\partial b_1'} = \delta_1'$.

$w_{11}^2_{\text{new}} = w_{11}^2 - \eta \frac{\partial \mathcal{L}}{\partial w_{11}^2}$, $w_{22}^2_{\text{new}} = w_{22}^2 - \eta \frac{\partial \mathcal{L}}{\partial w_{22}^2}$, $b^2_{\text{new}} = b^2 - \eta \frac{\partial \mathcal{L}}{\partial b^2}$

$w_{11}^1, w_{12}^1, w_{21}^1, w_{22}^1$. b_1^1, b_2^1 同样更新.



激活函数: non-linear: sigmoid, tanh.

$$\frac{e^x - e^{-x}}{e^x + e^{-x}} \in (-1, 1) = \sigma(2x) - 1 \quad \text{vanishing gradient} \Rightarrow \text{ReLU, dropout. (overfitting)}$$

六. 决策树:

熵: $\text{Entropy}(S) = -\sum_{i=1}^C p_i \log p_i$. $\text{Gain}(S, a) = \text{Entropy}(S) - \sum_{v \in \text{values}(a)} \frac{|S_v|}{|S|} \text{Entropy}(S_v)$.

例:

ID	Age	Gender	Income	Buy?
1	27	M	15W	X
②	47	F	30W	✓
3	32	M	12W	X
4	24	M	45W	✓
* ⑤	45	M	30W	X
* ⑥	56	M	32W	✓
7	31	M	15W	X
⑧	23	F	30W	✓

Age: 20-30, 30-40, 40+
Income: 10-20, 20-40, 40+
 $a \in \{\text{Age, Gender, Income}\}$

(4+, 4-)
 $\text{Entropy}(S) = -\frac{1}{2} \log \frac{1}{2} - \frac{1}{2} \log \frac{1}{2} = 1$.
Age: 20-30: 2+, 1-. $E = -\frac{2}{3} \log \frac{2}{3} - \frac{1}{3} \log \frac{1}{3} = 0.918$
30-40: 0+, 2-. $E = 0$
40+: 2+, 1-. $E = 0.918$

$\text{Gain}(S, \text{Age}) = 1 - (\frac{3}{8} \times 0.918 + \frac{2}{8} \times 0 + \frac{2}{8} \times 0.918) = 0.312$.

Income: 10-20: 0+, 3-. $E = 0$
20-40: 3+, 1-. $E = 0.811$
40+: 1+, 0-. $E = 0$

Age: 20-30,
30-40: 无法划分
40+:

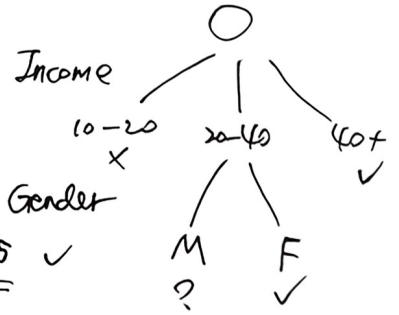
Gender: M: 1+, 1- $E = 1$
F: 2+, 0- $E = 0$

$\text{Gain} = 0.811 - \frac{2}{4} \times 1 = 0.311$

同理, 计算

$\text{Gain}(S, \text{Gender}) = 0.312$

$\text{Gain}(S, \text{Income}) = 0.595$ ✓



$\text{Gain} = 0.811 - (\frac{1}{4} \times 0 + \frac{3}{4} \times 0.918) = 0$

CART (Classification and Regression Tree) = partition the input space for continuous features (greedy method)

$\min \left[\min_{C_1} \sum_{x_i} (y_i - C_1)^2 + \min_{C_2} \sum_{x_i} (y_i - C_2)^2 \right]$
feature - threshold

七. Bagging & Boosting / Ensemble Method

1. Bagging

选择决策树作为 base learner: 1. non-linear, 2. easy to use, 3. easy to overfit.

定义: Bagging 是一种并行集成方法, 通过重采样投票(分类) (取平均(回归)) 测试预测结果

Random Forest: 1. random sampling

★ 2. randomized feature selection at each split step

no pruning for overfitting

Vote for prediction



2. Boosting.

定义: 顺序聚合, 每次提高被错分类样本的权重.

AdaBoost: 1. $w_n^{(1)} = \frac{1}{N}$

2. for $m=1, \dots, M$:

3. fit classifier $y^{(m)}(x)$, $J_m = \sum_{n=1}^N w_n^{(m)} I(y^{(m)}(x_n) \neq t_n)$

4. $\Sigma_m = \frac{J_m}{\sum_{n=1}^N w_n^{(m)}}$, $\alpha_m = \ln \frac{1 - \Sigma_m}{\Sigma_m}$

5. update w_n^{m+1}

6. $Y_M(x) = \text{sign}(\sum_{m=1}^M \alpha_m y^{(m)}(x))$ # final model

Gradient Boosting Decision Tree (GBDT): predict remaining data & sum

1. KNN.

stored record
distance metrics $\Rightarrow$ $\begin{cases} 1. \text{ compute the distance to other record} \\ 2. \text{ find the } k \text{ nearest neighbor} \\ 3. \text{ vote} \end{cases}$

7. Clustering

K-Means: 1. Randomly pick k data points as μ_k

2. Iterate until converge:

3. $Z_i = \arg \min_k \|x_i - \mu_k\|^2$, # assign each point to the cluster

4. $\mu_k = \frac{1}{N_k} \sum_{i=1}^{N_k} x_i^{(k)}$ # update seed point.

8. dimensionality reduction.

1. PCA (unsupervised)

1st PC: $\tilde{z}^{(1)} = a_1^T x_i$, s.t. $\max \text{Var}(\tilde{z}^{(1)})$

Solution: $\text{Var}(\tilde{z}^{(1)}) = E((\tilde{z}^{(1)} - \bar{\tilde{z}}^{(1)})^2) = \frac{1}{n} \sum_{i=1}^n (a_1^T x_i - a_1^T \bar{x})^2 = \frac{1}{n} \sum_{i=1}^n a_1^T (x_i - \bar{x})(x_i - \bar{x})^T a_1 = a_1^T S a_1$

$\Rightarrow \max_{a_1} a_1^T S a_1$, s.t. $a_1^T a_1 = 1$. $\mathcal{L} = a_1^T S a_1 - \lambda(a_1^T a_1 - 1)$. $\frac{\partial \mathcal{L}}{\partial a_1} = 0 \Rightarrow \lambda a_1 = S a_1$

$\Rightarrow a_1$ is the eigenvector of S corresponds to the largest eigenvalue λ_1 , $S = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})(x_i - \bar{x})^T$

2nd PC: $\max_{a_2} a_2^T S a_2$, s.t. $a_2^T a_2 = 1$, $\text{Cov}(\tilde{z}^{(1)}, \tilde{z}^{(2)}) = 0$.

例: Given $(90, 85, 80); (60, 65, 70); (85, 80, 75); (70, 75, 80); (95, 90, 85)$. $\bar{x} = (80, 77, 78)$

$\Rightarrow (10, 6, 2); (-20, -14, -8); (5, 1, -3); (-10, -4, 1); (15, 11, 7)$

$\Rightarrow S = \frac{1}{5} \left(\begin{pmatrix} 10 \\ 6 \\ 2 \end{pmatrix} \begin{pmatrix} 10 & 6 & 2 \end{pmatrix} + \begin{pmatrix} -20 \\ -14 \\ -8 \end{pmatrix} \begin{pmatrix} -20 & -14 & -8 \end{pmatrix} + \dots \right) = \begin{pmatrix} 170 & 110 & 50 \\ 110 & 74 & 38 \\ 50 & 38 & 26 \end{pmatrix}$

$(S - \lambda I) = \begin{vmatrix} 170-\lambda & 110 & 50 \\ 110 & 74-\lambda & 38 \\ 50 & 38 & 26-\lambda \end{vmatrix} = 0 \Rightarrow \lambda_1 = 258.875, \lambda_2 = 11.125, \lambda_3 = 0$

$(S - \lambda_1 I) v_1 = 0$
 $v_1 = \begin{pmatrix} 0.805 \\ 0.533 \\ 0.260 \end{pmatrix} = \hat{PC}_1$

~~$(x_i - \bar{x}) \cdot v_1$~~ $(x_i - \bar{x}) \cdot v_1$



2. LDA (supervised).

Rayleigh quotient: $J(a) = \frac{a^T S_B a}{a^T S_W a}$, S_W : 类内协方差, S_B : 类间协方差

$\max J(a) \Leftrightarrow S_W^{-1} S_B$

例:

类别	数	类
1	90	70
2	85	65
3	80	60
4	60	85
5	65	90
6	70	80

$\mu_1 = [85 \ 65]$
 $\mu_2 = [65 \ 85]$
 $\mu = [75 \ 75]$

S_W : 例: $\begin{pmatrix} 5 & 5 \\ 0 & 0 \\ -5 & -5 \end{pmatrix}$, $S_1 = \sum_i (X_i - \mu_1)(X_i - \mu_1)^T = \begin{pmatrix} 5 & 5 \\ 0 & 0 \end{pmatrix} + \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix} + \begin{pmatrix} -5 & -5 \\ -5 & -5 \end{pmatrix} = \begin{pmatrix} 5 & 5 \\ 0 & 0 \end{pmatrix}$

2: $\begin{pmatrix} -5 & 0 \\ 0 & 5 \\ 5 & -5 \end{pmatrix}$, $S_2 = \begin{pmatrix} 50 & -25 \\ -25 & 50 \end{pmatrix}$, $S_W = \begin{pmatrix} 100 & 25 \\ 25 & 100 \end{pmatrix}$

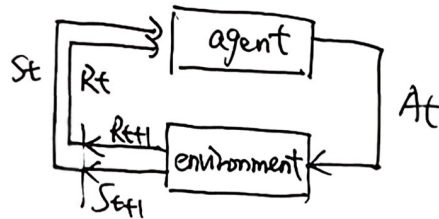
$S_B = S_B = \sum_{i=1}^2 n_i (\mu_i - \mu)(\mu_i - \mu)^T = 3 \begin{pmatrix} 10 & -10 \\ -10 & 10 \end{pmatrix} + 3 \begin{pmatrix} -10 & 10 \\ 10 & -10 \end{pmatrix} = \begin{pmatrix} 600 & -600 \\ -600 & 600 \end{pmatrix}$

$(S_B w = \lambda S_W w) \Rightarrow \begin{pmatrix} 600 & -600 \\ -600 & 600 \end{pmatrix} w = \lambda \begin{pmatrix} 100 & 25 \\ 25 & 100 \end{pmatrix} w$

$S_W^{-1} S_B = \begin{pmatrix} 8 & -8 \\ -8 & 8 \end{pmatrix} \Rightarrow \lambda_1 = 16, \lambda_2 = 0$
 $w_1 = (1, -1)$

+- 强化学习.

agent-environment interaction



(State, Action, Reward) = (S, A, R).

goal: optimal policy to max reward R.

policy $\pi(s)$
reward function
value function $V_\pi(s)$

Markov Decision Process (MDP): $P_r(S_{t+1}=s', R_{t+1}=r | S_t, A_t, R_t, S_{t-1}, A_{t-1}, \dots, S_0, A_0) = P_r(S_{t+1}=s', R_{t+1}=r | S_t, A_t)$

Bellman equation: $Q_\pi(s, a) = \sum_{s', r} p(s', r | s, a) (r + \gamma V_\pi(s'))$
 $V_\pi(s) = \sum_a (\pi(a|s) \sum_{s', r} p(s', r | s, a) (r + \gamma V_\pi(s')))$

例:

start	→	→	+
↑	///	↑	-
↑	←	←	←

0.812	0.868	0.912	+
0.762	///	0.660	-
0.705	0.655	0.611	0.388

$P = 0.8 \checkmark$
 $P = 0.1 \leftarrow, 0.1 \rightarrow$
 -0.04 each step, $\gamma = 1$

$V_\pi(s) = \sum_a \pi(a|s) \sum_{s', r} p(s', r | s, a) (r + \gamma V_\pi(s'))$
 $= 0.8 \times 1 \times (-0.04 + 1 \times 0.868) + 0.1 \times 1 \times (-0.04 + 1 \times 0.812) + 0.1 \times 1 \times (-0.04 + 0.868)$

$Q_\pi(s, a) = \sum_{s', r} p(s', r | s, a) (r + \gamma V_\pi(s'))$
 $= 1 \times (-0.04 + 1 \times 0.868)$

6.



DP: $V_{\pi}(s_t) = E_{\pi}[R_{t+1} + \gamma V(s_{t+1})]$

policy evaluation: $V_{\pi}(s) = \sum_{a \in A} (\pi(a|s) \sum_{s', r} p(s', r|s, a) (r + \gamma V_{\pi}(s'))) \quad \{V_{\pi}(s)\} \Rightarrow \{V_{\pi}(s)\}$

1. improvement: $q_{\pi}(s, \pi'(s)) > V_{\pi}(s) \Leftrightarrow V_{\pi'}(s) > V_{\pi}(s)$

$\pi'(s) = \arg \max_a \sum_{s', r} p(s', r|s, a) (r + \gamma V_{\pi}(s'))$: greedy
 $q_{\pi}(s', a)$

Monte Carlo: $q_{\pi}(s, a) = E[G_t]$. $V(s_t) \leftarrow V(s_t) + \alpha (G_t - V(s_t))$

exploration starts: first S-A pair return at the end

Temporal-Difference Learning: $V(s_t) \leftarrow V(s_t) + \alpha (R_{t+1} + \gamma V(s_{t+1}) - V(s_t))$

Sarsa: ϵ -greedy: Σ for exploration (all-actions) return in one step

$1-\epsilon$ for exploitation ($\pi(s)$)

$Q(s_t, A_t) \leftarrow Q(s_t, A_t) + \alpha [R_{t+1} + \gamma Q(s_{t+1}, A_{t+1}) - Q(s_t, A_t)]$

Q-learning: fully-greedy exploration. $\pi(s) = \arg \max_a Q(s, a)$

$Q(s_t, A_t) \leftarrow Q(s_t, A_t) + \alpha [R_{t+1} + \gamma \max_{a'} Q(s_{t+1}, A_{t+1}) - Q(s_t, A_t)]$

