

کد کنترل

528

A

528A

آزمون ورودی دوره دکتری (نیمه متمرکز) - سال ۱۴۰۰

دفترچه شماره (۱)

صبح جمعه

۹۹/۱۲/۱۵



جمهوری اسلامی ایران
وزارت علوم، تحقیقات و فناوری
سازمان سنجش آموزش کشور

«اگر دانشگاه اصلاح شود مملکت اصلاح می شود.»
امام خمینی (ره)

رشته مهندسی کامپیوتر - هوش مصنوعی - (کد ۲۳۵۶)

مدت پاسخ گویی: ۱۵۰ دقیقه

تعداد سؤال: ۴۵

عنوان مواد امتحانی، تعداد و شماره سؤالات

ردیف	مواد امتحانی	تعداد سؤال	از شماره	تا شماره
۱	مجموعه دروس تخصصی: - ساختمان داده ها و طراحی الگوریتم ها - شناسایی الگو - یادگیری ماشین	۴۵	۱	۴۵

استفاده از ماشین حساب مجاز نیست.

این آزمون نمره منفی دارد.

* داوطلب گرامی، عدم درج مشخصات و امضا در مندرجات جدول ذیل، به منزله عدم حضور شما در جلسه آزمون است.

اینجانب با شماره داوطلبی با آگاهی کامل، یکسان بودن شماره صندلی خود را با شماره داوطلبی مندرج در بالای کارت ورود به جلسه، بالای پاسخنامه و دفترچه سؤالات، نوع و کد کنترل درج شده بر روی دفترچه سؤالات و پائین پاسخنامه‌ام را تأیید می‌نمایم.

امضا:

۱- اگر به یک گراف جهت‌دار یک یال اضافه کنیم، تعداد اجزای قویا همبند در گراف چه مقدار تغییر می‌کند؟

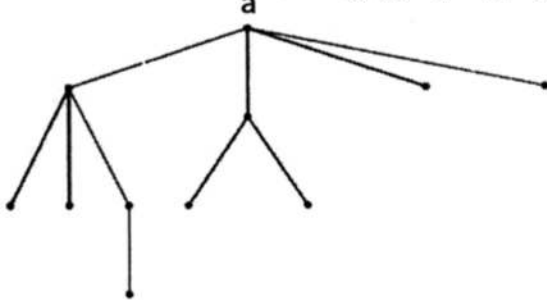
(۱) حداکثر یک واحد کمتر می‌شود. (۲) حداکثر دو واحد کمتر می‌شود.

(۳) ممکن است بیش از دو واحد کم شود. (۴) تغییر نمی‌کند یا ممکن است افزایش پیدا کند.

۲- فرض کنید شکل زیر، درخت حاصل از اجرای DFS روی یک گراف همبند و بدون جهت G است. با فرض این‌که

می‌دانیم گراف G دقیقاً یک رأس برشی دارد، کدام گزاره در مورد درجه رأس a در این گراف صحیح است؟

(یک رأس گراف G برشی است، اگر حذف آن تعداد مؤلفه‌های همبندی G را افزایش دهد.)



(۱) درجه رأس a در G هر عددی بین ۴ تا ۷ می‌تواند باشد.

(۲) درجه رأس a در G هر عددی بین ۸ و ۱۰ است.

(۳) درجه رأس a در G هر عددی بین ۹ و ۱۰ است.

(۴) درجه رأس a در G دقیقاً برابر ۴ است.

۳- آرایه‌ای شامل n عدد متمایز داده شده است. می‌خواهیم از روی این اعداد یک درخت دودویی بسازیم، با این

خاصیت که به ازای هر رأس درخت ساخته شده این خصوصیت را نیز داشته باشد که پیمایش میان ترتیب آن

دقیقاً معادل ترتیب عناصر در آرایه شود. کدام گزاره درست است؟

(۱) چنین درختی لزوماً به ازای هر آرایه وجود دارد، اما یکتا نیست.

(۲) چنین درختی لزوماً به ازای هر آرایه شامل n عدد متمایز وجود ندارد.

(۳) بهترین زمان برای ساخت چنین درختی از روی یک آرایه $O(n)$ است.

(۴) بهترین زمان برای ساخت چنین درختی از روی یک آرایه $O(n \log n)$ است.

۴- فرض کنید می‌خواهیم n تومان را با استفاده از سکه‌های a و b و c تومانی خرد کنیم. به ازای چه تعداد از

سه تایی‌های (a, b, c) زیر، الگوریتم حریصانه، n تومان را با کمترین تعداد سکه خرد می‌کند؟

• (۵, ۲, ۱)

• (۵, ۴, ۱)

• (۶, ۳, ۱)

(۴) ۳

(۳) ۲

(۲) ۱

(۱) ۰

۵- آرایه A حاصل ترکیب دو زیر آرایه B و C است که B یک آرایه صعودی و C یک آرایه نزولی است. به‌عنوان

نمونه، A می‌تواند به‌صورت $[۱, ۲, ۳, ۴, ۵, ۶, ۷, ۸, ۹]$ باشد، که در واقع ترکیب آرایه صعودی $[۲, ۴, ۶, ۸]$ و آرایه

نزولی $[۹, ۷, ۵, ۱]$ است. این آرایه را در چه زمانی می‌توان مرتب کرد؟ (بهترین گزینه را انتخاب کنید.)

(۲) $O(n \log n)$

(۱) $O(n)$

(۴) $O(n \log \log n)$

(۳) $O(n / \log n)$

۶- فرض کنید یک آرایه مرتب از اعداد طبیعی به طول n داریم، که در آن هر عدد به غیر از یکی دقیقاً دو بار ظاهر شده است. عضو غیر تکراری را در چه زمانی می‌توان پیدا کرد؟

$$O(n \log n) \quad (۱) \quad O(\log n) \quad (۲)$$

$$O(n) \quad (۳) \quad O(1) \quad (۴)$$

۷- یک هرم کمینه شامل n عنصر داده شده است. سومین کوچک‌ترین عنصر این آرایه را در چه زمانی می‌توان پیدا کرد؟

$$O(1) \quad (۱) \quad O(n) \quad (۲)$$

$$O(\log n) \quad (۳) \quad O(n \log n) \quad (۴)$$

۸- اگر در الگوریتم هافمن نویسه‌ای بیش از $\frac{2}{5}$ کل متن تکرار شود، در آن صورت کد این نویسه چند بیت می‌تواند باشد؟

$$(۱) \text{ هر عددی بین } ۱ \text{ تا } ۳ \quad (۲) \text{ فقط } ۱$$

$$(۳) \text{ فقط } ۲ \quad (۴) \text{ } ۱ \text{ یا } ۲$$

۹- در یک درخت T با n گره، فرض کنید تعداد برگ‌ها B و تعداد فرزندان هر گره غیربرگ ۲ باشد. همچنین فرض کنید $E[T]$ و $I[T]$ به ترتیب مجموع عمق برگ‌ها و مجموع عمق عناصر غیربرگ T باشند. اگر $n = ۹۹۹۹$ باشد، کدام گزینه همیشه درست است؟

$$B = ۵۰۰۰۰ \quad (۱) \quad E[T] = ۹۹۹۹۹ \quad (۲)$$

$$I[T] = ۹۹۹۹۸ \quad (۳) \quad E[T] - I[T] = ۱۰۰۰۰۰ \quad (۴)$$

۱۰- چند تا از گزاره‌های زیر صحیح است؟

- اگر پس از اتمام الگوریتم بلمن فورد، به‌روزرسانی فاصله‌ها را ادامه دهیم و فاصله مربوط به یک رأس v باز هم به روز شود، v در یک دور منفی قرار دارد.

- اگر در جست‌وجوی عمق اول گراف جهت‌دار G تنها یک یال بازگشتی (back edge) e وجود داشته باشد، آنگاه یال e' به جز یال e وجود دارد که $e' - G$ بدون دور است.

- اگر در الگوریتم دایکسترا که روی یک DAG و از مبدأ اجرا شده، فقط برخی از یال‌های خروجی s وزن منفی داشته باشند، الگوریتم ممکن است به درستی فاصله‌ها را محاسبه نکند.

$$(۱) \quad ۰ \quad (۲) \quad ۱$$

$$(۳) \quad ۲ \quad (۴) \quad ۳$$

۱۱- در یک مجموعه از اعداد صحیح به اندازه n دنبال یک ۴ تایی‌هایی مثل (a, b, c, d) هستیم که

$$d = \sqrt{a^2 + b^2 + c^2}. \text{ این کار در چه زمانی قابل انجام است؟ (بهترین گزینه را انتخاب کنید).}$$

$$O(n^3) \quad (۱) \quad O(n^4) \quad (۲)$$

$$O(n \log n) \quad (۳) \quad O(n^2 \log n) \quad (۴)$$

۱۲- با فرض درهم‌سازی یکنواخت ساده، احتمال این‌که سه عضو متفاوت a و b و c به یک خانه جدول درهم‌سازی

نگاشت شوند برابر کدام گزینه است؟ (فرض کنید اندازه جدول m است).

$$(۱) \quad \sqrt{m} \quad (۲) \quad \sqrt{m^2}$$

$$(۳) \quad \sqrt{m^3} \quad (۴) \quad \text{بستگی به مقادیر } a \text{ و } b \text{ و } c \text{ دارد.}$$

۱۳- خانواده $H = \{h_1, \dots, h_k\}$ از توابع درهم ساز را در نظر بگیرید که $h_i : \{a, b, c, d\} \rightarrow \{0, 1\}$ ، برای آن که این خانواده یک خانواده درهم ساز سراسری باشد، k حداقل چقدر باید باشد؟ (خانواده توابع H سراسری است اگر و فقط

اگر به ازای هر دو مقدار u و v داشته باشیم. $\Pr_{h \in H}[h(u) = h(v)] \leq \frac{1}{m}$ که m اندازه جدول درهم سازی است.)

(۱) ۱۶ (۲) ۴

(۳) ۲ (۴) ۱

۱۴- آرایه $A[1..n]$ از اعداد صحیح داده شده است. زیر دنباله متوالی $A[i..j]$ یک «بازه مثبت» نامیده می شود، اگر جمع اعضای $A[i]$ تا $A[j]$ مثبت (بزرگ تر از ۰) باشد. می خواهیم کمترین تعداد بازه های مثبت که تمام اعداد مثبت آرایه را پوشش می دهد پیدا کنیم. اگر ورودی آرایه زیر باشد، جواب کدام است؟

$A[1..15] = \langle 3, -5, 4, 1, -9, -8, 2, 3, 4, -10, 1, -2, -3, 6, -1 \rangle$

(۱) ۵ (۲) ۴

(۳) ۳ (۴) ۲

۱۵- یک گراف ۵ رأسی همبند و بدون جهت داریم که رأس های آن با شماره های ۱ تا ۵ شماره گذاری شده اند. فرض کنید از رأس ۱ الگوریتم BFS را اجرا می کنیم و تمام حالت هایی که BFS می تواند رؤس را ملاقات کند عبارتند از $\langle 1, 2, 3, 4, 5 \rangle$ ، $\langle 1, 3, 2, 4, 5 \rangle$ ، $\langle 1, 3, 2, 5, 4 \rangle$ و $\langle 1, 2, 3, 5, 4 \rangle$. حال اگر از رأس ۵ الگوریتم DFS را اجرا کنیم، کدام گزینه نمی تواند ترتیب ملاقات ها رؤس گراف باشد؟

(۱) ۵, ۴, ۳, ۱, ۲ (۲) ۵, ۴, ۳, ۲, ۱

(۳) ۵, ۳, ۴, ۱, ۲ (۴) ۵, ۴, ۱, ۲, ۳

۱۶- برنامه زیر چه کاری می کند و زمان اجرای آن کدام است؟

SS(A[0 .. n-1])

If $n = 2$ and $A[0] > A[1]$ then

Swap (A[0] , A[1])

else if $n > 2$

$m = \lceil 2n/3 \rceil$

SS(A[0 .. m-1])

SS(A[n-m .. n-1])

SS(A[0 .. m-1])

(۱) آرایه A را مرتب می کند و زمان اجرای آن $\theta(n^{\log_{3/2} 3})$ است.

(۲) آرایه A را مرتب می کند و زمان اجرای آن $\theta(n^{\log_{2/3} 3})$ است.

(۳) آرایه A را لزوماً مرتب نمی کند اما زمان اجرای آن $\theta(n^{\log_{3/2} 3})$ است.

(۴) آرایه A را لزوماً مرتب نمی کند اما زمان اجرای آن $\theta(n^{\log_{2/3} 3})$ است.

۱۷- دنباله $(۳, ۱, ۴, ۱, ۵, ۹, ۲, ۶, ۵, ۳, ۸, ۹, ۷, ۹, ۳, ۲, ۳, ۸, ۴, ۶, ۲, ۷, ۹)$ را در نظر بگیرید. چند عضو متوالی این دنباله را می توان به صورت یک عدد تصور کرد. مثلاً سه عنصر متوالی ۵ و ۳ و ۸ را عدد ۵۳۸ تصور کرد. دو عدد به این شکل را مجزا گوییم، اگر هیچ یک از عناصر دنباله در ساخت هر دوی آن ها نقش نداشته باشند. حداکثر چند عدد مجزا به این شکل می توان ساخت که به ترتیب از چپ به راست صعودی باشند؟

- (۱) ۱۱ (۲) ۱۰ (۳) ۹ (۴) ۸

۱۸- فرض کنید تابعی داریم که به عنوان ورودی دو دنباله گرفته و به عنوان خروجی طول بزرگ ترین زیر دنباله مشترک آن ها را برمی گرداند. با حداکثر یک بار فراخوانی این تابع به علاوه هزینه $O(n)$ چند مورد زیر را می توان محاسبه کرد؟

- محاسبه طول بزرگ ترین زیر دنباله آینه ای یک دنباله
 - تشخیص این که آیا یک دنباله زیر دنباله یک دنباله دیگر است.
 - تشخیص این که آیا یک دنباله آینه ای است.
- (۱) ۰ (۲) ۱ (۳) ۲ (۴) ۳

۱۹- درخت فراگیر T از گراف وزن دار G یک درخت گلوگاهی است، اگر سنگین ترین یال آن در بین تمامی درخت های فراگیر G سبک ترین باشد، کدام گزینه در خصوص گزاره های زیر درست است؟

- (الف) هر درخت گلوگاهی یک درخت فراگیر کمینه است.
(ب) هر درخت فراگیر کمینه یک درخت گلوگاهی است.
- (۱) (الف) درست - (ب) درست (۲) (الف) درست - (ب) نادرست
(۳) (الف) نادرست - (ب) درست (۴) (الف) نادرست - (ب) نادرست

۲۰- گراف جهت دار G با وزن یال های صحیح و دو رأس خاص s و t از گراف داده شده است. فرض کنید شار بیشینه از s به t در گراف داده شده است. کدام گزینه در خصوص گزاره های زیر درست است؟

- (الف) اگر ظرفیت یکی از یال های G یک واحد افزایش داده شود، شار بیشینه در گراف جدید در زمان خطی قابل محاسبه است.
(ب) اگر ظرفیت یکی از یال های G یک واحد کاهش داده شود، شار بیشینه در گراف جدید در زمان خطی قابل محاسبه است.

- (۱) (الف) درست - (ب) درست (۲) (الف) درست - (ب) نادرست
(۳) (الف) نادرست - (ب) درست (۴) (الف) نادرست - (ب) نادرست

۲۱- کدام یک از روش های کاهش ابعاد زیر از تکنیک «کشف ساختار درونی با بعد کمتر» یا «intrinsic low-dimensional structure detection» برای کاهش ابعاد داده ها استفاده می نماید؟

- (۱) ISOMAP
(۲) Locally Linear Embedding (LLE)
(۳) Multidimensional Scaling (MDS)
(۴) Principal Components Analysis (PCA)

۲۲- از معیار واگرایی (Divergence) برای اندازه گیری میزان جدایی پذیری دو کلاس می توان استفاده کرد. برای حالتی که توابع توزیع ویژگی x در دو کلاس w_i و w_j به صورت توابع گوسی یک متغیره $p(x | w_i) \sim N(\mu_i, \sigma_i^2)$ و $p(x | w_j) \sim N(\mu_j, \sigma_j^2)$ باشد، معیار واگرایی d_{ij} از کدام رابطه به دست می آید؟ نکته: معیار واگرایی d_{ij} توسط رابطه زیر محاسبه می شود.

$$d_{ij} = \int_{-\infty}^{\infty} (p(x | w_i) - p(x | w_j)) \ln \left(\frac{p(x | w_i)}{p(x | w_j)} \right) dx$$

$$d_{ij} = \frac{1}{2} \left(\frac{1}{\sigma_i^2} + \frac{1}{\sigma_j^2} - 2 \right) + \frac{1}{2} (\mu_i - \mu_j)^2 \left(\frac{\sigma_j^2}{\sigma_i^2} + \frac{\sigma_i^2}{\sigma_j^2} \right) \quad (۱)$$

$$d_{ij} = \frac{1}{2} (\mu_i - \mu_j)^2 \left(\frac{\sigma_j^2}{\sigma_i^2} + \frac{\sigma_i^2}{\sigma_j^2} - 2 \right) + \frac{1}{2} \left(\frac{1}{\sigma_i^2} + \frac{1}{\sigma_j^2} \right) \quad (۲)$$

$$d_{ij} = \frac{1}{2} (\mu_i - \mu_j)^2 \left(\frac{\sigma_j^2}{\sigma_i^2} + \frac{\sigma_i^2}{\sigma_j^2} \right) + \frac{1}{2} \left(\frac{1}{\sigma_i^2} + \frac{1}{\sigma_j^2} - 2 \right) \quad (۳)$$

$$d_{ij} = \frac{1}{2} \left(\frac{\sigma_j^2}{\sigma_i^2} + \frac{\sigma_i^2}{\sigma_j^2} - 2 \right) + \frac{1}{2} (\mu_i - \mu_j)^2 \left(\frac{1}{\sigma_i^2} + \frac{1}{\sigma_j^2} \right) \quad (۴)$$

۲۳- چنانچه مجموعه داده های زیر برای دسته بندی دودسته ای در اختیار باشند، نمونه آزمایشی $x = 5$ با استفاده از دسته بند بیز و با فرض توزیع گاوس برای هر دسته و دسته بند NN-1، به ترتیب از راست به چپ به کدام دسته تعلق خواهد داشت؟

$$C_1 = \{1, 3\}$$

$$C_7 = \{6, 7, 8, 9, 10\}$$

(۴) دو و دو

(۳) دو و یک

(۲) یک و دو

(۱) یک و یک

۲۴- اگر (x_1, x_2) به طور توأم دارای توزیع نرمال با میانگین $\mu = \begin{bmatrix} \mu_1 \\ \mu_2 \end{bmatrix}$ و ماتریس کوواریانس $\Sigma = \begin{bmatrix} \sigma_1^2 & \sigma \\ \sigma & \sigma_2^2 \end{bmatrix}$ باشد، کدام گزینه میانگین توزیع شرطی $P(x_1 | x_2 = a)$ را نشان می دهد؟

$$\mu_1 + \frac{\sigma^2}{\sigma_2^2} (a - \mu_2) \quad (۱)$$

$$\mu_1 + \frac{\sigma}{\sigma_2^2} (a - \mu_2) \quad (۲)$$

$$\mu_1 + \frac{\sigma}{\sigma_2^2} (a - \mu_2) \quad (۳)$$

$$\mu_1 + \frac{\sigma_1}{\sigma_2} (a - \mu_2) \quad (۴)$$

۲۵- در جدول زیر با فرض اینکه متغیرهای باینری X و Y مستقل باشند، مقادیر p_1 و p_2 مربوط به احتمال توأم $P(X, Y)$ کدام است؟

X	Y	P(X, Y)
+	+	$\frac{3}{5}$
+	-	$\frac{1}{5}$
-	+	$p_1 = ?$
-	-	$p_2 = ?$

$$p_2 = \frac{1}{20} \text{ و } p_1 = \frac{3}{20} \quad (۱)$$

$$p_2 = \frac{2}{15} \text{ و } p_1 = \frac{1}{15} \quad (۲)$$

$$p_2 = \frac{3}{20} \text{ و } p_1 = \frac{1}{20} \quad (۳)$$

$$p_2 = \frac{1}{15} \text{ و } p_1 = \frac{2}{15} \quad (۴)$$

۲۶- چهار نمونه که هر کدام دارای ۲ ویژگی هستند و به صورت ماتریس X نشان داده شده است.

$$X = \begin{bmatrix} 6 & -4 \\ -3 & 5 \\ -2 & 6 \\ 7 & -3 \end{bmatrix}$$

می‌خواهیم این داده‌ها را توسط الگوریتم PCA به یک بعد کاهش بدهیم، کدام بردار تبدیل می‌تواند برای این منظور به کار رود؟

$$\left(\frac{1}{\sqrt{2}}, -\frac{1}{\sqrt{2}} \right)^T \quad (۲)$$

$$\left(\frac{1}{\sqrt{2}}, \frac{1}{\sqrt{2}} \right)^T \quad (۱)$$

$$\left(-\frac{1}{\sqrt{2}}, -\frac{1}{\sqrt{2}} \right)^T \quad (۴)$$

$$\left(-\frac{1}{\sqrt{2}}, \frac{1}{\sqrt{2}} \right)^T \quad (۳)$$

۲۷- یک مجموعه آموزشی داریم که نمونه‌های آن برچسب‌های ۱+ و ۲+ دارند و هر نمونه دارای m ویژگی با مقدار حقیقی است. برای هر مقدار توزیع گاوسی برای ویژگی در نظر می‌گیریم، پس $P(x_k | c_j) = N(\mu_{jk}, \sigma_{jk}^2)$ که μ_{jk} میانگین ویژگی k برای دسته j و σ_{jk}^2 واریانس ویژگی k برای دسته j است. به این دسته‌بند، دسته‌بند بیز ساده گاوسی می‌گویند، کدام گزینه در مورد این دسته‌بند نادرست است؟

(۱) اگر $\sigma_{1k} = \sigma_{2k}$ باشد، دسته‌بند خطی است.

(۲) قدرت تفکیک دسته‌بند خطی کمتر از قدرت تفکیک دسته‌بند بیز ساده گاوسی است.

(۳) بدون هیچ شرطی روی σ_{1k} و σ_{2k} ، دسته‌بند بیز ساده گاوسی، یک دسته‌بند خطی است.

(۴) دسته‌بند بیز ساده گاوسی و logistic regression در حالت مجانبی وقتی که شرط $\sigma_{1k} = \sigma_{2k}$ برقرار باشد، یک مدل را می‌سازند.

۲۸- در رگرسیون چند جمله‌ای همراه با منظم‌سازی، با افزایش λ کدام یک از گزاره‌های زیر درست است؟

$$E = (w | x) = \frac{1}{2} \sum_{i=1}^N [y^i - g(x^i | w)]^2 + \lambda \sum_i w_i^2$$

(۲) بایاس افزایش می‌یابد.

(۱) واریانس افزایش می‌یابد.

(۴) به داده‌های پرت اهمیت بیشتری داده می‌شود.

(۳) مدل دچار بیش برازش می‌شود.

۲۹- آزمونی داده شده است که پیش‌بینی می‌کند یک شخص بیمار است (+) یا بیمار نیست (-). این آزمون توسط تابع زیر نشان داده می‌شود.

$$p(y = c+ | x) = \begin{cases} 0 & \text{if } x < 0 \\ x & \text{if } 0 \leq x < 1 \\ 1 & \text{if } x \geq 1 \end{cases}$$

همچنین می‌دانیم که هزینه خطای FN سه برابر هزینه خطای FP است. ما می‌خواهیم مقدار x^* را پیدا کنیم که احتمال دسته‌بندی اشتباه و ریسک را کمینه نماید. در این خصوص کدام گزینه درست است؟ (مقدار سمت راست احتمال دسته‌بندی اشتباه و مقدار سمت چپ ریسک را نشان می‌دهد).

$$\begin{array}{ll} (1) \quad \frac{1}{2} \text{ و } \frac{1}{2} & (2) \quad \frac{1}{2} \text{ و } \frac{1}{4} \\ (3) \quad \frac{1}{4} \text{ و } \frac{1}{2} & (4) \quad \frac{1}{4} \text{ و } \frac{1}{4} \end{array}$$

۳۰- فرض کنید که x یک بردار تصادفی d -بعدی با میانگین μ و ماتریس کوواریانس Σ است. اگر $Y = Ax + b$ ، که A یک ماتریس $n \times d$ و b یک بردار n -بعدی باشد، ماتریس کوواریانس Y کدام است؟

$$\begin{array}{ll} (1) \quad A \sum A^T & (2) \quad A \sum^{-1} A^T \\ (3) \quad A^T \sum A & (4) \quad A^T \sum^{-1} A \end{array}$$

۳۱- مسئله دسته‌بندی برای داده‌های دو دسته $S=1$ و $S=2$ با احتمال پیشین برابر $p(S=1) = p(S=2)$ را در نظر بگیرید. ورودی دسته‌بند بردار $X = (x_1, x_2)^T$ با دو ویژگی نامنفی و مستقل از هم x_1 و x_2 است. در هر دسته ویژگی‌های x_1 و x_2 دارای تابع توزیع نمایی به صورت زیر هستند.

$$f_{x_k|S}(x_k | i) = \begin{cases} \lambda_{ik} e^{-\lambda_{ik} x_k} & , x_k \geq 0 \\ 0 & , x_k < 0 \end{cases}$$

پارامترهای توزیع برای این دو دسته مشخص و به صورت زیر است.

$$S=1: \begin{cases} \lambda_{11}=1 \\ \lambda_{12}=2 \end{cases} \quad S=2: \begin{cases} \lambda_{21}=2 \\ \lambda_{22}=1 \end{cases}$$

فرض کنید برای این مسئله، دسته‌بند بهینه خطی طراحی کرده‌ایم. احتمال دسته‌بندی صحیح به کمک دسته‌بند بهینه کدام است؟

$$\begin{array}{llll} (1) \quad \frac{1}{3} & (2) \quad \frac{1}{2} & (3) \quad \frac{2}{3} & (4) \quad \frac{1}{1+e} \end{array}$$

۳۲- کدام گزینه نادرست است؟

- (۱) بعد VC پرسپترون کوچک‌تر از بعد VC در SVM خطی است.
- (۲) احتمال رویداد بیش‌برازش (Overfitting)، برای مجموعه داده‌های کوچک، بیشتر است.
- (۳) اگر عمق یک درخت تصمیم به اندازه طول بردار ویژگی باشد، نویز موجود در داده‌ها بیشتر تأثیر خواهد گذاشت.
- (۴) اگر داده آموزشی شامل نویز در برجسب داده‌ها باشد، روش ادغام Boosting صدمه بیشتری نسبت به روش Bagging خواهد داشت.

۳۳- تعداد همسایه‌ها (یعنی k) در روش نزدیک‌ترین همسایه (KNN)، چه اثری دارد؟

- (۱) افزایش k باعث کاهش هزینه محاسباتی می‌شود.
- (۲) افزایش k باعث افزایش حساسیت به نویز ویژگی‌ها می‌شود.
- (۳) افزایش k باعث کاهش حساسیت به نویز دسته‌بندی می‌شود.
- (۴) کاهش k باعث کاهش تعداد ویژگی‌های مؤثر در دسته‌بندی می‌شود.

۳۴- کدام عبارت صحیح است؟

- (۱) Naïve Bayes همیشه یک دسته‌بند خطی است.
- (۲) یکی از مشکلات روش خوشه‌بندی سلسله‌مراتبی با معیار فاصله Single-link، مشکل Chain effect است.
- (۳) دسته بند A دقت 90% روی داده آموزشی و 75% روی داده آزمون و دسته‌بند B دقت 78% روی هر دو داده آموزشی و آزمون دارد. می‌توان نتیجه گرفت که دسته‌بند A از B بهتر است، چون میانگین دقت آن بالاتر است.
- (۴) یک مجموعه داده خطی جدایی‌پذیر داریم و دو پرسپترون که یکی با روش نزول در امتداد گرادیان خطا و دیگری با قانون پرسپترون آموزش داده شده است. می‌توان گفت که هر دو پرسپترون دقت یکسانی بر روی مجموعه داده آموزشی و آزمون دارند.

۳۵- در یک دسته‌بند ماشین بردار پشتیبان با حاشیه سخت (Hard Margin SVM) در فضای دوبعدی

$x = (x_1, x_2)^T$ مرز جداساز آن به وسیله قاعده $4z_1 + 9z_2 + 4z_3 + 16z_4 = 0$ مشخص شده است و از تابع زیر برای نگاشت بردار ویژگی به فضای با ابعاد بالاتر استفاده می‌کند.

$$z = (z_1, z_2, z_3, z_4)^T = \phi(x) = (x_1^2, x_2^2, x_1 x_2, -x_1)^T$$

آیا نمونه‌های $x = [+1, -1]^T$ و $y = [+1, -1]^T$ به ترتیب می‌توانند بردارهای پشتیبان این دسته‌بند باشند؟

- (۱) خیر - خیر
- (۲) خیر - بلی
- (۳) بلی - خیر
- (۴) بلی - بلی

۳۶- فرض کنید در یک مجموعه داده، ستون TID نشان‌دهنده کلید مربوط به نمونه‌ها باشد. (برای هر نمونه یک TID

منحصر بفرد وجود داشته باشد). اگر یک درخت تصمیم غیرباینری روی این داده، آموزش دهیم و متغیر TID را هم به عنوان یک ویژگی وارد الگوریتم یادگیری کنیم، در صورت استفاده از معیار بهره اطلاعاتی

(Information Gain) برای انتخاب ویژگی، کدام گزینه صحیح است؟

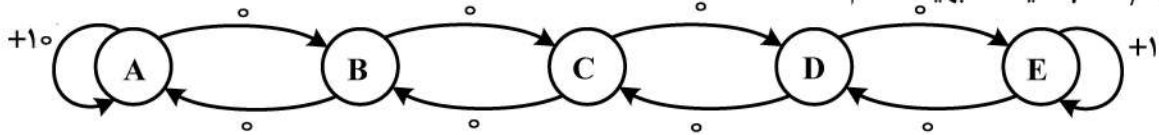
- (۱) درخت حاصل دارای عمق پایین و مشکل کم برازش (Underfitting) خواهد بود.
- (۲) درخت حاصل دارای عمق بالا و مشکل بیش برازش (Overfitting) خواهد بود.
- (۳) متغیر TID دارای بهره اطلاعاتی بالا بوده و به عنوان ریشه درخت انتخاب خواهد شد.
- (۴) در صورت مناسب بودن مابقی ویژگی‌ها، الگوریتم قادر به یادگیری درختی با قدرت تعمیم بالا خواهد بود.

۳۷- بعد VC یک درخت تصمیم با سه گره دودویی (گره‌ها دارای دو فرزند هستند) در R^1 کدام است؟

- (۱) ۵
- (۲) ۴
- (۳) ۳
- (۴) ۲

۳۸- یک عامل یادگیری تقویتی در محیط زیر که شامل ۵ حالت است، فعالیت می‌کند. اعداد روی کمان‌ها، پاداش را نشان می‌دهند. عامل تنها دو کنش چپ و راست که به صورت قطعی هستند را دارد.

برای $\gamma = 0.1$ سیاست بهینه کدام است؟



$$\pi^* = \{B : \text{left}, C : \text{right}, D : \text{right}\} \quad (2) \quad \pi^* = \{B : \text{left}, C : \text{left or right}, D : \text{right}\} \quad (1)$$

$$\pi^* = \{B : \text{left}, C : \text{left}, D : \text{right}\} \quad (4) \quad \pi^* = \{B : \text{left}, C : \text{left}, D : \text{left}\} \quad (3)$$

۳۹- یک کارخانه محصول A را تولید می‌کند. می‌خواهیم احتمال خراب بودن این محصول را تخمین بزنیم. برای این کار مجموعه‌ای از آزمایش‌ها را به صورت زیر انجام می‌دهیم. در هر آزمایش چندین محصول A را برداشته و بررسی می‌کنیم که آیا این محصول سالم است یا خراب، تا به یک محصول A سالم برسیم. در هر آزمایش تعداد محصول A آزمایش شده خراب که با k نشان می‌دهیم را یادداشت می‌کنیم. فرض کنید که احتمال خرابی محصول A برابر P باشد و m آزمایش مستقل انجام دهیم و تعداد محصولات خراب $k_1, k_2, \dots, k_m$ را یادداشت کنیم. تخمین P با استفاده از تخمین درست‌نمایی بیشینه کدام است؟

$$\frac{1}{m + \sum_{i=1}^m K_i} \quad (2) \quad m + 1 + \sum_{i=1}^m K_i \quad (1)$$

$$\frac{\sum_{i=1}^m K_i}{m + \sum_{i=1}^m K_i} \quad (4) \quad \frac{\sum_{i=1}^m K_i}{m + 1 + \sum_{i=1}^m K_i} \quad (3)$$

۴۰- مجموعه داده‌های زیر را داریم.

$$S = \{(1, 0), (2, 1)\}$$

مؤلفه اول هر زوج، بردار ویژگی x و مؤلفه دوم برچسب آن است. اگر از یک دسته‌بند **logistic regression** با پارامترهای $\omega_0 = -4$ و $\omega_1 = 3$ استفاده کنیم، مقدار **likelihood** کدام است؟

$$\begin{array}{ll} 0.27 \quad (1) & 0.64 \quad (2) \\ 0.73 \quad (3) & 0.88 \quad (4) \end{array}$$

۴۱- تابع توزیعی به صورت زیر داریم، اما مقدار ω را نداریم و می‌خواهیم با یک نمونه از توزیع مقدار آن (ω) را تخمین بزنیم. یک نمونه از توزیع گرفته می‌شود، که مقدار آن $x = 3$ است. تخمین درست‌نمایی بیشینه ω برابر کدام گزینه است؟

$$P(x) = \begin{cases} 0 & \text{if } x < 0 \\ \frac{2}{\omega} - \frac{2x}{\omega^2} & \text{if } 0 \leq x \leq \omega \\ 0 & \text{if } x > \omega \end{cases}$$

۳ (۴)

۴ (۳)

۵ (۲)

۶ (۱)

۴۲- فرض کنید فضای فرضیه (Hypothesis space) شامل درخت‌های تصمیم با عمق ۲ بر روی ۲ ویژگی است. همچنین فرض کنید هر داده شامل $n = 10$ ویژگی دودویی باشد. چه تعداد داده (m) کافی است تا مطمئن باشیم که یک درخت تصمیم با مشخصه ذکر شده می‌یابیم که به احتمال حداقل 0.99 خطای آن کمتر از 0.05 است؟

$$(1) m \geq 7 \quad (2) m \geq 100$$

$$(3) m \geq 160 \quad (4) m \geq 224$$

۴۳- یک ژن به شکل دو آلل A و a ظاهر می‌شود، به گونه‌ای که احتمال رخداد آلل A در جمعیت برابر با θ است. بدین معنی که یک کپی تصادفی از این ژن به احتمال θ آلل A و به احتمال $1-\theta$ آلل a است. یک ژنوتیپ از دو ژن تشکیل شده است و احتمال هر ژنوتیپ به صورت زیر است.

ژنوتیپ	AA	Aa	aa
احتمال	θ^2	$2\theta(1-\theta)$	$(1-\theta)^2$

فرض کنید که یک نمونه برداری تصادفی از جمعیت شامل k_1 ژنوتیپ AA، k_2 ژنوتیپ Aa و k_3 ژنوتیپ aa است. تخمین بیشینه درست‌نمایی (MLE) برای θ کدام است؟

$$(1) \hat{\theta} = \frac{2k_1 + k_2}{2k_1 + 2k_2 + 2k_3} \quad (2) \hat{\theta} = \frac{2k_1 + k_2}{k_1 + k_2 + k_3}$$

$$(3) \hat{\theta} = \frac{k_1 + k_2}{k_1 + k_2 + k_3} \quad (4) \hat{\theta} = \frac{k_1}{k_1 + k_2 + k_3}$$

۴۴- در مورد روش‌های خوشه‌بندی کدام مورد درست است؟

- (۱) تابع هزینه روش k -means ممکن است در برخی از گام‌های الگوریتم افزایش یابد.
- (۲) الگوریتم EM برای پیدا کردن پارامترهای GMM، پارامترهایی را پیدا می‌کند که متناظر با بیشینه سراسری تابع (Likelihood) هستند.
- (۳) در روش EM ممکن است حالتی از خوشه‌بندی پیدا شود که مقدار درست‌نمایی (Likelihood) را بی‌نهایت کند، ولی خوشه‌بندی مطلوبی نباشد.
- (۴) اگر خوشه‌ها توزیع گاوسی متقارن (با ماتریس کوواریانس αI) باشند، جواب k -means و EM روی GMM لزوماً یکی می‌شود.

۴۵- کدام یک از موارد زیر درست است؟

- (۱) در تشخیص داده‌های پرت، روش‌های پارامتری نسبت به روش‌های ناپارامتری حساسیت کمتری در برابر نویز دارند.
- (۲) تخمین‌زن k -NN توانایی محاسبه تابع چگالی احتمال را دارد.
- (۳) در k -NN با افزایش k ، واریانس مدل افزایش می‌یابد.
- (۴) در k -NN با کاهش k ، بعد VC افزایش می‌یابد.

