

مثال: طبقه‌بندی در دو دسته

$\alpha_1$ : deciding  $\omega_1$

$\alpha_2$ : deciding  $\omega_2$

$$\lambda_{ij} = \lambda(\alpha_i | \omega_j)$$

خط شرطی برای هر یک از دسته‌ها:

$$R(\alpha_1 | x) = \lambda_{11} P(\omega_1 | x) + \lambda_{12} P(\omega_2 | x)$$

$$R(\alpha_2 | x) = \lambda_{21} P(\omega_1 | x) + \lambda_{22} P(\omega_2 | x)$$

قاعده‌ی تصمیم‌گیری: برزی برابر حداقل کردن خطر:

decide  $\omega_1$  if  $R(\alpha_1 | x) < R(\alpha_2 | x)$

$$\Rightarrow \lambda_{11} P(\omega_1 | x) + \lambda_{12} P(\omega_2 | x) < \lambda_{21} P(\omega_1 | x) + \lambda_{22} P(\omega_2 | x)$$

$$\Rightarrow (\lambda_{11} - \lambda_{21}) P(\omega_1 | x) < (\lambda_{22} - \lambda_{12}) P(\omega_2 | x)$$

$$\Rightarrow (\lambda_{11} - \lambda_{21}) p(x | \omega_1) P(\omega_1) < (\lambda_{22} - \lambda_{12}) p(x | \omega_2) P(\omega_2)$$

$$\Rightarrow \frac{p(x | \omega_1)}{p(x | \omega_2)} > \underbrace{\frac{\lambda_{22} - \lambda_{12}}{\lambda_{11} - \lambda_{21}} \frac{P(\omega_2)}{P(\omega_1)}}_{\theta_\lambda}, \quad \lambda_{11} < \lambda_{21}$$

یک مقدار آستانه ثابت مستقل از  $x$    
 نسبت درست‌نمایی‌ها

Duda\_fig2.3

$\theta_\lambda$

مثال: تابع ضرر صفر و یک  $\Lambda = [\lambda(\alpha_i | \omega_j)] = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}$

$$\theta_\lambda = \frac{0 - 1}{0 - 1} \frac{P(\omega_2)}{P(\omega_1)} = \frac{P(\omega_2)}{P(\omega_1)}$$

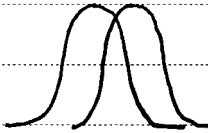
مثال:  $\Lambda = \begin{bmatrix} 0 & 2 \\ 1 & 0 \end{bmatrix}$

$$\theta_\lambda = \frac{0 - 2}{0 - 1} \frac{P(\omega_2)}{P(\omega_1)} = 2 \frac{P(\omega_2)}{P(\omega_1)}$$

$\omega_1, \omega_2$ مثال: سند طبقه بندی دو کلاس با ویژگی  $x$ 

$$\omega_1: p(x|\omega_1) = \frac{1}{\sqrt{\pi}} \exp(-x^2), \quad P(\omega_1) = 1/2$$

$$\omega_2: p(x|\omega_2) = \frac{1}{\sqrt{\pi}} \exp(-(x-1)^2), \quad P(\omega_2) = 1/2$$



الف) قاعدهی تقسیم گیری بیز برای تصمیم سازی احتمال خطا:

decide  $\omega_1$  if  $P(\omega_1|x) > P(\omega_2|x)$ 

$$\Rightarrow P(\omega_1) p(x|\omega_1) > P(\omega_2) p(x|\omega_2)$$

$$\Rightarrow \frac{1}{2} \frac{1}{\sqrt{\pi}} \exp(-x^2) > \frac{1}{2} \frac{1}{\sqrt{\pi}} \exp(-(x-1)^2)$$

$$\Rightarrow \exp(-x^2) > \exp(-(x-1)^2)$$

$$\Rightarrow -x^2 > -(x-1)^2$$

$$\Rightarrow x^2 < (x-1)^2$$

$$\Rightarrow x^2 < x^2 - 2x + 1$$

$$\Rightarrow x < \frac{1}{2}$$

ب) قاعدهی تقسیم گیری بیز برای تصمیم سازی ریسک (خطا) با ماتریس هزینه

$$\Lambda = \begin{bmatrix} 0 & 0.5 \\ 1 & 0 \end{bmatrix}$$

$$R(\alpha_i|x) = \sum_{j=1}^c \lambda(\alpha_i|\omega_j) P(\omega_j|x)$$
decide  $\omega_1$  if  $R(\omega_1|x) < R(\omega_2|x)$ 

$$\Rightarrow \lambda(\omega_1|\omega_1) P(\omega_1|x) + \lambda(\omega_1|\omega_2) P(\omega_2|x) <$$

$$\lambda(\omega_2|\omega_1) P(\omega_1|x) + \lambda(\omega_2|\omega_2) P(\omega_2|x)$$

$$\Rightarrow 0.5 P(\omega_2|x) < 1 \times P(\omega_1|x)$$

$$\Rightarrow 0.5 \exp(-(x-1)^2) < \exp(-x^2)$$

$$\Rightarrow \exp(-(x-1)^2) < 2 \exp(-x^2)$$

$$\Rightarrow \exp(-(x-1)^2) < \exp(\ln 2 - x^2)$$

$$\Rightarrow (x-1)^2 > x^2 - \ln 2 \Rightarrow -2x + 1 > -\ln 2 \Rightarrow x < \frac{1 + \ln 2}{2}$$

شال : (2-2)

مسئله طبقه بندی دو کلاس دو بعدی (برداری ویرگنی  $x \in \mathbb{R}^2$ )  
برداری های ویرگنی با دو توزیع نرمال با یک ماتریس کوواریانس مشترک

$$\Sigma = \begin{bmatrix} 1.1 & 0.3 \\ 0.3 & 1.9 \end{bmatrix}$$

$$\omega_1 : \mu_1 = [0, 0]^T$$

$$\omega_2 : \mu_2 = [3, 3]^T$$

و برداری های بیانگین :

(الف) بردار  $[1.0, 2.2]^T$  را با استفاده از طبقه بندی کننده ی بی زی ، طبقه بندی کند .

\* کافی است فاصله ی ماکزیمم این بردار را از دو بردار بیانگین مربوط به دو طبقه محاسبه کنیم :

$$\text{decide } \omega_1 \text{ if } d_m^2(\mu_1, x) < d_m^2(\mu_2, x)$$

$$\Rightarrow (x - \mu_1)^T \Sigma^{-1} (x - \mu_1) < (x - \mu_2)^T \Sigma^{-1} (x - \mu_2)$$

$$\Rightarrow [1, 2.2] \begin{bmatrix} 0.95 & -0.15 \\ -0.15 & 0.55 \end{bmatrix} \begin{bmatrix} 1 \\ 2.2 \end{bmatrix} < [-2, -0.8] \begin{bmatrix} 0.95 & -0.15 \\ -0.15 & 0.55 \end{bmatrix} \begin{bmatrix} -2 \\ -0.8 \end{bmatrix}$$

$$\Rightarrow 2.952 < 3.672$$

نکته : اگر از فاصله ی آقلیدس استفاده می کردیم ، بردار ویرگنی به  $\mu_2$  نزدیک تر می بود .

(ب) محورهای اصلی بیضی با مرکز  $[0, 0]^T$  متناظر با فاصله ی ماکزیمم  $d_m = \sqrt{2.952}$  از مرکز را محاسبه کند .

\* برای این هدف مقادیر ویژه ی  $\Sigma$  را محاسبه می کنیم :

$$\det(\Sigma - \lambda I) = 0 \Rightarrow$$

$$\left| \begin{bmatrix} 1.1 - \lambda & 0.3 \\ 0.3 & 1.9 - \lambda \end{bmatrix} \right| = 0 \Rightarrow \lambda^2 - 3\lambda + 2 = 0 \Rightarrow \lambda_1 = 1, \lambda_2 = 2$$

برای یافتن بردارهای ویژه ی متناظر با این مقادیر ، آنها را در معادله ی  $(\Sigma - \lambda I)v = 0$  قرار می دهیم :

$$v_1 = \left[ \frac{3}{\sqrt{10}}, -\frac{1}{\sqrt{10}} \right]^T, \quad v_2 = \left[ \frac{1}{\sqrt{10}}, \frac{3}{\sqrt{10}} \right]^T$$

و لا و لا برهم عود هستند .

$$\|v_1\| = 3.436, \quad \|v_2\| = 4.859$$

## Receiver Operating Characteristics

شخصی ROC :

مشخصه‌ی عملکرد گیرنده

## Signal Detection Theory and Operating Characteristics

معیاری برای سنجش فاصله میان دو توزیع گادوسی

کامبرد در: روان شناسی تجربی، آشکارسازی رادار و ...

$$p(x|w_1) \sim N(\mu_1, \sigma^2)$$

$$p(x|w_2) \sim N(\mu_2, \sigma^2)$$

$$d' = \frac{|\mu_2 - \mu_1|}{\sigma}$$

- تمایز پذیری

discriminability

\* مطلوب است که  $d'$  بزرگ باشد.

دقیقی  $\mu_1$ ،  $\mu_2$ ،  $\sigma$  و یا  $x^*$  را ندانیم فرض می‌کنیم که حالت طبعیت و تقسیم سیستم را می‌دانیم. این اطلاعات کمک می‌کند

$d'$  را بیابیم.

چرا احتمال زیر را در نظر می‌گیریم:

$$\text{hit} \quad P(x > x^* | x \in w_2)$$

احتمال اینکه سیگنال داخلی بالاتر از  $x^*$  باشد و داشته باشیم سیگنال خارجی موجود است.

$$\text{false alarm} \quad P(x > x^* | x \in w_1)$$

احتمال اینکه سیگنال داخلی بالاتر از  $x^*$  باشد، علیرغم اینکه سیگنال خارجی حاضر نباشد.

$$\text{miss} \quad P(x < x^* | x \in w_2)$$

احتمال اینکه سیگنال داخلی پایین‌تر از  $x^*$  باشد، و داشته باشیم سیگنال خارجی موجود است.

$$\text{correct rejection} \quad P(x < x^* | x \in w_1)$$

احتمال اینکه سیگنال داخلی پایین‌تر از  $x^*$  باشد، علیرغم اینکه سیگنال خارجی حاضر نباشد.

$\left. \begin{array}{l} \text{سیگنال داخلی بالاتر از } x \\ \text{با یابگین } \mu_1 \text{ وقتی سیگنال خارجی موجود باشد (کلاس } w_1) \\ \text{با یابگین } \mu_2 \text{ وقتی سیگنال خارجی موجود باشد (کلاس } w_2) \end{array} \right\}$   
 $x^*$  مقدار آستانه برای تشخیص وجود/عدم وجود سیگنال خارجی + فوئر تصادفی

Subject:

Year . Month . Date . ( )

بمقدار زیادی آزمایش انجام می دهیم (بافرض ثابت)

اما این احتمالات را به طور تجربی به دست آوریم: به طور خاص  $hit$  و  $false\ alarm$

این نقاط را بر روی یک نمودار رسم می کنیم (نمودار دوبعدی)

اگر چنانچه ثابت باشند اما آستانه  $\theta$  تغییر باشد، نرخ های  $hit$  و  $false\ alarm$  نیز تغییر می کند.

↔ برای مقدار تکلیک پذیری  $d'$  داده شده، نقطه ای مادر راستای یک منحنی هوار (ROC) حرکت می کند.

مثال:  $N$  داده‌ی تولید شده با pdf گاوسی یک بعدی داریم که میانگین آن  $\mu$  و واریانس آن مجهول است.

$$x_1, x_2, \dots, x_N$$

MLE واریانس را بدست آورید:

حل: تابع log-likelihood عبارت است از:

$$L(\sigma^2) = \ln \prod_{k=1}^N p(x_k; \sigma^2) = \sum_{k=1}^N \ln p(x_k; \sigma^2)$$

$$= \sum_{k=1}^N \ln \left( \frac{1}{\sqrt{2\pi\sigma^2}} \exp \left( -\frac{(x_k - \mu)^2}{2\sigma^2} \right) \right)$$

$$= -\frac{N}{2} \ln(2\pi\sigma^2) - \frac{1}{2\sigma^2} \sum_{k=1}^N (x_k - \mu)^2$$

$$\frac{\partial L(\sigma^2)}{\partial \sigma^2} = 0 \Rightarrow -\frac{N}{2\sigma^2} + \frac{1}{2\sigma^4} \sum_{k=1}^N (x_k - \mu)^2 = 0 \Rightarrow \sigma^2 = \hat{\sigma}_{ML}^2 = \frac{1}{N} \sum_{k=1}^N (x_k - \mu)^2$$

این برآورد برای  $N$  مشاهده‌ای، اریب (biased) است زیرا:

$$E\{\hat{\sigma}_{ML}^2\} = \frac{1}{N} \sum_{k=1}^N E\{(x_k - \mu)^2\} = \frac{N-1}{N} \sigma^2 \neq \sigma^2 = \text{واریانس واقعی تابع چگالی}$$

الاقته  $N \rightarrow \infty$ ، نااریب می‌شود:

$$\lim_{N \rightarrow \infty} E\{\hat{\sigma}_{ML}^2\} = \left(1 - \frac{1}{N}\right) \sigma^2 = \sigma^2$$

بردارهای  $x_1, x_2, \dots, x_N$  که از یک توزیع نرمال با ماتریس کوواریانس معلوم  $\Sigma$  و میانگین مجهول بیرون کشیده شده اند را در نظر بگیرید:

$$p(x_k; \mu) = \frac{1}{(2\pi)^{1/2} |\Sigma|^{1/2}} \exp\left(-\frac{1}{2}(x_k - \mu)^T \Sigma^{-1}(x_k - \mu)\right)$$

برآورد ML بردار میانگین مجهول را بیابید.

حل: برای  $N$  نمونه بی وجود داریم:

$$\begin{aligned} L(\mu) &= \ln \prod_{k=1}^N p(x_k; \mu) = \sum_{k=1}^N \ln p(x_k; \mu) \\ &= -\frac{N}{2} \ln((2\pi)^1 |\Sigma|) - \frac{1}{2} \sum_{k=1}^N (x_k - \mu)^T \Sigma^{-1} (x_k - \mu) \end{aligned}$$

$$\frac{\partial L(\mu)}{\partial \mu} = \begin{bmatrix} \frac{\partial L}{\partial \mu_1} \\ \frac{\partial L}{\partial \mu_2} \\ \vdots \\ \frac{\partial L}{\partial \mu_l} \end{bmatrix} = \sum_{k=1}^N \Sigma^{-1} (x_k - \mu) = 0 \Rightarrow \mu = \frac{1}{N} \sum_{k=1}^N x_k = \hat{\mu}_{ML}$$

$$A = A^T \Rightarrow \frac{\partial (\alpha^T A \alpha)}{\partial \alpha} = 2A\alpha \quad \text{تذکر:}$$

پس MLE برای میانگین در یکسانی گاوسی، میانگین نمونه ای است.

این برآورد طبیعی برای توابع یکگالی غیر گاوسی لزوماً بهترین ML نیست.

مثال: بردارهای  $x_1, x_2, \dots, x_N$  از یک توزیع نرمال با پارامترهای  $\mu$  و  $\sigma^2$  معلوم  $\Sigma$  میانگین مجهول  $\mu$  گرفته شده اند. توزیع بردار میانگین مجهول  $\mu$  به صورت زیر است:

$$p(\mu) = \frac{1}{(2\pi)^{1/2} \sigma_\mu^2} \exp\left(-\frac{1}{2} \frac{\|\mu - \mu_0\|^2}{\sigma_\mu^2}\right)$$

برآورد MAP به صورت زیر است:

$$\frac{\partial}{\partial \mu} \ln\left(\prod_{k=1}^N p(x_k|\mu) p(\mu)\right) = 0$$

برای  $\Sigma = \sigma^2 I$  داریم:

$$\sum_{k=1}^N \frac{1}{\sigma^2} (x_k - \hat{\mu}) - \frac{1}{\sigma_\mu^2} (\hat{\mu} - \mu_0) = 0 \Rightarrow \hat{\mu} = \hat{\mu}_{MAP} = \frac{\mu_0 + \frac{\sigma_\mu^2}{\sigma^2} \sum_{k=1}^N x_k}{1 + \frac{\sigma_\mu^2}{\sigma^2} N}$$

شاهد می شود که:

(۱) اگر  $\frac{\sigma_\mu^2}{\sigma^2} \gg 1$  باشد، یعنی واریانس  $\sigma_\mu^2$  خیلی زیاد است  $\Leftarrow$  کمبودی تناظر بسیار عریض (با تغییرات کم روی بازه‌ی مورد نظر) می باشد  $\Leftarrow$

$$\hat{\mu}_{MAP} \approx \hat{\mu}_{ML} = \frac{1}{N} \sum_{k=1}^N x_k$$

(۲) برای  $N \rightarrow \infty$  بدون توجه به مقدار واریانس، برآورد به میانگین تبدیل می شود.

(۳) برآورد MAP به طور مجانبی به برآورد ML میل می کند:

$$N \rightarrow \infty \Rightarrow \hat{\theta}_{MAP} \rightarrow \hat{\theta}_{ML}$$

(۴) اگر  $\frac{\sigma_\mu^2}{\sigma^2} \ll 1$  باشد، و  $N$  کوچک باشد، پارامتر به سمت پارامتر پیشین میل می کند:

$$\hat{\theta}_{MAP} \rightarrow \theta_0$$

(مربوط به توزیع پیشین تغییر)

## Bayesian Inference (BI)

استنتاج بیزی

روش‌های ML و MAP یک برآورد خاص از بردار پارامتر مجهول  $\theta$  را ارائه می‌دهند.

لازمه BI می‌گیرد و دنبال می‌شود:

با داشتن مجموعه‌ی  $X$  از  $N$  بردار آموزشی و

اطلاعات پیشین در مورد pdf  $p(\theta)$ :

هدف: محاسبه‌ی  $p(x|X)$  است.

$$p(x|X) = \int p(x|\theta) p(\theta|X) d\theta \quad (70) \quad \text{داریم:}$$

$$p(\theta|X) = \frac{p(X|\theta) p(\theta)}{p(X)} = \frac{p(X|\theta) p(\theta)}{\int p(X|\theta) p(\theta) d\theta} \quad (71)$$

$$p(X|\theta) = \prod_{k=1}^N p(x_k|\theta) \quad (72)$$

توزیع شرطی  $p(\theta|X)$  نیز به عنوان برآورد pdf پسین معلوم است:

(زیرا "دانش" به روز شده در مورد خواص آماری  $\theta$  پس از مشاهده‌ی مجموعه داده‌ی  $X$  است).

محاسبه‌ی  $p(x|X)$  نیاز به آنکوال گیری دارد (70): عمدتاً به صورت تقریب عددی قابل محاسبه است.

راه‌های: اگر  $p(\theta|X)$  معلوم باشد، آن‌گاه  $p(x|X)$  یا کمین  $p(x|\theta)$  نسبت به  $\theta$  است:

$$p(x|X) = E_{\theta} \{p(x|\theta)\}$$

با فرض بزرگ بودن تعداد نمونه‌های  $\{\theta_i\}_{i=1}^L$  از بردار تصادفی  $\theta$ ، می‌توانیم مقادیر قساطر  $p(x|\theta_i)$  را محاسبه کنیم و امید را به عنوان مقدار متوسط تقریب بزنیم:

$$p(x|X) \approx \frac{1}{L} \sum_{i=1}^L p(x|\theta_i)$$

انگزن مسئله تبدیل شده به تولید یک مجموعه نمونه  $\{\theta_i\}_{i=1}^L$

مثلاً اگر  $p(\theta|X)$  گادوسی باشد، می‌توان با مولد اعداد شبه تصادفی گادوسی،  $L$  نمونه تولید کرد.

شکل اینجاست که عموماً بزنم دقیق  $p(\theta|X)$  معلوم نیست. به حل شکل با مجموعه روش‌های MCMC

مثال:

$$p(x|\mu) \sim N(\mu, \sigma^2)$$

$$p(\mu) \sim N(\mu_0, \sigma_0^2) \quad \mu \text{ پارامتر مجهول}$$

داریم:

$$p(\mu|X) = \frac{p(X|\mu)p(\mu)}{p(X)} = \frac{1}{\alpha} \prod_{k=1}^N p(x_k|\mu)p(\mu)$$

که در آن  $\alpha = P(X)$  با توجه به مجموع داده‌ی آموزشی داده شده ثابت است.

$$p(\mu|X) = \frac{1}{\alpha} \prod_{k=1}^N \frac{1}{\sigma\sqrt{2\pi}} \exp\left(-\frac{(x_k - \mu)^2}{\sigma^2}\right) \frac{1}{\sigma_0\sqrt{2\pi}} \exp\left(-\frac{(\mu - \mu_0)^2}{2\sigma_0^2}\right)$$

$$= \frac{1}{\sigma_N\sqrt{2\pi}} \exp\left(-\frac{(\mu - \mu_N)^2}{2\sigma_N^2}\right) \quad \mu_N = \frac{N\sigma_0^2\bar{x}_N + \sigma^2\mu_0}{N\sigma_0^2 + \sigma^2}$$

$$\sigma_N^2 = \frac{\sigma^2\sigma_0^2}{N\sigma_0^2 + \sigma^2}$$

$$\bar{x}_N = \frac{1}{N} \sum_{k=1}^N x_k$$

با داشتن  $N$  نمونه،  $p(\mu|X)$  نیز گاوسی خواهد بود (نیاز به اثبات دارد).دقیق  $p(\mu|X)$  محاسبه شده، با جایگزینی در ۷۰ خواهیم داشت:

$$p(x|X) = \frac{1}{\sqrt{2\pi(\sigma^2 + \sigma_N^2)}} \exp\left(-\frac{1}{2} \frac{(x - \mu_N)^2}{\sigma^2 + \sigma_N^2}\right) \sim N(\mu_N, \sigma^2 + \sigma_N^2)$$

## Maximum Entropy Estimation (ME)

برآورد ماکزیمم آنتروپی

آنتروپی : معیاری برای عدم اطمینان یک رویداد است.  
 معیار تصادفی بودن پیام‌ها که در خروجی یک سیستم رخ می‌دهد. (پیام = بردار ویژگی)

$$H = - \int_{\mathcal{X}} p(x) \ln p(x) dx$$

برآورد ME برای pdf مجهول  $p(x)$  : آن که آنتروپی را ماکزیمم می‌کند (با دقت داشتن قیدهای داده شده).  
 \* مناظر با توزیعی که بالاترین تصادفی بودن ممکن را (دنبت به قیدهای موجود) نشان می‌دهد.

مثال : متغیر تصادفی  $x$  برای  $x_1 \leq x \leq x_2$  غیر منفی و در سایر موارد صفر است.  
 برآورد ماکزیمم آنتروپی را برای pdf آن ماکس کنید :

$$\max H = \max - \int p(x) \ln p(x) dx$$

$$\text{subject to } \int_{x_1}^{x_2} p(x) dx = 1$$

با استفاده از ضرایب لاگرانژ داریم :

$$H_L = - \int_{x_1}^{x_2} p(x) (\ln p(x) - \lambda) dx$$

$$\frac{\partial H_L}{\partial p(x)} = - \int_{x_1}^{x_2} [(\ln p(x) - \lambda) + 1] dx = 0 \Rightarrow p(x) = \exp(\lambda - 1) \triangleq \hat{p}(x)$$

برای ماکسیمی  $\lambda$ ، آن را در معادله قید جایگذاری می‌کنیم :

$$\exp(\lambda - 1) = \frac{1}{x_2 - x_1}$$

پس

$$\hat{p}(x) = \begin{cases} \frac{1}{x_2 - x_1} & , x_1 \leq x \leq x_2 \\ 0 & , \text{otherwise} \end{cases}$$

\* برآورد ماکزیمم آنتروپی برای pdf مجهول، توزیع یکواخت است.

نکته : اگر قیدهای دوم و سوم را برابر  $\mu$  و  $\sigma^2$  مشخص قرار می‌دهیم، برآورد ماکزیمم آنتروپی، تابع گاوسی می‌شود.

fig. 2.17

شال:

 $N=100$  نقطه در فضای دوبعدی

از توزیع چندمنه استخراج شده اند: multimodal

نمونه ها با دو بولد تقادین کادسی تولید شده اند:

$$N(\mu_1, \Sigma_1) \quad N(\mu_2, \Sigma_2) \quad \mu_1 = \begin{bmatrix} 1 \\ 1 \end{bmatrix} \quad \mu_2 = \begin{bmatrix} 2 \\ 2 \end{bmatrix} \quad \Sigma_1 = \Sigma_2 = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}$$

در هر زمان یک نمونه  $x_k$  ( $k=1, 2, \dots, N$ ) تولید می شود.

$$\begin{cases} P(H) = P = 0.8 \rightarrow \textcircled{1} \\ P(T) = 1 - P = 0.2 \rightarrow \textcircled{2} \end{cases}$$

یک سکه انداخته می شود تا مشخص شود  $x_k$  باید توسط کدام کادسی تولید شود:

(پس داده ها با مرکزیت  $\begin{bmatrix} 1 \\ 1 \end{bmatrix}$  چکان تر از داده ها با مرکزیت  $\begin{bmatrix} 2 \\ 2 \end{bmatrix}$  است.)

پس:

$$p(x) = g(x; \mu_1, \sigma_1^2) P + g(x; \mu_2, \sigma_2^2) (1-P)$$

$$\Sigma = \sigma^2 I = \text{diag}(\sigma^2) \quad \text{pdf کادسی با پارامترهای } \mu \text{ و ماتریس کوارایانس}$$

هدف: یافتن MLE با بردار پارامترهای مجهول

$$\Theta^T = [P, \mu_1^T, \sigma_1^2, \mu_2^T, \sigma_2^2]$$

برمبای  $N=100$  نقطه  $\{(x_k, z_k)\}_{k=1}^N$  که در آن  $z_k \in \{1, 2\}$ 

الای ها (اطلاعات بر حسب) موجود نیستند.

$$x_k = \alpha_k x_k^1 + (1 - \alpha_k) x_k^2$$

$$\swarrow N(\mu_1, \Sigma_1) \quad \searrow N(\mu_2, \Sigma_2)$$

$$L(\Theta; \alpha) = \sum_{k=1}^N \alpha_k \ln \{g(x_k; \mu_1, \sigma_1^2) P\} + \sum_{k=1}^N (1 - \alpha_k) \ln \{g(x_k; \mu_2, \sigma_2^2) (1-P)\}$$

$$E\{\alpha_k | x_k; \hat{\Theta}\} = 1 \times P(1 | x_k; \hat{\Theta}) + 0 \times (1 - P(1 | x_k; \hat{\Theta})) = P(1 | x_k; \hat{\Theta})$$

انتخاب مقادیر اولیه و استفاده از EM تا همگرا شدن.