

**PART A: Grammar**

Directions: Choose the word or phrase (1), (2), (3), or (4) that best completes the blank. Then, mark the correct choice on your answer sheet.

1- The rate that bright comets enter the solar system implies there should be around 3000 of them buzzing around, only 25 are known.

- 1) nonetheless 2) regardless of the fact 3) and yet 4) as there are

2- Contemporary theories of interpretation require that, in our analyses of texts, we consider not only what the text says“made.”

- 1) also its meaning gets and 2) but also gets the meaning of it
3) but its meaning also gets 4) but how its meaning gets

3- individual behavior is influenced by social networks is beyond dispute.

- 1) That 2) An 3) The 4) It is that

4- Plant scientists have been trying for years to genetically modify flowers for aesthetic purposes. The first to go on sale were blue carnations in Australia, in 1996.

- 1) were produced 2) produced 3) had been produced 4) to produce

5- Weapons have been carried and delivered by a wide variety of vehicles, weapon platforms.

- 1) they are often called 2) often called 3) called they are often 4) that are called often

6- Articulating what the difference between humans and other creatures consists of behind it have formed a large and difficult project tackled by biologists, anthropologists, psychologists, and philosophers.

- 1) uncovering the biology 2) the biology of uncovering
3) the biology uncovering 4) and uncovering the biology

7- Most healthcare professionals view depression as “just part of getting old and argue that this illness,, can have serious, even fatal consequences.

- 1) untreated then 2) untreated whether it is 3) if untreated 4) that is untreated

8- Ted had a terrible habit of boasting so much about his smallest accomplishments his vainglory became renowned throughout the small college campus.

- 1) that 2) as 3) in that 4) as though

**PART B: Vocabulary**

Directions: Choose the word or phrase (1), (2), (3), or (4) that best completes the blank. Then, mark the correct choice on your answer sheet.

🐞 9- Dogs growl and show their teeth in an attempt to frighten the animal or person they perceive as a

- 1) habitat 2) prey 3) suspicion 4) threat

🐞 10- Based on his recent poor decisions, it was obvious that Seth lacked even a modicum of good

- 1) sentiment 2) sense 3) sensation 4) sensitivity

🐞 11- The judge the extraneous evidence because it was not pertinent to the trial.

- 1) disclosed 2) distended 3) dismissed 4) distorted

🐞 12- The more frequently employees take time to exercise during working hours each week, the fewer sick days they

- 1) expend 2) save 3) take 4) recall

🐞 13- Classic psychology experiments have shown that when rats are first with an electrical shock to fear a tone when it sounds, they later fear the tone even without the associated shock.

- 1) conditioned 2) sparkled 3) displayed 4) intended

🐞 14- In 1998 Gordon Sinclair, the owner of a well-known restaurant, was struggling with a problem that all restaurateurs. Patrons frequently reserve a table but, without notice, fail to appear.

- 1) delegates 2) afflicts 3) intensifies 4) evades

🐞 15- Despite what the scientist said, the volcano eruption is not , so do not be concerned!

- 1) impassive 2) negotiable 3) vulnerable 4) imminent

🐞 16- At the landfill, the process is in full swing, turning much of the garbage into gasses.

- 1) conversion 2) restoration 3) decomposition 4) pressurization

🐞 17- Because I am an extreme planner who needs to control everything, I never engage in

- 1) justification 2) pretention 3) coincidence 4) spontaneity

🐞 18- The roads in our town already have too much traffic; building a new shopping mall will the problem.

- 1) frustrate 2) exacerbate 3) preserve 4) exploit

🐞 19- The movie *Close Encounters of the Third Kind* tells the story of the first contact between beings from outer space and creatures, that is, those living on earth.

- 1) terrestrial 2) dominant 3) ingenious 4) affable

🐞 20- There is agreement that an airport is needed; no one disputes that, but there is fundamental disagreement about where to build it.

- 1) uniform 2) utilitarian 3) unique 4) unanimous

بخش اول: دستور زبان

در سؤالات زیر، از بین گزینه‌های (۱)، (۲)، (۳) و (۴) پاسخی را انتخاب کنید که به بهترین نحو جای خالی را پر کند. آنگاه پاسخ‌تان را روی پاسخنامه علامت بزنید.

۱- گزینه «۳» با توجه به سرعت و تعداد ورود ستاره‌های دنباله‌دار به منظومه شمسی می‌توان حدس زد که باید تقریباً ۳۰۰۰ مورد از آن‌ها وجود داشته؛ با این حال تنها ۲۵ عدد از آنها شناسایی شده‌اند.

توضیح: همان‌طور که می‌دانید nonetheless قید ربط است؛ یعنی قبل از آن باید نقطه یا نقطه‌ویرگول و بعد از آن باید حتماً کاما بیاید. با این حساب گزینه (۱) نادرست است. گزینه (۲) در صورتی ارزش بررسی کردن دارد که طراح بعد از fact از حرف ربط that استفاده می‌کرد. گزینه (۳) صحیح است، هم با توجه به مفهوم جمله و هم با توجه به اینکه قبل از and کاما می‌آید. و گزینه (۴) نادرست است چون بعد از as دو تا فعل داریم؛ یکی are و یکی know.

۲- گزینه «۴» نظریه‌پردازان معاصر در زمینه ترجمه شفاهی باور دارند ما در آنالیز متن، علاوه بر چیزی که متن می‌گوید، باید به نحوه شکل‌گیری معنی آن نیز توجه داشته باشیم.

توضیح: همان‌طور که می‌بینیم این تست با مبحث not only ... but also سروکار دارد. اول از همه اینکه در این ساختار but also می‌تواند به صورت but یا also هم به کار برود. پس امیدوارم فوری گزینه (۲) را نزده باشید.

ضمناً می‌دانیم ساختار (also) but ... not only مستلزم رعایت ساختار موازی است؛ بنابراین چون بعد از not only کلمه پرسشی what را داریم باید بعد از but(also) هم از کلمه‌ی پرسشی how استفاده کنیم:

.... not only what the text says but how its meaning gets made.

۳- گزینه «۱» اینکه شبکه‌های اجتماعی بر روی رفتار افراد تاثیرگذار هستند، قابل تردید نمی‌باشد.

توضیح: تست بسیار ساده‌ای است. توی مبحث جمله‌واره‌ی اسمی گفتیم یکی از کاربردهای that clause این است که قبل از فعل be به عنوان فاعل استفاده شوند. گفتیم در این موارد that به صورت «اینکه» ترجمه می‌شود:

That individual behavior is influenced by social networks is beyond dispute.

مثال بیشتر:

That coffee grows in Brazil is well known.

۴- گزینه «۲» گیاه‌شناسان سال‌هاست که با استفاده از اصلاح ژنتیک به دنبال زیباتر ساختن گل‌ها هستند. گل میخک آبی اولین موردی بود که برای فروش عرضه شد. این گل در سال ۱۹۹۶ در استرالیا تولید شد.

توضیح: این تست از دو جمله تشکیل شده که برای پاسخگویی به آن فقط به جمله دوم نیاز داریم. جمله دوم دارای فعل اصلی were می‌باشد، با این حساب به هیچ فعل اصلی دوم دیگری نیاز نداریم چون هر جمله فقط و فقط باید یک فعل اصلی داشته باشد. این یعنی حذف همزمان گزینه‌های (۱) و (۳). گزینه (۴) نادرست است چون قصد بیان هدف نداریم.

ضمناً شکل اولیه گزینه (۲) این‌طوری بوده:

The first to go on sale were blue carnations that were produced in Australia, in 1996.

اگر that were را حذف کنیم، به گزینه (۲) می‌رسیم.



۵- گزینه «۲» سلاح‌ها از طریق وسایل نقلیه مختلفی حمل و تحویل داده می‌شوند. این وسایل نقلیه اغلب با نام پلتفرم سلاح شناخته می‌شوند.
توضیح: تقریباً هر سال از این بحث سؤال می‌آید و ما هم هر سال می‌گوییم بعد از کاما کاربرد that ممنوع است. (این یعنی حذف گزینه (۴)).
گزینه (۱) در صورتی صحیح است که کاما به نقطه تبدیل بشود و they هم به They. مهم‌ترین دلیل رد گزینه (۳) کاربرد they بعد از called است.
ضمناً شکل اولیه‌ی گزینه‌ی ۲ این طوری بوده:

Weapons have been carried and delivered by a wide variety of vehicles, **which are often called** weapon platforms.

اگر which are حذف کنیم، به گزینه (۲) می‌رسیم.

۶- گزینه «۴» درک تفاوت بین انسان و سایر موجودات و مسائل بیولوژیکی نهفته در آن باعث بوجود آمدن مباحث و تحقیقات دشوار و گسترده‌ای شده است که دانشمندانی از رشته‌های مختلف مانند زیست‌شناسی، انسان‌شناسی، روانشناسی و فلسفه به آن می‌پردازند.
توضیح: توی تست‌هایی که این‌قدر طولانی هستند، اولین کار این است که به دنبال فعل اصلی باشیم. فعل اصلی سوال ما have formed است. پس به خاطر حضور have باید فاعلمون جمع باشد. اما articulating به تنهایی به فعل مفرد نیاز دارد، این یعنی باید articulating را با and به یک ساختار ing دار موازی دیگر متصل کنیم تا آن موقع کاربرد فعل have هم معنی پیدا کند. و چون فقط گزینه (۴) است که دارای and می‌باشد، می‌توانیم باقی گزینه‌ها را رد کنیم.

۷- گزینه «۳» اکثر متخصصین حوزه بهداشت و درمان، افسردگی را بخشی از پروسه افزایش سن می‌دانند و اعتقاد دارند که در صورت عدم درمان می‌تواند عواقب بسیار وخیمی داشته و یا حتی باعث مرگ بیمار شود.
توضیح: اول از همه اینکه طراح سؤال ظاهراً یادش رفته آن (") را که باز کرده ببندد. باید این علامت را قبل از and بیاورد. حالا می‌رسیم به رد گزینه‌ها. کاربرد that بعد از کاما ممنوع است (یعنی رد گزینه (۴)). گزینه ۱ نادرست است چون معلوم نیست طراح سوال آن then را بابت چی استفاده کرده. گزینه (۲) هم کنار می‌رود چون بعد از is هیچ عبارت کامل‌کننده‌ای نداریم. اما برای اینکه ببینیم چرا گزینه (۳) صحیح است باید اصل جمله را پیدا کنیم.

Most healthcare professionals argue that this illness, **if it is untreated**, can have serious, even fatal consequences.

چون it به this illness برمی‌گردد، می‌توانیم با فرض اینکه فاعل‌ها یکسان هستند، فاعل جمله‌واره‌ی وابسته یعنی it و فعل is را حذف کنیم و یک وجه وصفی بسازیم:

Most healthcare professionals argue that this illness, **if untreated**, can have serious, even fatal consequences.

۸- گزینه «۱» تد اخلاق بسیار زشتی داشت و به خاطر کوچک‌ترین موفقیت‌هایش به قدری فخر فروشی می‌کرد که عادت خودستایی او در سرتاسر محوطه‌ی کوچک دانشگاه زبانزد عام و خاص بود.
توضیح: از ساختار that so استفاده شده.

....so much about that

بخش دوم: واژگان

دستورالعمل: در سؤالات زیر، از بین گزینه‌های (۱)، (۲)، (۳) و (۴) پاسخی را انتخاب کنید که به بهترین نحو جای خالی را پر کند. آنگاه پاسخ‌تان را روی پاسخنامه علامت بزنید.

۹- گزینه «۴» سگ‌ها در هنگام مواجهه با خطر / تهدید، پارس می‌کنند و دندان‌های خود را نشان می‌دهند تا حیوان یا شخص مورد نظر را بترسانند.

- (۱) زیستگاه، زیست‌بوم (۲) طعمه (۳) سوءظن، تردید (۴) خطر، تهدید

۱۰- گزینه «۲» ضعف تصمیمات اخیرِ سِت نشان می‌دهد که کوچکترین درکی نسبت به مسائل مختلف ندارد.

- (۱) تمایل، گرایش، احساس (۲) شعور، معنی، ادراک (۳) احساس، هیجان (۴) حساسیت

۱۱- گزینه «۳» قاضی شواهد غیرضروری را مردود اعلام کرد زیرا ارتباط چندانی با روال دادرسی نداشت.

- (۱) افشاء کردن، فاش کردن (۲) بزرگ کردن، منبسط کردن (۳) مردود شمردن، رد کردن (۴) کج کردن، تحریف کردن

۱۲- گزینه «۳» کارمندان هرچقدر در طول هفته بیشتر ورزش کنند، کمتر به مرخصی استعلاجی نیاز پیدا می‌کنند.

توضیح: جواب این سؤال عیناً توی خود سؤال آمده. یعنی اول بوده take time حالا شده take days.

اصطلاح take sick days یعنی «استعلاجی گرفتن».

۱۳- گزینه «۱» طبق آزمایشات روانشناسی کلاسیک، وقتی موش‌ها برای ترسیدن از یک صدای بخصوص به وسیله شوک الکتریکی شرطی شوند،

بعدها بدون وجود شوک الکتریکی هم از آن صدا می‌ترسند.

- (۱) شرطی کردن (۲) درخشیدن، برق زدن (۳) نمایش دادن (۴) قصد داشتن

۱۴- گزینه «۲» گوردن سینکлер، صاحب یکی از رستوران‌های مشهور، در سال ۱۹۹۸ با یکی از مشکلاتی که تمامی صاحب رستوران‌ها را آزار می‌داد،

دست و پنجه نرم می‌کرد. مشتریان، میز رزرو می‌کردند اما بدون اطلاع در رستوران حاضر نمی‌شدند.

- (۱) تفویض کردن، سپردن (۲) رنجاندن، آزردن (۳) تشدید کردن (۴) طفره رفتن، شانه خالی کردن

۱۵- گزینه «۴» علیرغم دیدگاه آن دانشمند، فوران آتشفشان در آینده‌ای نزدیک اتفاق نخواهد افتاد. بنابراین جایی برای نگرانی نیست!

- (۱) خونسرد، بی‌عاطفه (۲) قابل مذاکره (۳) آسیب‌پذیر (۴) قریب‌الوقوع

۱۶- گزینه «۳» در مکان دفن زباله، پروسه‌ی تجزیه با قدرت در حال اجرا است. در این پروسه بخش زیادی از زباله به گاز تبدیل می‌شود.

- (۱) تبدیل، تغییر (۲) ترمیم، بازسازی (۳) تجزیه (۴) تحت فشار قرار دادن

۱۷- گزینه «۴» من، به دلیل اینکه برنامه‌ریز بسیار دقیقی هستم، باید همه چیز را تحت نظر داشته باشم. بنابراین در زندگی من هرگز چیزی

بدون برنامه‌ریزی (فی‌البداهه) اتفاق نمی‌افتد.

- (۱) توجیه، دلیل آوردن (۲) تظاهر، وانمود (۳) اتفاق، تصادف، تقارن (۴) بدون برنامه، خودبخودی

سؤالات مهندسی کامپیوتر - هوش مصنوعی

مجموعه دروس تخصصی (ساختمان داده‌ها و طراحی الگوریتم‌ها، شناسایی الگو، یادگیری ماشین)

کله ۱- الگوریتم فلوید - وارشال از یک الگوریتم استفاده می‌کند. برای حل مسئله کوتاه‌ترین مسیرهای تمام جفت رؤس در یک گراف جهت‌دار $G = (V, E)$ در زمان استفاده می‌کند.

(۱) حریصانه، $\Theta(v^3)$ (۲) حریصانه، $\Theta(V^2 \log E)$ (۳) برنامه‌نویسی پویا، $\Theta(v^3)$ (۴) برنامه‌نویسی پویا، $\Theta(V^2 \log E)$

کله ۲- با فرض اینکه $P \neq NP$ باشد، کدام مورد درست است؟

(۱) $NP - hard = NP$ (۲) $NP - complete = P$ (۳) $NP - complete = NP$ (۴) $NP - complete \cap P = \emptyset$

کله ۳- کمترین تعداد مقایسه موردنیاز برای تعیین اینکه یک عدد صحیح بیش از $\frac{n}{4}$ مرتبه در یک آرایه مرتب از اعداد صحیح به طول n ظاهر می‌شود، از کدام مرتبه است؟

(۱) $\Theta(1)$ (۲) $\Theta(\log n)$ (۳) $\Theta(n)$ (۴) $\Theta(n \log n)$

کله ۴- n آرایه نامرتب $A_1, \dots, A_n$ را در نظر بگیرید (n عددی فرد است). هر کدام از این آرایه‌ها دارای n عنصر متمایز است. هیچ عنصر مشترکی میان هیچ دو آرایه‌ای وجود ندارد. کمترین پیچیدگی زمانی الگوریتمی برای محاسبه میانه این آرایه‌ها از چه مرتبه‌ای است؟

(۱) $\Theta(n)$ (۲) $\Theta(n \log n)$ (۳) $\Theta(n^2)$ (۴) $\Omega(n^2 \log n)$

کله ۵- فرض کنید $W(n)$ و $A(n)$ ، به ترتیب، نشان‌دهنده بدترین حالت و میانگین زمان اجرای الگوریتم اجراشده بر روی ورودی با اندازه n باشند. کدام مورد همواره درست است؟

(۱) $A(n) = O(W(n))$ (۲) $A(n) = \Theta(W(n))$ (۳) $A(n) = \Omega(W(n))$ (۴) $A(n) = o(W(n))$

کله ۶- یک آرایه مرتب‌شده از اعداد داریم. می‌خواهیم دو عدد در این آرایه پیدا کنیم که جمع آن دو عدد مساوی یک عدد داده‌شده x باشد. کمترین پیچیدگی زمانی حل این مسئله کدام است؟

(۱) $\Theta(n)$ (۲) $\Theta(n^2)$ (۳) $\Theta(\log n)$ (۴) $\Theta(n \log n)$

کله ۷- فرض کنید آرایه‌ای از اعداد صحیح $A = [a_1; a_2; \dots; a_n]$ داده شود. فرض کنید یک اندیس (k ناشناخته) وجود دارد به‌طوری‌که زیر آرایه $A = [a_1; a_2; \dots; a_k]$ به ترتیب اکیداً افزایشی مرتب شده است و زیر آرایه $A = [a_k; a_{k+1}; \dots; a_n]$ به ترتیب اکیداً نزولی مرتب شده است (یعنی اگر $1 \leq i < j \leq k$ ، آنگاه $a_i < a_j$ ، و اگر $k \leq i < j \leq n$ ، آنگاه $a_i > a_j$ ، هدف شما تعیین k است. یک الگوریتم بهینه برای حل این مسئله چه زمان اجرایی دارد؟

(۱) $\Theta(n^2 \log n)$ (۲) $\Theta(n \log n)$ (۳) $\Theta(n)$ (۴) $\Theta(\log n)$

کله ۸- کدام یک از موارد زیر درست است؟

(۱) آرایه $A = [10; 3; 5; 1; 4; 2]$ یک max heap است.

(۲) هر مسئله محاسباتی با اندازه ورودی n را می‌توان با یک الگوریتمی با زمان چندجمله‌ای برحسب n حل کرد.

(۳) برای تمام توابع مثبت $f(n)$ ، $g(n)$ و $h(n)$ ، اگر $f(n) = O(g(n))$ و $f(n) = \Omega(h(n))$ باشد، آنگاه $g(n) + h(n) = \Omega(f(n))$ است.

(۴) اگر هر رقم جداگانه در RADIX SORT را با استفاده از INSERTION SORT به‌جای COUNTING SORT مرتب کنیم، آنگاه RADIX SORT به‌درستی کار نمی‌کند (یعنی خروجی صحیح را تولید نمی‌کند).

کله ۹- کدام مورد زیر مطمئناً عبارت $f(n) = \Omega(g(n))$ را پشتیبانی می‌کند؟

(۱) $f(n) \leq 4 \times g(n)$ برای تمام $n \geq 1$ (۲) $f(n) \geq 4 \times g(n)$ برای تمام $n \geq 136$

(۳) $f(n) \leq 4 \times g(n)$ برای تمام $n \geq 100$ (۴) $\lim_{n \rightarrow \infty} \frac{f(n)}{g(n)} = 0$



۱۰- کدام یک از گزاره‌های زیر درست است؟

الف) اگر مسئله P_1 بتواند به مسئله P_2 در زمان خطی کاهش (reduce) یابد، آنگاه اگر P_2 یک مسئله NP-hard باشد، می‌توان نتیجه گرفت P_1 نیز NP-hard است.

ب) یک Clique در یک گراف بدون جهت لزوماً یک vertex cover در گراف مکمل نیست.

(۱) فقط گزاره «الف» درست است.

(۳) هر دو گزاره «الف» و «ب» درست است.

(۴) هر دو گزاره «الف» و «ب» نادرست است.

۱۱- کدام یک از گزاره‌های زیر درست است؟

الف) اینکه زمان حل یک مسئله P حد پایین $\Omega(n^2)$ دارد به این معنی است که برای هر الگوریتم A که P را حل می‌کند فقط برخی از نمونه‌های P وقتی به‌عنوان ورودی به A داده شوند، باعث می‌شود A زمان $\Omega(n^2)$ صرف کند.

ب) اینکه زمان حل یک مسئله P حد پایین $\Omega(n^2)$ دارد به این معنی است که برای هر الگوریتم A که P را حل می‌کند هر نمونه از P که به‌عنوان ورودی به A داده شود، باعث می‌شود A زمان $\Omega(n^2)$ صرف کند.

(۱) فقط گزاره «الف» درست است.

(۳) هر دو گزاره «الف» و «ب» درست هستند.

(۴) هر دو گزاره «الف» و «ب» نادرست هستند.

۱۲- کدام یک از گزاره‌های زیر درست است؟

الف) اگر یک الگوریتم زمان چندجمله‌ای برای یک مسئله که NP-hard است ارائه شود، آنگاه می‌توان نتیجه گرفت $P = NP$ است.

ب) اگر یک مسئله NP-complete است، آنگاه می‌توان نتیجه گرفت که آن مسئله هیچ راه‌حلی ندارد.

(۱) فقط گزاره «الف» درست است.

(۳) هر دو گزاره «الف» و «ب» درست هستند.

(۴) هر دو گزاره «الف» و «ب» نادرست هستند.

۱۳- کدام یک از گزاره‌های زیر درست است؟

الف) اگر یک مسئله که به کلاس NP تعلق دارد یک راه‌حل زمان چندجمله‌ای داشته باشد، آنگاه $P = NP$ است.

ب) اگر کسی یک حد پایین زمان نمایی برای یک مسئله که NP-complete است بدهد، آنگاه $P \neq NP$ است.

(۱) فقط گزاره «الف» درست است.

(۳) هر دو گزاره «الف» و «ب» درست هستند.

(۴) هر دو گزاره «الف» و «ب» نادرست هستند.

۱۴- مسئله کوله‌پشتی ۱-۰ را در نظر بگیرید که n شیء با وزن صحیح داریم و گنجایش کوله‌پشتی عدد صحیح M است. این مسئله دارای یک الگوریتم مبتنی بر روش برنامه‌ریزی پویا با زمان $O(M, n)$ است. این مرتبه زمانی برحسب اندازه مسئله چگونه است؟

(۱) خطی (۲) درجه دو (۳) نمایی (۴) شبه چندجمله‌ای

۱۵- برای پیدا کردن k امین عدد در میان n عدد که به‌عنوان کلید در گره‌های یک درخت جستجوی دودویی متوازن ذخیره شده‌اند، کمترین پیچیدگی زمانی ممکن کدام است؟ (هر گره درخت فقط شامل کلید و اشاره‌گر به پدر و فرزند چپ و راست است.)

(۱) $\Theta(n)$ (۲) $\Theta(\log n)$ (۳) $\Theta(n \log n)$ (۴) $\Theta(\log \log n)$

۱۶- n عدد به‌عنوان کلید در گره‌های یک درخت جستجوی دودویی متوازن ذخیره شده‌اند. هر گره علاوه بر کلید و اشاره‌گر به پدر و فرزند چپ و راست، تعداد گره‌های زیر درخت خود را هم نگهداری می‌کند. برای پیدا کردن rank کلید یک گره (یعنی اینکه کلید گره چندمین عدد در بین n عدد است) کمترین پیچیدگی زمانی ممکن کدام است؟

(۱) $\Theta(\log \log n)$ (۲) $\Theta(\log n)$ (۳) $\Theta(n)$ (۴) $\Theta(n \log n)$

۱۷- در یک درخت قرمز-سیاه، طول طولانی‌ترین مسیر ساده از یک گره x به یک برگ در زیر درخت خودش حداکثر چند برابر طول کوتاه‌ترین مسیر از گره x به یک برگ در زیر درخت خودش است؟

(۱) ۴ (۲) ۳ (۳) ۲ (۴) ۱

۱۸- کدام یک از گزاره‌های زیر درست است؟

الف) هر درخت جستجوی دودویی دلخواه با n گره می‌تواند به یک درخت جستجوی دودویی دلخواه دیگر با n گره با انجام عمل rotation تبدیل شود.

ب) برای هر دو تابع $f(n)$ و $g(n)$ یکی از سه حالت $f(n) \in o(g(n))$ ، $f(n) \in \Theta(g(n))$ و $f(n) \in \omega(g(n))$ برقرار است.

(۱) فقط گزاره «الف» درست است.

(۳) هر دو گزاره «الف» و «ب» درست هستند.

(۴) هر دو گزاره «الف» و «ب» نادرست هستند.

کله ۱۹- کمترین پیچیدگی زمانی ممکن برای مرتب‌سازی n عدد طبیعی که مقادیر کمتر از n^2 دارند، کدام است؟

$$\Theta(n \log^2 n) \quad (۱) \quad \Theta(n^2) \quad (۲) \quad \Theta(n \log n) \quad (۳) \quad \Theta(n) \quad (۴)$$

کله ۲۰- اگر عدد ۳۶۳ را در یک درخت جستجوی دودویی، جستجو کنیم، کدام دنباله زیر نمی‌تواند دنباله‌ای از کلید گره‌هایی باشد که بررسی می‌شوند؟ (ترتیب از راست به چپ است.)

$$\begin{aligned} (۱) \quad & ۳۶۳, ۳۹۷, ۳۴۴, ۳۳۰, ۳۹۸, ۴۰۱, ۲۵۲, ۲ \\ (۲) \quad & ۳۶۳, ۳۶۲, ۲۵۸, ۸۹۸, ۲۴۴, ۹۱۱, ۲۲۰, ۹۲۴ \\ (۳) \quad & ۳۶۳, ۲۷۸, ۳۸۱, ۳۸۲, ۲۶۶, ۲۱۹, ۳۸۷, ۳۹۹, ۲ \\ (۴) \quad & ۳۶۳, ۲۴۵, ۹۱۲, ۲۴۰, ۹۱۱, ۲۰۲, ۹۲۵ \end{aligned}$$

کله ۲۱- نزدیک‌ترین فاصله نقطه z از ابرصفحه $W^T x = 0$ کدام است؟

$$W^T z \quad (۱) \quad \frac{W^T z}{\|W\|} \quad (۲) \quad \|W\| \cdot \|z\| \quad (۳) \quad \frac{W^T z}{\|W\|^2} \quad (۴)$$

کله ۲۲- در دسته‌بندی یک مسئله M دسته‌ای به کمک روش بیزین، کدام رابطه در مورد احتمال خطای دسته‌بندی دسته‌بند بهینه یا همان P_e صادق است؟

$$\frac{1}{M} \leq P_e \leq 1 - \frac{1}{M} \quad (۱) \quad P_e \leq 1 - \frac{M}{1+M} \quad (۲) \quad P_e \leq \frac{1}{M} \quad (۴) \quad P_e \leq 1 - \frac{1}{M} \quad (۳)$$

کله ۲۳- در ساخت یک جنگل تصادفی هر درخت را براساس نمونه‌گیری Bootstrap از داده‌های آموزشی می‌سازیم. اگر n تعداد کل نمونه‌های آموزشی و d تعداد ویژگی باشد، به‌طور متوسط هر درخت براساس چه تعداد نمونه آموزشی ساخته می‌شود؟

$$(1 - \frac{1}{n})^n \quad (۱) \quad (1 - \frac{1}{n})^d \quad (۲) \quad (1 - \frac{1}{n})^d \quad (۳) \quad 1 - (1 - \frac{1}{n})^n \quad (۴)$$

کله ۲۴- فرض کنید در یک مسئله دسته‌بندی، داده‌های دارای دو ویژگی x_1 و x_2 هستند که هر کدام مقادیر -1 ، 0 یا 1 را می‌گیرند و برچسب طبقه هم مقدار صفر یا یک را می‌گیرد. سه فرضیه h_1, h_2, h_3 به‌صورت زیر تعریف شده‌اند:

$$h_1(x) = \begin{cases} 1 & \text{if } x_1 = 1 \\ 0 & \text{else} \end{cases} \quad h_2(x) = \begin{cases} 1 & \text{if } x_1 \neq x_2 \\ 0 & \text{else} \end{cases} \quad h_3(x) = \begin{cases} 1 & \text{if } x_1, x_2 = 0 \\ 0 & \text{else} \end{cases}$$

کدام یک از فرضیه‌های فوق را می‌توان از طریق یک پرسپترون (با گرفتن وزن‌های مناسب) پیاده‌سازی کرد؟

$$h_1 \quad (۱) \quad h_2 \quad (۲) \quad h_3 \quad (۳) \quad h_3 \text{ و } h_2 \quad (۴)$$

کله ۲۵- در دسته‌بندی یک مسئله دو دسته‌ای با داده‌های آموزش به شرح جدول زیر، دو دسته‌بند Naïve Bayes و Bayes داده آزمون $(Weight = ۸۰, Height = ۱۸۰)$ را در کدام دسته، دسته‌بندی می‌کنند؟

$$Naïve Bayes \rightarrow -, Bayes \rightarrow + \quad (۱)$$

$$Naïve Bayes \rightarrow +, Bayes \rightarrow - \quad (۲)$$

$$\text{هر دو مثبت} \quad (۳)$$

$$\text{هر دو منفی} \quad (۴)$$

#	Height	Weight	Class
۱	۸۰	۱۸۰	+
۲	۷۰	۱۸۰	+
۳	۸۰	۱۸۰	-
۴	۸۰	۱۸۰	+
۵	۷۰	۱۷۰	+
۶	۵۰	۱۷۰	-
۷	۴۰	۱۶۰	-
۸	۳۰	۱۴۰	-
۹	۱۰۰	۱۷۰	+
۱۰	۱۱۰	۱۶۰	+
۱۱	۵۰	۱۶۰	-
۱۲	۱۱۰	۱۶۰	+



پاسخنامه مهندسی کامپیوتر – هوش مصنوعی

مجموعه دروس تخصصی (ساختمان داده‌ها و طراحی الگوریتم‌ها، شناسایی الگو، یادگیری ماشین)

۱- گزینه «۳» الگوریتم فلوید - وارشال یک روش حل برنامه‌نویسی پویاست (از حافظه استفاده می‌کند) که هزینه محاسباتی آن $\Theta(v^3)$ است.

۲- گزینه «۴» اگر P برابر با NP نباشد، P و NP -Complete اشتراکی ندارند.

۳- گزینه «۱» این کار را می‌توان با یک جدول درهم‌سازی انجام داد. توجه داشته باشید که تعداد مقایسه مدنظر است نه هزینه محاسباتی. تعداد مقایسه‌های مورد نیاز $\Theta(1)$ است.

۴- گزینه «۳» یافتن میانه میانه‌ها برای n^2 عنصر هزینه $\Theta(n^2)$ دارد.

۵- گزینه «۱» اگر بخواهیم بگوییم بهترین حالت زمان اجرای یک الگوریتم، در بدترین حالت، بیشتر از بدترین حالت زمان اجرای آن نخواهد بود می‌توانیم بگوییم:

$$A(n) = O(W(n))$$

۶- گزینه «۱» این دو آرایه را A و B در نظر بگیرید که هر دو صعودی هستند. اعضای B را از b به $x-b$ تغییر می‌دهیم و آن را C می‌نامیم. برای حفظ صعودی بودنش، از آخر به اول می‌چینیم. حال روالی مانند merge از مرتب‌سازی ادغامی را روی آن می‌یابیم تا یک عنصر تکراری در A و C بیابیم. در صورت یافتن این دو عنصر با فرض اینکه مقدار این عنصر برابر a است، این دو عنصر برابر a از مجموعه A و $x-a$ از مجموعه B هستند.

۷- گزینه «۴» به روشی مشابه با جست‌وجوی دودویی می‌توان عنصری که از دو عنصر مجاورش بزرگ‌تر باشد را بیابیم. هزینه این کار $\Theta(\log n)$ است.

۸- گزینه «۳» عبارت گزینه (۱) نادرست است، زیرا عنصر ۴ نمی‌تواند فرزند عنصر ۳ در هرم بیشینه باشد. عبارت گزینه (۲) نادرست است و هر مسئله‌ای راه‌حل چندجمله‌ای ندارد. عبارت گزینه (۳) صحیح است، زیرا اگر $f(n) = O(g(n))$ ، آنگاه $g(n) = \Omega(f(n))$ خواهد بود. یعنی رشد g از f کمتر نیست. اضافه شدن یک مقدار مثبت به g نیز این شرط را رد نمی‌کند. عبارت گزینه (۴) نادرست است. استفاده از مرتب‌سازی درجی منجر به افزایش هزینه مرتب‌سازی می‌شود و خروجی نادرست نمی‌دهد.

۹- گزینه «۲» اینکه رشد f ناکمتر از g است تنها در عبارت گزینه (۲) قابل مشاهده است.

۱۰- گزینه «۲» عبارت اول نادرست است و مسئله A می‌تواند NP -Hard نباشد. عبارت دوم صحیح است و یک vertex cover الزاماً یک clique نیست. clique یک زیرگراف کامل است ولی vertex cover مجموعه رئوسی است که هر یال از گراف حداقل به یکی از آن‌ها متصل است.

۱۱- گزینه «۲» تعریف «ب» تعریف درستی برای نماد Ω است.

۱۲- گزینه «۱» عبارت «الف» درست است. عبارت «ب» نادرست است. منظور از NP -Complete غیرقابل حل بودن مسئله نیست.

۱۳- گزینه «۲» عبارت «الف» نادرست است. مسئله NP راه حل چندجمله‌ای ندارد.

۱۴- گزینه «۴» مرتبه $O(M.n)$ شبه چندجمله‌ای است.

۱۵- گزینه «۱» نیاز به پیمایش درخت داریم که هزینه آن $\Theta(n)$ است.

۱۶- گزینه «۲» با جست‌وجوی دودویی در درخت متوازن با هزینه $\Theta(\log n)$ به عنصر مطلوب می‌رسیم. rank این گره با توجه به تعداد گره‌های زیردرخت اجداد و فرزند چپ گره موردنظر تعیین می‌شود. اجداد گره موردنظر و گره‌های فرزند چپ آن که گره موردنظر در زیردرخت راست آن قرار دارد، قبل از گره موردنظر هستند. برای مثال اگر گره هدف عدد ۵۰ باشد و یکی از اجداد آن که گره ۵۰ در زیردرخت راست آن قرار دارد ۳۰ باشد، تمام فرزندان زیردرخت چپ ۳۰ و خود گره ۳۰ قبل از ۵۰ خواهند بود. همچنین فرزندان زیردرخت چپ ۵۰ نیز قبل از ۵۰ خواهند بود. تعداد گره‌های این مسیر $\Theta(\log n)$ است.

۱۷- گزینه «۳» درخت قرمز - سیاه یک درخت متوازن است. نسبت بزرگ‌ترین و کوچک‌ترین مسیر از یک گره تا برگ‌ها حداکثر برابر ۲ است.

۱۸- گزینه «۱» عبارت «الف» درست است. ولی عبارت «ب» با توجه به اینکه توابع می‌توانند نوسانی باشند نادرست است.

۱۹- گزینه «۴» می‌توان از مرتب‌سازی radix sort استفاده کرد. هزینه این کار معادل با دو مرتبه مرتب‌سازی شمارشی برای اعداد ۰ تا n است که یکبار باقی‌مانده و یکبار خارج‌قسمت اعداد به n را برای مرتب‌سازی در نظر می‌گیریم. هزینه محاسبات $\Theta(n)$ خواهد بود.

۲۰- گزینه «۴» عدد ۹۱۲ نمی‌تواند بعد از ۹۱۱ و ۲۴۰ بیاید. به همین دلیل دنباله عبارت گزینه (۴) مورد قبول نیست.

۲۱- گزینه «۲» در معادله ابر صفحه جداکننده، W برابر بردار وزن و X نیز برابر بردار ویژگی می‌باشد. ضرب نقطه‌ای $W^t \cdot Z$ بیانگر میزان شباهت نقطه Z با ابر صفحه می‌باشد. نتیجه حاصل از این ضرب اگر مقداری مثبت باشد، این نقطه در یک طرف ابر صفحه بوده و اگر مقدار منفی، نقطه در طرف دیگر ابر صفحه می‌باشد. حال اگر مقدار به‌دست آمده تقسیم بر بزرگی بردار وزن $(\|W\|)$ گردد، نتیجه به‌دست آمده بیانگر نزدیک‌ترین فاصله نقطه Z نسبت به ابر صفحه جداکننده می‌باشد.

در گزینه (۱) فاصله به‌دست آمده بدون در نظر گرفتن بزرگی بردار می‌باشد.

در گزینه (۳) تنها بزرگی دو بردار در یکدیگر ضرب شده است که ارتباطی با فاصله دو بردار ندارد.

در گزینه (۴) نیز فرمول به صورت نادرست نوشته شده است.

۲۲- گزینه «۳» یک مسئله کلاس‌بندی با دو کلاس را در نظر بگیرید. در صورتی که مدل به صورت شانس و تصادفی هر داده ورودی را به یک کلاس تعلق دهد بیانگر عدم یادگیری مدل و بالاترین میزان خطا می‌باشد. در این حالت احتمال انتخاب هر کلاس $\frac{1}{2}$ می‌باشد. برعکس این قضیه، در صورتی که مدل کاملاً آموزش دیده باشد و بتواند تمامی نمونه‌های ورودی را به کلاس صحیح اختصاص دهد، احتمال دقت مدل برابر ۱ می‌گردد و احتمال خطا برابر ۰. حال این قضیه قابل تعمیم به مسائلی با تعداد کلاس‌های بیشتر نیز می‌باشد.

با توجه مثال فوق، در بدترین حالت احتمال تشخیص درست مدل برابر $\frac{1}{M}$ می‌باشد و بنابراین در این حالت احتمال خطا برابر $1 - \frac{1}{M}$ می‌باشد. در نتیجه گزینه (۳) صحیح است.

قابل ذکر است که گزینه (۱) نیز نادرست می‌باشد. زیرا در صورتی که مدل تمام نمونه‌ها را درست دسته‌بندی نماید، مقدار خطا برابر صفر می‌گردد اما در این گزینه این اتفاق رخ نمی‌دهد.

۲۳- گزینه «۴» مفهوم سؤال کمی مبهم مطرح شده است. زیرا روش Bootstrap روشی با جایگذاری می‌باشد و در این صورت برای ساخت هر درخت به طور میانگین از کل داده‌ها استفاده می‌گردد که برابر است با n . با توجه به گزینه‌های مطرح شده، می‌توان به تعبیری توضیحات زیر را ارائه نمود:

در روش نمونه‌گیری Bootstrap، احتمال انتخاب هر نمونه $\frac{1}{n}$ می‌باشد و در این صورت احتمال این که آن نمونه انتخاب نگردد نیز برابر $1 - \frac{1}{n}$ است. از آنجایی که ما n نمونه را با جایگذاری نمونه‌برداری می‌کنیم، احتمال این که یک مثال خاص در هیچ‌یک از n نمونه انتخاب نشده باشد، برابر می‌گردد با $(1 - \frac{1}{n})^n$ ؛ بنابراین در نقطه مقابل، احتمال این که یک مثال خاص حداقل یک بار انتخاب گردد برابر می‌گردد با $1 - (1 - \frac{1}{n})^n$. پس به طور میانگین هر

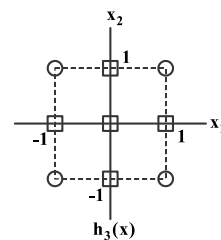
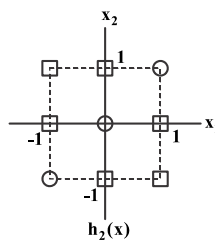
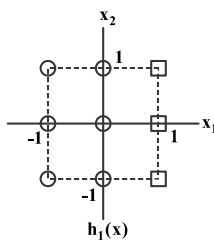
مثال به اندازه $1 - (1 - \frac{1}{n})^n$ بار در نمونه‌ی Bootstrap انتخاب می‌گردد.

از آنجاییکه نمونه Bootstrap در کل شامل n مثال است، تعداد مورد انتظار نمونه‌های متمایز آن $n \times [1 - (\frac{1}{n})^n]$ می‌باشد. زمانی که n مقداری بزرگ است، مقدار این معادله نیز برابر با n می‌گردد.

در نهایت به طور خلاصه، میانگین اندازه‌ی نمونه Bootstrap برای ساخت هر درخت n می‌باشد که به طور تقریبی، مقدار $1 - (\frac{1}{n})^n$ نیز قابل نمایش است.

نکته: تعداد ویژگی‌ها (d) تأثیری در اندازه‌ی نمونه‌های Bootstrap ندارد.

۲۴- گزینه «۱» یک پرسپترون یک دسته‌بند خطی است، یعنی تنها قادر به دسته‌بندی خطی داده‌ها می‌باشد. بنابراین داده‌ها می‌بایست خطی جداپذیر باشند. برای بررسی این موضوع کافی است هر سه فرضیه داده شده را مصورسازی نماییم سپس مشخص می‌گردد کدام فرضیه‌ها خطی جداپذیر نیستند. بر طبق اطلاعات سؤال، تنها فرضیه ۱ خطی جداپذیر بوده و دو فرضیه دیگر به صورت خطی قابل تفکیک نیستند. (در نمودارها یک کلاس با استفاده از دایره و دیگری با استفاده از مربع رسم شده‌اند)



۲۵- گزینه «۳» در الگوریتم بیز، ویژگی‌ها به صورت وابسته در نظر گرفته شده و احتمال توأمان برای آن‌ها محاسبه می‌گردد، در صورتی که در روش بیز ساده، ویژگی‌ها مستقل در نظر گرفته شود، پس از محاسبه احتمال هر ویژگی به صورت جداگانه، آن‌ها را در یکدیگر ضرب می‌نماییم. در نهایت در هر دو روش، داده جدید به کلاسی متعلق می‌گردد که احتمال Posterior بیشتری داشته باشد.

ابتدا احتمال Prior هر کلاس را محاسبه می‌نماییم:

$$p(+) = \frac{7}{12} \quad p(-) = \frac{5}{12}$$

نکته: در جدول داده شده برای این سؤال، به اشتباه عنوان ویژگی Weight و Height جابه‌جا نوشته شده است.

$$P(+ | \text{Weight} = \lambda^0, \text{Height} = \lambda^0) = \frac{P(\text{Weight} = \lambda^0, \text{Height} = \lambda^0 | +)P(+)}{P(\text{Weight} = \lambda^0, \text{Height} = \lambda^0)} \quad (1)$$

روش بیز:

چون مخرج معادله در پاسخ تأثیری ندارد آن را نادیده می‌گیریم. حال به صورت جداگانه مقدار بیشینه‌ی درست‌نمایی (ML) توأمان را برای هر دو کلاس محاسبه می‌نماییم:

$$P(\text{Weight} = \lambda^0, \text{Height} = \lambda^0 | +) = \frac{2}{7}$$

$$P(\text{Weight} = \lambda^0, \text{Height} = \lambda^0 | -) = \frac{1}{5}$$

حال با جایگذاری مقادیر فوق در رابطه‌ی (۱) مقدار احتمال Posterior را برای نقطه داده شده محاسبه می‌نماییم (در صورت تبدیل + در رابطه (۱) به -، احتمال کلاس منفی محاسبه می‌گردد):

$$P(+ | \text{Weight} = \lambda^0, \text{Height} = \lambda^0) = \frac{2}{7} \times \frac{7}{12} \approx 0.16$$

$$P(- | \text{Weight} = \lambda^0, \text{Height} = \lambda^0) = \frac{1}{5} \times \frac{5}{12} \approx 0.08$$

با توجه به مقدار احتمال Posterior بیشتر برای کلاس +، نتیجه می‌گیریم با روش بیز، داده به کلاس مثبت تعلق دارد.

$$P(+ | \text{Weight} = \lambda^0, \text{Height} = \lambda^0) = \frac{P(\text{Weight} = \lambda^0 | +) \times P(\text{Height} = \lambda^0 | +) \times P(+)}{P(\text{Weight} = \lambda^0, \text{Height} = \lambda^0)} \quad (2)$$

روش بیز ساده:



شناسایی الگو - یادگیری ماشین

کدام ۲۱- برای یک طبقه‌بند Hard-SVM دو کلاسه، ۴ عدد داده آموزشی متعادل (balanced) از هر دو کلاس، که به‌طور خطی جدایی‌پذیر هستند موجود است. کدام یک از مقادیر زیر نمی‌تواند تعداد بردارهای پشتیبان آن پس از اتمام فاز آموزش باشد؟

(۴) ۳

(۳) ۴

(۲) ۱

(۱) ۲

کدام ۲۲- در آموزش پرسپترون که داده‌ها یکی یکی در هر مرحله به آن داده می‌شوند و ضریب یادگیری η با افزایش مراحل i با فرمول $\eta_i = \frac{1}{1+i}$ کاهش می‌یابد، کدام مورد در خصوص تأثیر تغییرات ضریب یادگیری η با i درست است؟

(۱) اصلاً ارتباطی بر تأثیر داده‌ها ندارد. (۲) تأثیر داده‌های ابتدایی، کمتر از داده‌های انتهایی می‌شود.

(۳) تأثیر داده‌های انتهایی، کمتر از داده‌های ابتدایی می‌شود. (۴) تأثیر داده‌های انتهایی و داده‌های ابتدایی، تقریباً برابر می‌شود.

کدام ۲۳- کدام مورد زیر، می‌تواند ماتریس کواریانس یک مجموعه داده باشد؟

$$\begin{bmatrix} -3 & 3 \\ 3 & 3 \end{bmatrix} \quad (۴)$$

$$\begin{bmatrix} 1 & 4 \\ 1 & 2 \end{bmatrix} \quad (۳)$$

$$\begin{bmatrix} 2 & 3 \\ 3 & 1 \end{bmatrix} \quad (۲)$$

$$\begin{bmatrix} 3 & 2 \\ 2 & 4 \end{bmatrix} \quad (۱)$$

کدام ۲۴- فرض کنید می‌خواهیم سیستمی مبتنی بر یادگیری برای تبدیل صوت به متن ایجاد کنیم که برای هر صوت، متن مرتبط با آن را نیز داریم. داده‌های ما صوت‌هایی به زبان فارسی هستند، ولی لهجه‌های مختلفی در آنها وجود دارد که در مجموعه آموزشی، برجستگی برای لهجه‌ها وجود ندارد. اگر بخواهیم از تأثیر لهجه‌ها در یادگیری استفاده شود، فرایند یادگیری در این مسئله باید چگونه باشد؟

(۱) ابتدا مبتنی بر دسته‌بندی و سپس مبتنی بر خوشه‌بندی و سپس مبتنی بر دسته‌بندی

(۳) تنها مبتنی بر دسته‌بندی (۲) ابتدا مبتنی بر خوشه‌بندی و سپس مبتنی بر خوشه‌بندی

(۴) تنها مبتنی بر خوشه‌بندی

کدام ۲۵- در کدام یک از طبقه‌بندهای زیر، برای فاز تست لازم است همه داده‌های آموزشی یا بخشی از آنها استفاده شود؟

K-nn (۴)

Hard-SVM (۳)

MLP (۲) شبکه عصبی

(۱) درخت تصمیم

کدام ۲۶- در الگوریتم C-SVM، مسئله بهینه‌سازی مقید زیر را داریم:

$$J(\underline{w}, w_0, \underline{\alpha}) = \frac{1}{2} \|\underline{w}\|^2 + c \sum_{i=1}^N \alpha_i$$

$$\text{subject to: } \begin{cases} y_i(\underline{w}^T \underline{x}_i + w_0) - 1 + \alpha_i \geq 0 & i = 1, 2, \dots, N \\ \alpha_i \geq 0 & i = 1, 2, \dots, N \end{cases}$$

تابع لاگرانژ هم به‌صورت زیر نوشته شده است:

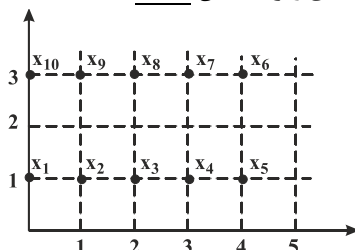
$$L(\underline{w}, w_0, \underline{\alpha}, \underline{\lambda}, \underline{\mu}) = \frac{1}{2} \underline{w}^T \underline{w} + c \sum_{i=1}^N \alpha_i - \sum_{i=1}^N \lambda_i (y_i(\underline{w}^T \underline{x}_i + w_0) - 1 + \alpha_i) - \sum_{i=1}^N \mu_i \alpha_i$$

اگر برای $C = 5$ مسئله حل شده باشد و برای داده $\underline{x}_8$ داشته باشیم $\lambda_8 = 2$ ، آنگاه کدام گزاره نادرست است؟

(۱) $\alpha_8 = 0$ و این داده درست طبقه‌بندی شده است. (۲) $\mu_8 > 0$ و این داده درست طبقه‌بندی شده است.

(۳) $\alpha_8 > 0$ و این داده روی مارجین (حاشیه) قرار دارد. (۴) $\mu_8 = 3$ و این داده روی مارجین (حاشیه) قرار دارد.

کدام ۲۷- در فضای دوبعدی، ۱۰ داده داریم که می‌خواهیم با روش K-means با معیار فاصله اقلیدسی یک خوشه‌بندی دوکلاسی انجام دهیم. برای شروع کار، دو نقطه دلخواه را به‌عنوان مرکز دو خوشه قرار می‌دهیم و الگوریتم را اجرا می‌کنیم و داده‌ها به‌جز دو نقطه اولیه، به‌ترتیب شماره اندیس به الگوریتم وارد می‌شود. نتیجه نهایی خوشه‌بندی با دو نقطه اولیه x_1 و x_{10} با نتیجه خوشه‌بندی با کدام دو نقطه اولیه در گزینه‌های زیر، یکسان نیست؟

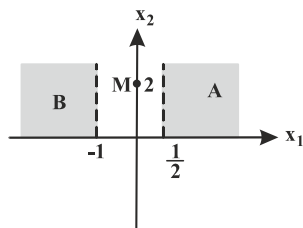
(۱) x_1 و x_6 (۲) x_5 و x_6 (۳) x_4 و x_7 (۴) x_2 و x_9

کدام ۲۸- برای ارزیابی یک تخمین‌زننده، بایاس و واریانس آن را محاسبه می‌کنند. کدام تخمین‌زننده بهتر است؟

(۱) بایاس آن کم و واریانس آن زیاد باشد. (۲) بایاس و واریانس آن کم باشد.

(۳) بایاس آن زیاد و واریانس آن کم باشد. (۴) بایاس و واریانس آن زیاد باشد.

کجه ۲۹- در یک طبقه‌بندی دو کلاسی با دو ویژگی، از کلاس ۱ به تعداد ۱۰۰ داده آموزشی داریم که در ناحیه A در شکل زیر و از کلاس ۲ به تعداد ۱۰۰ داده آموزشی داریم که در ناحیه B در شکل زیر هر دو به صورت تصادفی قرار دارند. می‌خواهیم به داده‌ای که در شکل با نقطه M مشخص شده است، با طبقه KNN با $K = 1$ و فاصله اقلیدسی یک برچسب بدهیم. کدام مورد درست است؟



(۱) بدون دیدن داده‌های آموزش، نمی‌توان برچسب را تعیین کرد.

(۲) همواره احتمال تعلق داده به هر دو کلاس، برابر است.

(۳) برچسب این داده، کلاس ۱ است.

(۴) برچسب این داده، کلاس ۲ است.

کجه ۳۰- اگر در یک مسئله طبقه‌بندی دو کلاسی در فضای ویژگی دوبعدی، ماتریس کوواریانس و میانگین کلاس‌ها مقادیر زیر باشند:

$$\Sigma = \begin{bmatrix} 4 & 0 \\ 0 & 9 \end{bmatrix} \quad \mu_1 = \begin{bmatrix} -1 \\ 0 \end{bmatrix} \quad \mu_2 = \begin{bmatrix} 1 \\ 0 \end{bmatrix}$$

آنگاه نقاطی که فاصله ماکزیمم از هر دو کلاس خود دارند، حداکثر مقدار فاصله اقلیدسی آنها از همان میانگین، چقدر خواهد بود؟

۶۰ (۴)

۵۰ (۳)

۴۰ (۲)

۳۰ (۱)

کجه ۳۱- برای طبقه‌بندی داده‌های XOR در فضای دوبعدی، کدام یک از طبقه‌بندهای زیر مناسب هستند؟

(a) طبقه‌بند درخت تصمیم

(b) طبقه‌بند Hard SVM

(c) یک نرون عصبی مصنوعی پرسپترون

(d) طبقه‌بند شبکه عصبی MLP

(۴) c و d

(۳) d و a

(۲) c و b

(۱) a و b

کجه ۳۲- اگر برای متغیر تصادفی x تابع چگالی احتمال داده نشده باشد، ولی میانگین μ و واریانس σ^2 آن معلوم باشد، با روش حداکثر آنتروپی، کدام یک از توابع چگالی احتمال زیر به دست می‌آید؟

$$f(x) = \begin{cases} \frac{1}{x_2 - x_1} & x_1 \leq x \leq x_2 \\ 0 & \text{در غیر این صورت} \end{cases} \quad \text{تابع (۱)}$$

$$f(x) = \begin{cases} \frac{x - \mu + \sigma}{2\sigma^2} & \mu - \sigma \leq x \leq \mu + \sigma \\ 0 & \text{در غیر این صورت} \end{cases} \quad \text{تابع (۲) اگر}$$

$$f(x) = \frac{\mu}{2\sigma} e^x \quad \text{تابع (۳)}$$

$$f(x) = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{(x-\mu)^2}{2\sigma^2}} \quad \text{تابع (۴)}$$

کجه ۳۳- کدام مورد، در خصوص ماشین‌های بردار پشتیبان SVM درست است؟

(۱) در یک مسئله جدایی‌ناپذیر خطی با افزایش C در روش C-SVM حاشیه (margin) بین دو کلاس و تعداد نقاط خطا افزایش می‌یابد.

(۲) روش C-SVM و روش v-SVM، با انتخاب مناسب پارامترها ممکن است به جواب یکسانی برسند.

(۳) ضرایب لاگرانژ پاسخ سیستم برای یک پایگاه داده خاص D، همواره منحصر به فرد است.

(۴) با افزودن داده آموزشی جدید به سیستم، همواره پاسخ سیستم تغییر می‌کند.

کجه ۳۴- یک ماتریس کوواریانس کدام ویژگی یا ویژگی‌ها از ۵ ویژگی زیر را همیشه دارا است؟

(a) وارون‌پذیری

(b) قطری بودن

(c) همه درایه‌ها مثبت هستند.

(d) متقارن بودن

(e) داشتن حداقل یک مقدار ویژه منفی

(۴) e و c, b, a

(۳) e و d, b

(۲) d

(۱) a و d

۱۵- گزینه «۳» الگوریتم BFS از صف استفاده می‌کند. درخت‌های عمق اول و سطح اول یک گراف الزاماً مشابه نیستند. عبارت‌های دوم و سوم صحیح هستند.

۱۶- گزینه «۲» در هرم بیشینه، بزرگ‌ترین مقدار در ریشه قرار دارد و دومین بزرگ‌ترین مقدار، یکی از فرزندان ریشه است. هزینه یافتن این عنصر $O(1)$ خواهد بود.

۱۷- گزینه «۱» هر گره در مسئله MAZE یک مسیر از ابتدا تا یک خانه‌ای را نشان می‌دهد که طول آن برابر ارتفاع آن گره است. با توجه به اینکه از هر گره حداکثر ۴ خروجی وجود دارد، $O(4^n)$ حد بالای پیچیدگی زمانی خواهد بود.

۱۸- گزینه «۴» در این مسئله، هزینه مسیر معنی ندارد و تنها رسیدن به اولین جواب مهم است. در نتیجه روش عقب‌گرد انتخاب مناسبی برای این مسئله است و شاخه و حد راه‌حل خوبی نیست. با در نظر گرفتن صفحه $n \times n$ و نیاز به عبور دقیقاً یک مرتبه از تمام خانه‌ها، (n^2) تعداد حالت‌های مسئله را نشان می‌دهد که البته تعداد قابل توجهی از آن‌ها هرس می‌شوند.

۱۹- گزینه «۴» حالت‌های ممکن تقسیم ماتریس‌ها، شاخه‌های درخت هستند. مجموع ضرب‌ها نیز حدپایین هستند تا بتوان مسیرهایی که به جواب بهینه نخواهند رسید را هرس کرد.

۲۰- گزینه «۱» هر سه عبارت صحیح هستند.

شناسایی الگو - یادگیری ماشین

۲۱- گزینه «۲» در Hard-SVM برای داده‌های خطی جدایی‌پذیر، مرز تصمیم‌گیری به گونه‌ای رسم می‌شود که حاشیه بین دو کلاس بیشینه شود. بردارهای پشتیبان نقاطی هستند که دقیقاً روی حاشیه قرار می‌گیرند یا درون آن هستند. بنابراین حداقل تعداد بردارهای پشتیبان ۲ (یک نقطه از هر کلاس) می‌باشد. همچنین حداکثر تعداد در تئوری می‌تواند به اندازه کل داده‌ها باشد، اما در عمل معمولاً کمتر است. با توجه به داده‌های متعادل (۴ نقطه از هر کلاس) و خطی جدایی‌پذیر بودن، امکان وجود تنها یک بردار پشتیبان وجود ندارد، زیرا مرز تصمیم‌گیری نیاز به نقاطی از هر دو کلاس دارد تا حاشیه را تعریف کند.

۲۲- گزینه «۳» در الگوریتم پرسپترون با ضریب یادگیری کاهشی، مقدار ضریب یادگیری با افزایش مراحل (افزایش i در مخرج کسر) کاهش می‌یابد. در مراحل ابتدایی (مقدار i کوچک‌تر) ضریب یادگیری بزرگ‌تر است، بنابراین تغییرات وزن‌ها تحت تأثیر قوی‌تر داده‌های اولیه قرار می‌گیرند. در مراحل انتهایی (مقدار i بزرگ‌تر) ضریب یادگیری کوچک‌تر می‌شود و تغییرات ناشی از داده‌های جدید کاهش می‌یابد. در نتیجه، داده‌های ابتدایی تأثیر بیشتری بر مدل نهایی دارند، در حالی که داده‌های انتهایی تأثیر کمتری خواهند داشت. این استدلال بر گزینه (۳) دلالت دارد. سازمان سنجش گزینه (۴) را به عنوان پاسخ صحیح اعلام کرده است، اما با توجه به توضیحات فوق، گزینه (۳) صحیح می‌باشد.

۲۳- گزینه «۱» ماتریس کواریانس باید دو شرط اصلی داشته باشد:

۱- متقارن بودن: عناصر غیرقطری سطر و ستون یکسان باشند.

۲- نیمه معین مثبت (Positive Semi-Definite): همه مقادیر ویژه ماتریس باید غیرمنفی باشند.

حال به بررسی گزینه‌ها می‌پردازیم:

گزینه (۱): - متقارن است ($2 = 2$).

$$\lambda = \frac{7 \pm \sqrt{17}}{2} \approx 1/44,5 / 56 \text{ (هر دو مثبت)}$$

- مقادیر ویژه:

گزینه (۲): - متقارن است ($3 = 3$).

$$\lambda = \frac{3 \pm \sqrt{37}}{2} \approx 5/54, -2/54 \text{ (یکی منفی)}$$

- مقادیر ویژه:

گزینه (۳): - متقارن نیست ($4 \neq 1$).

گزینه (۴): - متقارن است ($3 = 3$).

$$\lambda = \pm \sqrt{-18} \text{ (موهومی)}$$

- مقادیر ویژه:

۲۴- گزینه «۲» از آنجا که لهجه‌ها در داده‌های آموزشی برچسب‌گذاری نشده‌اند، می‌توان با خوشه‌بندی، صوت‌ها را بر اساس شباهت‌های آوایی (مانند ویژگی‌های فرکانسی، الگوی تلفظ و غیره) به گروه‌هایی تقسیم کرد. این گروه‌ها احتمالاً نماینده لهجه‌های مختلف هستند. پس از شناسایی خوشه‌ها (لهجه‌ها)، می‌توان از این اطلاعات به‌عنوان ویژگی‌های اضافی در مدل دسته‌بندی (تبدیل صوت به متن) استفاده کرد. برای مثال، مدل می‌آموزد که چگونه هر لهجه بر تلفظ کلمات تأثیر می‌گذارد و دقت تبدیل صوت به متن را بهبود می‌بخشد.

۲۵- گزینه «۴» K-nn (K - نزدیک‌ترین همسایه): این الگوریتم مبتنی بر نمونه (Instance-based) است و در فاز تست نیاز دارد تمام یا بخشی از داده‌های آموزشی را در حافظه نگه دارد تا نزدیک‌ترین همسایه‌ها را برای نمونه‌های جدید محاسبه کند. بدون دسترسی به داده‌های آموزشی، K-nn نمی‌تواند عمل کند.

در رابطه با سایر گزینه‌ها (درخت تصمیم، Hard-SVM، MLP): این مدل‌ها پس از آموزش، پارامترهای نهایی (مانند ساختار درخت، ابرصفحه، یا وزن‌های شبکه عصبی) را ذخیره می‌کنند. در فاز تست، تنها از این پارامترها استفاده می‌شود و نیازی به داده‌های آموزشی اصلی نیست.

۲۶- گزینه «۳» شرایط بهینه‌سازی (KKT) شامل موارد زیر است:

$$\frac{\partial L}{\partial \alpha_i} = c - \lambda_i - \mu_i = 0 \Rightarrow \lambda_i + \mu_i = c$$

- شرط ایستایی نسبت به α_i :

$$\lambda_i [y_i (w^T x_i + w_0) - 1 + \alpha_i] = 0, \mu_i \alpha_i = 0$$

- شرط تکمیلی:

حال برای داده X_5 با شرایط داده شده تحلیل می‌کنیم. مقدار $C = 5$ است. داده X_5 برای آن $\lambda_5 = 2$ داده شده است. طبق شرط ایستایی:

$$\lambda_5 + \mu_5 = 5 \Rightarrow \mu_5 = 5 - 2 = 3$$

$$\mu_5 \alpha_5 = 0$$

همراه با شرط تکمیلی برای α_5 :

$$y_5 (w^T x_5 + w_0) - 1 + \alpha_5 = 0$$

چون $0 < \mu_5 = 3$ داریم، بنابراین به ناچار $\alpha_5 = 0$ همچنین، از تکمیلی بودن قید اصلی نتیجه می‌گیریم:

$$y_5 (w^T x_5 + w_0) = 1$$

با قرار دادن $\alpha_5 = 0$ ، خواهیم داشت:

بنابراین، داده X_5 دقیقاً روی مارجین قرار دارد و طبقه‌بندی آن صحیح است چون شرط $y_i (w^T x_i + w_0) \geq 1$ برقرار است.

حال به بررسی گزاره‌ها می‌پردازیم:

گزاره (۱) درست است، چون $\alpha_5 = 0$ و داده طبق شرط روی مارجین (که به‌عنوان طبقه‌بندی صحیح در نظر گرفته می‌شود) قرار دارد.

گزاره (۲) درست است، چون $0 < \mu_5 = 3$

گزاره (۳) نادرست است، چرا که محاسبات نشان می‌دهد $\alpha_5 = 0$ و نه بزرگ‌تر از صفر.

گزاره (۴) درست است، همانطور که محاسبه کردیم $\mu_5 = 3$ و داده روی مارجین قرار دارد.

۲۷- گزینه «۱» با فرض اینکه الگوریتم k-means به روش کلاسیک (ابتدا تخصیص داده به نزدیک‌ترین مرکز و سپس به‌روزرسانی مراکز) اجرا شود و داده‌های باقی‌مانده (x_7 تا x_9) به‌ترتیب شماره وارد شوند، روند کار به شرح زیر خواهد بود:

گام ۱: تعیین مراکز اولیه

- خوشه ۱: مرکز اولیه $x_1 = (0, 1)$

- خوشه ۲: مرکز اولیه $x_{10} = (0, 3)$

گام ۲: تخصیص نقاط به خوشه‌ها (با توجه به فاصله اقلیدسی)

برای هر نقطه باقی‌مانده، فاصله از مراکز فعلی را محاسبه می‌کنیم:

$$x_7 = (2, 1)$$

$$x_7 = (1, 1)$$

- فاصله از $(0, 1)$: ۲

- فاصله از $(0, 1)$: ۱

- فاصله از $(0, 3)$: $2 / \sqrt{2}$

$$\sqrt{(0-1)^2 + (3-1)^2} \approx 2 / \sqrt{2}$$

→ خوشه ۱

→ تخصیص به خوشه ۱

$$x_8 = (4, 1)$$

$$x_8 = (3, 1)$$

- فاصله از $(0, 1)$: ۴

- فاصله از $(0, 1)$: ۳

- فاصله از $(0, 3)$: $4 / \sqrt{2}$

- فاصله از $(0, 3)$: $3 / \sqrt{2}$

→ خوشه ۱

→ خوشه ۱

$$\Delta - X_6 = (4, 3)$$

- فاصله از $(0, 1) : 4 / 472 \approx 20$

- فاصله از $(0, 3) : 4$

→ خوشه ۲

$$\Delta - X_7 = (3, 3)$$

- فاصله از $(0, 1) : 3 / 606 \approx 20$

- فاصله از $(0, 3) : 3$

→ خوشه ۲

$$\Delta - X_8 = (2, 3)$$

- فاصله از $(0, 1) : 2 / 828 \approx 20$

- فاصله از $(0, 3) : 2$

→ خوشه ۲

$$\Delta - X_9 = (1, 3)$$

- فاصله از $(0, 1) : 2 / 236 \approx 20$

- فاصله از $(0, 3) : 1$

→ خوشه ۲

پس از این تخصیص اولیه، خوشه‌بندی اولیه به صورت زیر است:

- خوشه ۱: X_1, X_2, X_3, X_4, X_5

- خوشه ۲: X_6, X_7, X_8, X_9 و X_{10}

گام ۳: به‌روزرسانی مراکز خوشه‌ها

محاسبه میانگین (centroid) هر خوشه:

- برای خوشه ۱ $(X_1, X_2, X_3, X_4, X_5)$

$$x = \frac{(0+1+2+3+4)}{5} = \frac{10}{5} = 2 \text{ میانگین محور } x$$

$$y = \frac{(1+1+1+1+1)}{5} = \frac{5}{5} = 1 \text{ میانگین محور } y$$

→ مرکز جدید خوشه ۱: $(2, 1)$

- برای خوشه ۲ $(X_6, X_7, X_8, X_9, X_{10})$

$$x = \frac{(0+4+3+2+1)}{5} = \frac{10}{5} = 2 \text{ میانگین محور } x$$

$$y = \frac{(3+3+3+3+3)}{5} = \frac{15}{5} = 3 \text{ میانگین محور } y$$

→ مرکز جدید خوشه ۲: $(2, 3)$

گام ۴: بررسی تخصیص‌ها با مراکز جدید

با مراکز جدید $(2, 1)$ و $(2, 3)$ ، فاصله هر نقطه به مراکز محاسبه می‌شود. به‌عنوان مثال:

$$X_1 = (0, 1)$$

- فاصله تا $(2, 1) : 2$

- فاصله تا $(2, 3) : 2 / 828 \approx 20$

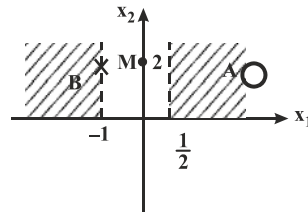
→ در خوشه ۱ باقی می‌ماند.

همین منطق برای سایر نقاط هم به همان نتیجه منتهی می‌شود. بنابراین تخصیص نهایی تغییر نمی‌کند و الگوریتم همگرا می‌شود. حال با بررسی هریک از گزینه‌ها با همین روش که توضیح داده شد، می‌توانیم دریابیم که تنها گزینه‌ای که نتیجه خوشه‌بندی با آن، متفاوت از نتایج فعلی می‌شود، گزینه (۱) می‌باشد.

۲۸- گزینه «۲» در ارزیابی یک تخمینگر، دو عامل اصلی وجود دارند که باید در نظر گرفته شوند: بایاس و واریانس. بایاس معیاری است که نشان می‌دهد تخمینگر به‌طور سیستماتیک چقدر از مقدار واقعی پارامتر فاصله دارد. به‌عبارتی، اگر میانگین تخمین‌های یک تخمینگر در نمونه‌های متعدد به مقدار واقعی نزدیک باشد، آن تخمینگر دارای بایاس کم است. در نتیجه، تخمینگری با بایاس کمتر، به‌طور میانگین، نزدیک به مقدار واقعی قرار می‌گیرد. واریانس نوسان تخمین‌ها را نسبت به هم در نمونه‌های مختلف بیان می‌کند. تخمینگری با واریانس پایین نسبت به تغییرات نمونه‌های انتخابی پایدارتر است؛ یعنی نتایج آن در نمونه‌های مختلف تغییر چندانی نمی‌کند. یک تخمینگر با واریانس بالا ممکن است در نمونه‌های مختلف بسیار متفاوت عمل کند، حتی اگر به‌طور میانگین خطای کمتری داشته باشد.

در انتخاب تخمینگر مناسب، هدف این است که هر دو معیار در بهترین حالت باشند؛ به این معنا که تخمینگر هم به‌صورت سیستماتیک نزدیک به مقدار درست باشد (بایاس پایین) و هم از نظر تغییرپذیری در نمونه‌های مختلف پایداری داشته باشد (واریانس پایین). به‌عبارت دیگر، تخمینگری که بتواند هم از خطاهای سیستماتیک جلوگیری کند و هم از نوسانات شدید در نتایج اجتناب کند، از نظر عملکرد آماری بهتر است.

۲۹- گزینه «۱» می‌دانیم که در روش 1-NN، با استفاده از نزدیک‌ترین همسایه به داده موردنظر، کلاس داده مورد نظر را تعیین می‌نماییم. با توجه به نمودار، داده M روی محور عمودی واقع شده است ($X_1 = 0$) و ناحیه A از $X_1 = \frac{1}{2}$ شروع می‌گردد و ناحیه B از $X_1 = -1$ در نگاه اول شاید به نظر برسد که چون مرز شروع ناحیه A به نقطه موردنظر نزدیک‌تر است تا مرز شروع ناحیه B به نقطه موردنظر، پس نقاطی که در ناحیه A وجود دارند به نقطه مورد نظر ما نزدیک‌تر هستند پس این نقاط کلاس نقطه موردنظر را تعیین می‌کنند. ولی با توجه به صورت سؤال «داده‌ها به صورت تصادفی در هر ناحیه قرار دارند». یعنی امکان دارد نزدیک‌ترین داده به داده موردنظر ما از کلاس B باشد. برای درک بهتر یک مثال نقض را در شکل زیر نشان می‌دهیم. در شکل زیر فرض شده نزدیک‌ترین داده به نقطه موردنظر، در کلاس A با دایره نشان داده شود و در کلاس B با X:



بنابراین بدون دیدن داده‌های آموزش، نمی‌توان برچسب را تعیین نمود.

۳۰- گزینه «۱» برای نوشتن معادله فاصله ماحالانوبیس، فرض می‌کنیم یک نقطه X نسبت به میانگین μ (هر کلاس) دارای فاصله ماحالانوبیس باشد. با توجه به داده‌های مسئله، فاصله ماحالانوبیس ثابت $d_M = 10$ است. بنابراین این معادله به شکل زیر نوشته می‌شود:

$$d_M(x, \mu) = \sqrt{(x - \mu)^T \Sigma^{-1} (x - \mu)}$$

$$(x - \mu)^T \Sigma^{-1} (x - \mu) = 10^2 = 100$$

$$\frac{y_1^2}{4} + \frac{y_2^2}{9} = 100 \quad \text{با توجه به اینکه} \quad \Sigma^{-1} = \begin{bmatrix} \frac{1}{4} & 0 \\ 0 & \frac{1}{9} \end{bmatrix} \quad \text{اگر } y = x - \mu \text{ قرار دهیم، معادله به صورت مقابل درمی‌آید:}$$

$$\frac{y_1^2}{400} + \frac{y_2^2}{900} = 1 \quad \text{برای به دست آوردن فرم استاندارد یک بیضی، کل معادله را بر 100 تقسیم می‌کنیم:}$$

در این فرم، طول‌های هون (نیم‌قطرهای اصلی) بیضی a و b به صورت زیر است:

$$\sqrt{400} = 20 \quad \text{- در جهت } y_1 \text{ (محور افقی):}$$

$$\sqrt{900} = 30 \quad \text{- در جهت } y_2 \text{ (محور عمودی):}$$

یادآوری: این نیم‌قطرها نشان‌دهنده بیشترین فاصله‌های اقلیدسی از میانگین در جهت‌های اصلی بیضی هستند.

پس از میانگین، حداکثر فاصله اقلیدسی (یعنی $\|y\|_2 = \sqrt{y_1^2 + y_2^2}$) برای نقاطی که دارای فاصله ماحالانوبیس ثابت 10 هستند، زمانی به دست می‌آید که نقطه روی محور دارای واریانس بزرگ‌تر قرار گیرد (در این مثال، محور y_2 که واریانس آن 9 است). به عبارت دیگر، بعد تبدیل از فاصله ماحالانوبیس به فاصله اقلیدسی در جهت یک واریانس خاص، برابر است با ضرب فاصله ماحالانوبیس در جذر واریانس مربوطه، بنابراین داریم:

$$\max \|y\| = 10 \times \sqrt{\max\{4, 9\}} = 10 \times \sqrt{9} = 10 \times 3 = 30$$

سازمان سنجش گزینه (۴) یعنی مقدار 60 را به عنوان گزینه صحیح انتخاب نموده است. در رد این گزینه می‌توان بیان نمود که این اشتباه احتمالاً به دلیل تفسیر «حداکثر فاصله» به عنوان «قطر» توسط سازمان سنجش می‌باشد. از آنجا که نقطه‌ای در یک جهت از میانگین فاصله 30 دارد، نقطه‌ای در جهت مخالف نیز 30 واحد فاصله دارد؛ بنابراین بیشینه فاصله بین دو نقطه‌ای که هر دو از میانگین به اندازه‌ی 30 فاصله دارند، برابر خواهد بود با 60. اما در سؤال خواسته شده «فاصله اقلیدسی آن‌ها از همان میانگین» که به معنای شعاع یا فاصله از مرکز (میانگین) است (که برابر با 30 می‌شود)، نه فاصله دو نقطه مقابل.

سازمان سنجش گزینه (۴) را به عنوان پاسخ صحیح اعلام کرده است، اما با توجه به توضیحات فوق، گزینه (۱) صحیح می‌باشد.

۳۱- گزینه «۳» با توجه به ماهیت مسئله XOR، متوجه می‌شویم که داده‌های XOR به طور خطی قابل تفکیک نیستند؛ به عبارت دیگر، مسئله XOR غیرخطی است. از این رو، طبقه‌بندهای خطی مانند یک پرسپترون تک‌لایه یا SVM خطی (که معمولاً به عنوان hard SVM نیز شناخته می‌شود) قادر به دسته‌بندی صحیح داده‌ها نخواهند بود.

پرسپترون تک‌لایه (مورد C): تنها قادر به پیدا کردن مرزهای تصمیم خطی است و برای مسئله غیرخطی XOR ناکافی است.

Hard SVM classifier (مورد b): اگر بدون استفاده از کرنل غیرخطی (مثلاً کرنل گاوسی یا چندجمله‌ای) به کار رود، مرز تصمیم خطی ایجاد می‌کند و نمی‌تواند مسئله XOR را تفکیک کند.

MLP (شبکه عصبی چندلایه) (مورد d): به دلیل وجود لایه‌های پنهان و توانایی ایجاد مرزهای تصمیم غیرخطی، می‌تواند مسئله XOR را به خوبی دسته‌بندی کند.

درخت تصمیم (مورد a): با تقسیم‌بندی بازگشتی فضای ویژگی‌ها حتی به کمک تقسیم‌بندی‌های محور و ایجاد مرزهای غیرخطی، قادر است مسئله XOR را تفکیک کند.

بنابراین، دو طبقه‌بندی که به خوبی می‌توانند داده‌های XOR را دسته‌بندی کنند، طبقه‌بند شبکه عصبی چندلایه (MLP) و طبقه‌بند درخت تصمیم هستند.

۳۲- گزینه «۴» برای یک متغیر تصادفی پیوسته که تنها میانگین (μ) و واریانس (σ^2) آن مشخص باشد، اصل حداکثر آنتروپی بیان می‌کند که تابع چگالی احتمالی با بیشترین آنتروپی همان توزیع نرمال (گاوسی) است. از میان چهار تابع ارائه‌شده، تنها تابع گزینه ۴ همان تابع چگالی نرمال است.

۳۳- گزینه «۲» بیاپید هر گزینه را بررسی کنیم:

گزینه (۱): در یک مسئله جدایی‌ناپذیر خطی با روش C-SVM، پارامتر C کنترل‌کننده مجازات (penalty) نقاطی است که داخل حاشیه یا اشتباه طبقه‌بندی شده‌اند. با افزایش C، مجازات‌ها بیشتر می‌شود؛ این بدان معناست که مدل تمایل دارد اشتباهات را کاهش دهد حتی به هزینه کاهش حاشیه. از این رو، با افزایش C معمولاً حاشیه (margin) کوچک‌تر و تعداد نقاط خطا کمتر می‌شود، نه اینکه افزایش پیدا کند. پس گزینه (۱) نادرست است.

گزینه (۲): روش‌های C-SVM و v-SVM هر دو هدف اصلی یکسانی دارند: یافتن یک hyperplane با حداکثر حاشیه و کنترل خطاهای مجاز. در واقع، این دو مدل با انتخاب مناسب پارامترهای C و ν می‌توانند به جواب‌های معادل دست پیدا کنند. بنابراین گزینه (۲) درست است.

گزینه (۳): اگرچه تابع تصمیم نهایی (که شامل ترکیب خطی از داده‌های پشتیبان و آستانه است) یکتا است، اما ضرایب لاگرانژ (Lagrange multipliers) ممکن است برای یک مجموعه داده خاص منحصر به فرد نباشند؛ دلیل آن وجود چندین جواب بهینه در مسئله بهینه‌سازی SVM در شرایط خاص (مانند همبستگی میان داده‌ها) است. پس ادعای «همواره منحصر به فرد» بودن ضرایب لاگرانژ صحیح نیست.

گزینه (۴): افزودن داده‌های آموزش جدید به سیستم همیشه باعث تغییر جواب نمی‌شود. اگر داده‌های افزوده شده در داخل حاشیه قرار نگیرند (یعنی نقاط جدید support vector نشوند)، جواب سیستم تحت تأثیر قرار نخواهد گرفت.

۳۴- گزینه «۲» برای یک مجموعه داده، ماتریس کوواریانس به صورت زیر تعریف می‌شود:

$$\Sigma = \text{cov}(X) \quad \text{که در آن } \Sigma_{ij} = \text{cov}(X_i, X_j)$$

از آنجا که کوواریانس بین X_i و X_j (یعنی $\text{cov}(X_i, X_j)$) برابر با کوواریانس بین X_j و X_i (یعنی $\text{cov}(X_j, X_i)$) است، نتیجه می‌شود که:

$$\Sigma_{ij} = \Sigma_{ji}$$

یعنی هر ماتریس کوواریانس دارای ویژگی متقارن بودن است.

حال به گزینه‌های ارائه‌شده نگاه می‌کنیم:

- (a) وارون‌پذیری: یک ماتریس کوواریانس در صورتی وارون‌پذیر است که تمام متغیرها خطی مستقل باشند؛ اما در بسیاری از موارد (مثلاً در حضور همبستگی کامل یا وابستگی خطی بین متغیرها) ممکن است ماتریس کوواریانس تکین (singular) شود. بنابراین وارون‌پذیری ویژگی عمومی و مشترکی نیست.

- (b) قطری بودن: یک ماتریس کوواریانس زمانی قطری است که متغیرهای موجود کاملاً بدون همبستگی باشند. اما در حالت کلی، متغیرها همبستگی‌هایی با یکدیگر دارند؛ بنابراین ماتریس کوواریانس لزوماً قطری نیست.

- (c) مثبت بودن همه درایه‌ها: این ویژگی مربوط به ماتریس‌های مثلثی (پایین یا بالا) است؛ درحالی‌که در ماتریس کوواریانس به دلیل متقارن بودن، هم درایه‌های بالایی و هم درایه‌های پایینی برابر هستند اما آن را به عنوان یک ماتریس مثلثی نمی‌توان در نظر گرفت.

- (d) متقارن بودن: همان‌طور که توضیح داده شد، این ویژگی همیشه در ماتریس‌های کوواریانس به قوت عمل می‌کند.

- (e) داشتن حداقل یک مقدار ویژه منفی: از آنجایی‌که ماتریس‌های کوواریانس مثبت نیم هستند، تمامی مقادیر ویژه آن‌ها صفر یا مثبت است؛ یعنی هیچ مقدار ویژه منفی‌ای وجود ندارد.

۳۵- گزینه «۱» با افزایش پیچیدگی مدل به طور کلی نتایج زیر به دست می‌آید:

- بایاس (Bias): با پیچیدگی بالاتر مدل، مدل قادر است الگوهای داده را بهتر به دست آورد و تقریب بهتری از تابع واقعی ارائه دهد؛ بنابراین بایاس کاهش می‌یابد.

- واریانس (Variance): از سوی دیگر، مدل‌های پیچیده معمولاً حساس‌تر به تغییرات داده‌های آموزشی شده و ممکن است نتایجی ناپایدار و متغیر داشته باشند، یعنی واریانس افزایش می‌یابد.