

حل تمرین های داده‌کاوی

مثال ۱: طبقه‌بندی با درخت تصمیم (Decision Tree)

مسئله: یک بانک می‌خواهد بر اساس سن، درآمد و سابقه‌ی اعتباری مشتریان، پیش‌بینی کند که آیا وام درخواستی تأیید می‌شود یا خیر.

گام‌های حل:

۱. ابتدا متغیر هدف را مشخص کنید: «تأیید وام» با دو مقدار بله/خیر.
 ۲. برای هر ویژگی (سن، درآمد، سابقه اعتباری)، میزان بهره‌ی اطلاعاتی (Information Gain) یا کاهش آنتروپی را محاسبه کنید.
 ۳. ویژگی‌ای که بیشترین بهره‌ی اطلاعاتی را دارد، به عنوان گره‌ی ریشه انتخاب می‌شود. در این مثال معمولاً «سابقه اعتباری» بیشترین تأثیر را دارد.
 ۴. درخت را به صورت بازگشتی برای هر شاخه ادامه دهید تا زمانی که گره‌ها خالص شوند یا معیار توقف برسد.
 ۵. برای جلوگیری از بیش‌برازش (Overfitting)، درخت را با روش Pruning هرس کنید.
- نکته‌ای که کمتر گفته می‌شود:** بسیاری در محاسبه‌ی آنتروپی دچار اشتباه محاسباتی می‌شوند چون فرمول لگاریتم پایه ۲ را با پایه ۱۰ اشتباه می‌گیرند. همیشه پیش از محاسبه، پایه‌ی مورد استفاده در سؤال را بررسی کنید.

حل عددی کامل:

وام تأیید شد؟	سابقه اعتباری	درآمد	سن	مشتری
خیر	بد	پایین	جوان	۱
بله	خوب	بالا	جوان	۲
بله	خوب	بالا	میان‌سال	۳
بله	خوب	پایین	مسن	۴
خیر	بد	بالا	مسن	۵
خیر	بد	پایین	میان‌سال	۶
خیر	بد	بالا	جوان	۷

بله	خوب	بالا	مسن	۸
-----	-----	------	-----	---

از ۸ نمونه، ۴ مورد «بله» و ۴ مورد «خیر» است، پس آنتروپی کل مجموعه برابر است با:

$$Entropy(S) = -(\frac{4}{8}) \times \log_2(\frac{4}{8}) - (\frac{4}{8}) \times \log_2(\frac{4}{8}) = -0.5 \times (-1) - 0.5 \times (-1) =$$

حالا داده را بر اساس «سابقه اعتباری» تقسیم می‌کنیم:

• شاخه‌ی «خوب»: مشتریان ۲، ۳، ۴، ۸ ← هر ۴ نفر «بله» = Entropy ← «»

• شاخه‌ی «بد»: مشتریان ۱، ۵، ۶، ۷ ← هر ۴ نفر «خیر» = Entropy ← «»

بهره‌ی اطلاعاتی:

$$Information\ Gain = 1 - [(\frac{4}{8}) \times 0 + (\frac{4}{8}) \times 0] =$$

چون بهره‌ی اطلاعاتی برابر با حداکثر مقدار ممکن (۱) است، «سابقه اعتباری» به‌تنهایی و بدون نیاز به هیچ ویژگی دیگری، کل داده را به‌صورت خالص تفکیک می‌کند و به عنوان گره‌ی ریشه‌ی درخت انتخاب می‌شود. در نتیجه درخت نهایی فقط یک سطح دارد: اگر سابقه اعتباری «خوب» باشد ← وام تأیید می‌شود؛ در غیر این صورت ← رد می‌شود.

مثال ۲: خوشه‌بندی با الگوریتم K-Means

مسئله: یک فروشگاه آنلاین می‌خواهد مشتریان را بر اساس میزان خرید و تعداد بازدید از سایت به سه گروه تقسیم کند.

گام‌های حل:

۱. تعداد خوشه‌ها ($k=3$) را انتخاب کنید. این عدد یا در صورت سؤال داده شده، یا با روش Elbow تعیین می‌شود.

۲. سه مرکز خوشه‌ی اولیه را به صورت تصادفی انتخاب کنید.

۳. فاصله‌ی اقلیدسی هر نقطه از مراکز خوشه را محاسبه کرده و هر نقطه را به نزدیک‌ترین مرکز اختصاص دهید.

۴. میانگین نقاط هر خوشه را دوباره محاسبه کرده و مرکز جدید را جایگزین کنید.

۵. مراحل ۳ و ۴ را تا زمانی که مراکز خوشه دیگر تغییر نکنند، تکرار کنید.

نکته‌ای که کمتر گفته می‌شود: انتخاب مراکز اولیه به صورت تصادفی می‌تواند نتیجه‌ی نهایی را تغییر دهد. در تمرین‌های امتحانی، همیشه فرض اولیه‌ی داده‌شده در سؤال را رعایت کنید، نه فرضی که خودتان انتخاب می‌کنید.

حل عددی کامل:

تعداد بازدید	خرید (واحد صد هزار تومان)	مشتري
۳	۲	C۱
۳	۳	C۲
۶	۶	C۳
۸	۸	C۴
۴	۲	C۵
۷	۷	C۶

فرض کنید $k=2$ و مراکز اولیه، (C_1, C_2, C_3) و C_4, C_5, C_6 باشند.

تکرار اول - محاسبه فاصله ی اقلیدسی هر نقطه تا دو مرکز:

خوشه	فاصله تا مرکز ۲	فاصله تا مرکز ۱	مشتري
۱	۷.۸۱	۰	C۱
۱	۷.۰۷	۱.۰۰	C۲
۲	۲.۸۳	۵.۰۰	C۳
۲	۰	۷.۸۱	C۴
۱	۷.۲۱	۱.۰۰	C۵
۲	۱.۴۱	۶.۴۰	C۶

مراکز جدید بر اساس میانگین اعضای هر خوشه:

• خوشه ۱: (C_1, C_2, C_5) مرکز جدید $= ((3)/2+3+2, (3)/3+3+4) = (2.33, 3.33)$

• خوشه ۲: (C_3, C_4, C_6) مرکز جدید $= ((3)/6+8+7, (3)/6+8+7) = (7, 7)$

تکرار دوم: با محاسبه ی مجدد فاصله ها از مراکز جدید، هر ۶ مشتری دقیقاً در همان خوشه ی قبلی باقی می ماند، یعنی الگوریتم همگرا شده است. نتیجه ی نهایی: مشتریان C۱، C۲ و C۵ یک گروه «کم خرید» و مشتریان C۳، C۴ و C۶ یک گروه «پر خرید و فعال» را تشکیل می دهند.



مثال ۳: قوانین انجمنی (Association Rules) و تحلیل سبد خرید

مسئله: بررسی تراکنش‌های یک سوپرمارکت برای یافتن الگوهایی مانند «مشتریانی که نان می‌خرند، اغلب کره هم می‌خرند.»

گام‌های حل:

۱. مقدار Support برای هر مجموعه آیتم را محاسبه کنید: تعداد تراکنش‌هایی که آن آیتم‌ها را با هم دارند، تقسیم بر کل تراکنش‌ها.

۲. مقدار Confidence قانون $A \rightarrow B$ را محاسبه کنید $\text{Support}(A, B)$: تقسیم بر $\text{Support}(A)$.

۳. مقدار Lift را برای سنجش قدرت واقعی رابطه محاسبه کنید. اگر Lift بیشتر از ۱ باشد، رابطه معنادار است.

۴. الگوریتم Apriori را برای حذف مجموعه‌های کم‌تکرار (زیر آستانه Minimum Support) اعمال کنید.

نکته‌ای که کمتر گفته می‌شود: خیلی‌ها Confidence بالا را با یک رابطه‌ی معنادار اشتباه می‌گیرند. یک قانون می‌تواند Confidence بالا داشته باشد اما Lift آن نزدیک یا کمتر از ۱ باشد، یعنی در واقع هیچ رابطه‌ی واقعی وجود ندارد.

حل عددی کامل:

آیتم‌های خریداری شده	تراکنش
نان، کره، شیر	T1

T _۲	نان، کره
T _۳	نان، شیر
T _۴	کره، شیر
T _۵	نان، کره، شیر

از مجموع ۵ تراکنش، محاسبات مربوط به قانون «نان ← کره» به این صورت است:

- Support(نان) = ۴ از ۵ تراکنش = ۰.۸
- Support(کره) = ۴ از ۵ تراکنش = ۰.۸
- Support(نان و کره با هم) = ۳ از ۵ تراکنش = ۰.۶ (T_۱, T_۲, T_۵)
- Confidence(نان ← کره) = Support(نان، کره) ÷ Support(نان) = ۰.۶ ÷ ۰.۸ = ۰.۷۵ (۷۵٪)
- Lift(نان ← کره) = Confidence ÷ Support(کره) = ۰.۷۵ ÷ ۰.۸ = ۰.۹۳۷

با اینکه Confidence این قانون ۷۵٪ و در نگاه اول بالا به نظر می‌رسد، مقدار Lift کمتر از ۱ است. این یعنی خرید نان عملاً احتمال خرید کره را افزایش نمی‌دهد - دو آیتم صرفاً به این دلیل با هم دیده می‌شوند که هر دو در کل تراکنش‌ها پرتکرار هستند، نه به این دلیل که واقعاً به هم مرتبط‌اند. این دقیقاً همان اشتباهی است که در نکته‌ی بالا به آن اشاره شد.

مثال ۴: رگرسیون خطی برای پیش‌بینی

مسئله: پیش‌بینی قیمت یک آپارتمان بر اساس متراژ آن.

گام‌های حل:

۱. داده‌ها را روی نمودار پراکندگی رسم کنید تا رابطه‌ی خطی بودن یا نبودن بررسی شود.
 ۲. با روش کمترین مربعات خطا (Least Squares) ضرایب خط رگرسیون (شیب و عرض از مبدأ) را محاسبه کنید.
 ۳. معادله‌ی خط را برای پیش‌بینی مقادیر جدید استفاده کنید.
 ۴. کیفیت مدل را با معیار R^2 ضریب تعیین (ارزیابی کنید).
- نکته‌ای که کمتر گفته می‌شود:** مقدار R^2 بالا همیشه به معنای مدل خوب نیست. اگر داده کم باشد یا متغیرهای دیگری در قیمت تأثیرگذار باشند (مثل موقعیت مکانی)، مدل ساده رگرسیون خطی می‌تواند گمراه‌کننده باشد.
- حل عددی کامل:**

آپارتمان	متراژ (متر مربع)	قیمت (میلیون تومان)
۱	۵۰	۵۰۰
۲	۶۰	۶۰۰
۳	۷۰	۶۵۰
۴	۸۰	۷۵۰
۵	۱۰۰	۹۰۰

میانگین متراژ = ۷۲ و میانگین قیمت = ۶۸۰. با فرمول کمترین مربعات خط:

$$\text{شیب خط} = \frac{\sum[(x-\bar{x})(y-\bar{y})]}{\sum[(x-\bar{x})^2]} = (b) = 1480 \div 11700 = 7.91$$

$$\text{عرض از مبدأ} = \bar{y} - b \times \bar{x} = (a) = 72 - 680 \times 7.91 = 110.8$$

$$\text{معادله ی خط رگرسیون: قیمت} = 110.8 + 7.91 \times \text{متراژ}$$

با این معادله، قیمت یک آپارتمان ۹۰ متری این گونه پیش بینی می شود:

$$\text{قیمت} = 110.8 + (90 \times 7.91) = 822.7 \approx 822.7 \text{ میلیون تومان}$$

با محاسبه ی مجموع مربعات خطا و مجموع مربعات کل، مقدار R^2 این مدل تقریباً ۰.۹۹۴ به دست می آید؛ یعنی حدود ۹۹.۴٪ از تغییرات قیمت با متراژ توضیح داده می شود که نشان دهنده ی برازش بسیار قوی مدل خطی روی این داده ی خاص است.



مثال ۵: تشخیص داده‌های پرت (Outlier Detection)

مسئله: در یک دیتاست تراکنش‌های بانکی، تراکنش‌های مشکوک به تقلب را شناسایی کنید.

گام‌های حل:

۱. میانگین و انحراف معیار داده را محاسبه کنید.
 ۲. با استفاده از روش Z-Score، هر داده‌ای که بیش از ۳ انحراف معیار از میانگین فاصله دارد را به عنوان داده‌ی پرت علامت‌گذاری کنید.
 ۳. به عنوان روش جایگزین، از $Q3 + 1.5 \times IQR$ (از $Q1 - 1.5 \times IQR$ تا $Q3 + 1.5 \times IQR$ پرت محسوب می‌شود).
 ۴. داده‌های پرت را بررسی کنید تا مشخص شود خطای ثبت داده هستند یا واقعاً یک الگوی غیرعادی (مثل تقلب) را نشان می‌دهند.
- نکته‌ای که کمتر گفته می‌شود:** حذف خودکار همه‌ی داده‌های پرت اشتباه است. در بسیاری از مسائل واقعی (مثل تشخیص تقلب)، خودِ داده‌ی پرت هدف اصلی تحلیل است، نه یک نویز که باید حذف شود.

حل عددی کامل:

مبلغ ۸ تراکنش (به هزار تومان): ۴۸، ۴۹، ۵۰، ۵۱، ۵۲، ۵۳، ۵۵، ۷۰۰

میانگین = ۱۳۲.۲۵ و انحراف معیار ≈ ۲۱۴.۶

Z-Score تراکنش مشکوک (۷۰۰):

$$Z = (۲۱۴.۶ - ۱۳۲.۲۵) \div ۲.۶۵ = ۳۰.۶۵$$

نکته‌ی مهم اینجا است: عدد ۲.۶۵ از آستانه‌ی رایج ۳ کمتر است، پس ممکن است این روش این تراکنش را «پرت» تشخیص ندهد! دلیل آن این است که همان مقدار ۷۰۰ خودش میانگین و انحراف معیار را به شدت بالا کشیده است - این یکی از ضعف‌های شناخته‌شده‌ی Z-Score در نمونه‌های کوچک است.

حالا همین داده را با روش QR ابررسی می‌کنیم:

- چارک اول $(Q_1) \approx ۴۹.۲۵$

- چارک سوم $(Q_3) \approx ۵۴.۵$

- $۵.۲۵ IQR = Q_3 - Q_1 =$

- سقف مجاز $= Q_3 + 1/5 \times IQR = ۵۴.۵ + ۷.۸۸ = ۶۲.۳۸$

چون مقدار ۷۰۰ به وضوح از سقف ۶۲.۳۸ بزرگ‌تر است، روش IQR این تراکنش را به درستی به عنوان داده‌ی پرت شناسایی می‌کند. این مثال دقیقاً نشان می‌دهد چرا در دیتاست‌های کوچک یا دارای پرت‌های شدید، IQR معمولاً قابل اعتمادتر از Z-Score است.

مثال ۶: کاهش ابعاد با تحلیل مؤلفه‌های اصلی (PCA)

مسئله: یک دیتاست با ۲۰ متغیر دارید و می‌خواهید آن را برای تجسم بهتر به دو بعد کاهش دهید.

گام‌های حل:

۱. داده‌ها را استانداردسازی کنید (میانگین صفر، واریانس یک).

۲. ماتریس کوواریانس بین متغیرها را محاسبه کنید.

۳. مقادیر ویژه (Eigenvalues) و بردارهای ویژه (Eigenvectors) را استخراج کنید.

۴. مؤلفه‌هایی که بیشترین واریانس داده را توضیح می‌دهند (معمولاً دو یا سه مؤلفه اول) را انتخاب کنید.

۵. داده‌ی اصلی را روی این مؤلفه‌های جدید تصویر کنید.

نکته‌ای که کمتر گفته می‌شود: PCA فقط یک روش کاهش بعد نیست، بلکه ابزاری قدرتمند برای حذف هم‌خطی چندگانه (Multicollinearity) پیش از اجرای مدل‌های رگرسیونی است؛ نکته‌ای که در اکثر تمرین‌های دانشگاهی نادیده گرفته می‌شود.

حل عددی کامل (نسخه‌ی ساده‌شده با ۲ متغیر):

متغیر Y	متغیر X	نمونه
---------	---------	-------

۱	۲.۵	۲.۴
۲	۰.۵	۰.۷
۳	۲.۲	۲.۹
۴	۱.۹	۲.۲
۵	۳.۱	۳.۰

میانگین $X = ۲.۰۴$ و میانگین $Y = ۲.۲۴$. پس از مرکزی کردن داده و محاسبه‌ی ماتریس کوواریانس:

$$Cov(X,X) = ۰.۹۳۸ \quad Cov(Y,Y) = ۰.۸۵۳ \quad Cov(X,Y) = ۰.۸۴۱$$

با حل معادله‌ی مقادیر ویژه روی این ماتریس، دو مقدار ویژه به دست می‌آید:

$$\bullet \quad \lambda_1 = ۱.۷۳۷ \text{ (مؤلفه‌ی اصلی اول)}$$

$$\bullet \quad \lambda_2 = ۰.۰۵۴ \text{ (مؤلفه‌ی اصلی دوم)}$$

درصد واریانس توضیح داده شده توسط مؤلفه‌ی اول:

$$\lambda_1 \div (\lambda_1 + \lambda_2) = ۱.۷۳۷ \div ۱.۷۹۱ \approx ۹۷\%$$

یعنی تنها با نگهداشتن یک بُعد (مؤلفه‌ی اول) به جای دو متغیر اصلی، حدود ۹۷٪ از اطلاعات داده حفظ می‌شود. در دیتاست‌های واقعی با ۲۰ متغیر، همین منطق برای انتخاب دو یا سه مؤلفه‌ی برتر از میان تمام مؤلفه‌ها استفاده می‌شود.



مقایسه ی ۶ روش داده کاوی

روش	نوع داده هدف	کاربرد اصلی	سطح پیچیدگی
درخت تصمیم	برچسب دار (Classification)	پیش بینی دسته	متوسط
K-Means	بدون برچسب (Clustering)	گروه بندی مشتریان	متوسط
قوانین انجمنی	تراکنشی	تحلیل سبد خرید	ساده تا متوسط
رگرسیون خطی	عددی پیوسته	پیش بینی مقدار	ساده
تشخیص داده پرت	آماري	شناسایی تقلب/خطا	متوسط
PCA	چندبعدي	کاهش ابعاد و تجسم	بالا

اشتباهات رایج در حل تمرین های داده کاوی

- انتخاب الگوریتم بدون بررسی نوع متغیر هدف (کیفی یا کمی بودن آن).
- نادیده گرفتن نرمال سازی داده قبل از اجرای الگوریتم های مبتنی بر فاصله مثل K-Means .
- تفسیر نکردن نتیجه ی نهایی و بسنده کردن به محاسبه ی عددی صرف.
- استفاده از کل داده برای آموزش مدل، بدون تفکیک داده آموزش و آزمون.
- اشتباه گرفتن همبستگی با رابطه ی علی در تفسیر نتایج رگرسیون یا قوانین انجمنی.

