

-- CONHECIMENTOS ESPECÍFICOS --

Julgue os itens a seguir, considerando que, em registros administrativos de servidores, um índice de risco de inconsistência (X) é uma variável aleatória descrita pela função de densidade de probabilidade $f(x) = \frac{e^{-x}}{(1+e^{-x})^2}$, em que $x \in \mathbb{R}$.

- 51 O primeiro quartil de X é igual a zero.
- 52 A mediana de X é igual a 1.
- 53 O momento de ordem 9, denotado por $E[X^9]$, é positivo.
- 54 O índice X é contínuo, assume valores reais, distribui-se simetricamente em torno de zero e possui desvio padrão superior a 1,5.
- 55 $P(X > 0) = 2 \cdot f(0)$.

O número diário de solicitações registradas no serviço de protocolo de certo órgão, denotado por N , segue uma distribuição de Poisson com média igual a 10. Em cada dia, a quantidade Q de solicitações protocoladas que tratam de requisição de informações segue uma distribuição condicional na forma $P(Q = q|N = n) = \frac{1}{2^n} \cdot \binom{n}{q}$, em que $q \in \{0, 1, \dots, n\}$.

Com base nessas informações, julgue os próximos itens.

- 56 $P(N = n|Q = q) = \frac{e^{-5}}{(n-q)!} \cdot 5^{n-q}$, para $n \geq q$.
- 57 $P(Q = 1) = 5 \cdot P(Q = 0)$.
- 58 $E[Q] = 5$.
- 59 $\text{Var}(Q = q|N = n) = 5$.
- 60 A quantidade Q segue uma distribuição de Poisson.

Considerando uma amostra aleatória simples $X_1, \dots, X_n$ extraída de uma distribuição uniforme contínua no intervalo $[a, a + 1]$, em que $a \in \mathbb{R}$, julgue os seguintes itens, a respeito da estatística $M_k = n^{-1} \sum_{i=1}^n X_i^k$, para $k > 0$.

- 61 À medida que $n \rightarrow +\infty$, se $a = -1/2$, $\sqrt{M_2}$ converge em probabilidade para 1.
- 62 $E[M_k] = a + 0,5$.
- 63 A variância de M_1 é igual a $1/12$.
- 64 De acordo com a lei fraca dos grandes números, se $a = 0$, M_k converge em probabilidade para $(k + 1)^{-1}$ à medida que $n \rightarrow +\infty$.
- 65 À medida que $n \rightarrow +\infty$, M_3 converge em distribuição para uma distribuição normal.

A respeito de inferência estatística, julgue os itens a seguir.

- 66 Amostra é um subconjunto da população.
- 67 População é um subconjunto de todos os elementos ou resultados sob investigação.
- 68 Mesmo sem se conhecer o formato da distribuição da população, o formato da distribuição de $\bar{x}$ de várias amostras coletadas será aproximadamente normal desde que o tamanho das amostras seja muito grande, isto é, $n \geq 20$.

Uma empresa fabricante de suco compra cerca de 15.000 laranjas por dia. Em uma amostra aleatória simples de 250 laranjas, coletada para o controle de qualidade, indicou-se que, em determinado dia, 75% das laranjas estavam adocicadas, boas para a produção de suco. No entanto, as laranjas adocicadas nesse dia correspondiam, na verdade, a apenas 70% do total.

A partir da situação hipotética precedente, e considerando que $\hat{p}$ represente a proporção da amostra de laranjas boas para a produção de suco, julgue os itens subsequentes.

- 69 O desvio padrão da distribuição amostral $\hat{p}$ pode ser calculado por $\sigma_{\hat{p}} = \sqrt{\frac{0,75 \cdot (0,25)}{250}}$.

- 70 A média da distribuição amostral $\hat{p}$ será $\mu_{\hat{p}} = 0,7$.

- 71 Se for coletada uma amostra aleatória simples de 30 laranjas, a distribuição amostral seguirá o formato de uma distribuição normal.

Espaço livre

Em um campeonato de basquete, aproximadamente 63% dos arremessos de 3 pontos são convertidos pelos atletas. Em uma temporada, mais de 1.000 partidas foram realizadas, tendo havido mais de 48.000 arremessos de 3 pontos. O técnico de um time, para estudar o desempenho dos jogos de sua equipe, resolveu tomar uma amostra aleatória em um único mês, analisando 600 arremessos de 3 pontos no período.

Com base nessas informações, e supondo que o percentual de acertos descrito seja preciso em toda a temporada, julgue o item a seguir.

72 O formato da distribuição amostral é viesado para a esquerda, considerada a distribuição da população.

Um pesquisador planeja perguntar a uma amostra aleatória de 100 alunos de uma escola se eles gostam da comida servida no refeitório escolar. O pesquisador utilizará os resultados da amostra para elaborar um intervalo de confiança de 90% para a proporção real de todo o corpo discente daquela escola que gosta da comida.

A partir dessas informações, julgue os próximos itens, em relação à interpretação do nível de confiança de 90%.

73 Se o pesquisador repetir esse processo e construir 20 intervalos a partir de amostras independentes, espera-se que 18 desses intervalos contenham a proporção real dos alunos que gostam da comida do refeitório.

74 Se o pesquisador repetir esse processo diversas vezes, cerca de 90% dos intervalos produzidos capturarão a proporção real dos alunos que gostam da comida do refeitório.

No que diz respeito à análise de regressão linear, julgue os itens que se seguem.

75 Os resíduos (a diferença entre o valor observado e o valor previsto pelo modelo) devem ser dependentes uns dos outros, com variância constante (ou homoscedasticidade), para que se estimem bons parâmetros do modelo.

76 Em um teste de hipóteses, o erro do tipo II consiste em rejeitar a hipótese alternativa quando esta for falsa.

77 Em um teste de hipóteses, o erro do tipo I consiste em rejeitar a hipótese nula quando esta for verdadeira.

78 O coeficiente de determinação é a razão da variação explicada sobre a variação total, de modo que, quando a dispersão em torno da reta de regressão for pequena em relação à variação total dos valores em torno de sua média, a variação explicada responderá por uma grande porcentagem da variação total e o coeficiente estará mais próximo de 1.

Nos dados a seguir, a variável y refere-se ao preço de venda de carros e a variável x , à sua quilometragem (em 1.000 km).

$$\begin{aligned}\sum x_i &= 505 \\ \sum y_i &= 21.600 \\ \sum x_i y_i &= 640.000 \\ \sum x_i^2 &= 21.825\end{aligned}$$

Com base nos dados precedentes, julgue os itens subsecutivos, considerando uma amostra $n = 14$.

79 Para um carro com determinada quilometragem, será obtido necessariamente o preço de venda exato indicado pela equação de regressão.

80 A equação de regressão resultante do método de mínimos quadrados é igual a $\hat{y} = 2.934 + 38,56x$.

Julgue os itens a seguir, a respeito de amostragem aleatória simples, estratificada, sistemática e por conglomerados.

81 Para ser representativa, uma amostra aleatória estratificada proporcional de 8 pessoas selecionadas de uma população composta por 15 homens e 25 mulheres deve ser constituída por 3 homens e 5 mulheres.

82 A aleatoriedade de uma amostra diz respeito ao fato de que os elementos da população devem ser selecionados por meio de sorteio não viciado, ou seja, qualquer elemento da população deve ter a mesma chance de fazer parte da amostra.

83 Se, para estudar o perfil dos consumidores de determinada área de uma cidade, uma equipe de consultoria dividir a área de estudo em conglomerados (casas), por meio de um esquema de amostragem por conglomerados, ela deve utilizar um esquema de amostragem acidental para entrevistar os participantes da pesquisa.

84 Erro puro é aquele que pode ocorrer quando se considera parte da população (amostra) para inferir o verdadeiro valor populacional (parâmetro).

85 O esquema de amostragem aleatória estratificada é utilizado quando a população em estudo é composta por unidades heterogêneas (subgrupos ou estratos) para a variável que se pretende estudar.

Espaço livre

O departamento de trânsito de determinada cidade deseja estimar o número médio de infrações registradas no período de um mês na via de maior fluxo da cidade. Registros anteriores apontam um desvio padrão de 8 casos mensais.

Considerando essa situação hipotética, com confiança de 95% ($Z_{\alpha/2} = 1,96$), e supondo que a média amostral não difira da média populacional em mais de 2 casos, julgue os próximos itens, em relação a tamanho amostral.

- 86** Na situação apresentada, o tamanho mínimo da amostra é inferior a 50.
- 87** Se o nível de confiança fosse reduzido e as demais quantidades, mantidas constantes, o número mínimo de infrações registradas por mês diminuiria.
- 88** Caso a margem de erro fosse reduzida pela metade, mantendo-se constante os demais parâmetros, o número mínimo de infrações duplicaria.

Julgue os itens a seguir, relativos à análise multivariada.

- 89** A função densidade de probabilidade de um vetor aleatório com dois elementos e que segue uma distribuição normal bivariada terá formato de esfera.
- 90** Na distribuição normal multivariada, os contornos da função densidade de probabilidade são elipsoides.
- 91** A função densidade de probabilidade relativa à distribuição normal multivariada é bimodal.
- 92** A função densidade de probabilidade relativa à distribuição normal multivariada é simétrica.

Uma análise multivariada que está sendo realizada envolve quatro variáveis padronizadas Z_1 , Z_2 , Z_3 e Z_4 , constituídas por duzentos registros administrativos. Para a identificação de estruturas latentes e a redução da dimensionalidade do problema, foram aplicadas a análise de componentes principais (ACP) e a análise fatorial exploratória (AFE) com extração por componentes principais. A seguir, são mostrados os autovalores para as componentes principais, que resultaram da decomposição espectral da matriz de correlações, bem como as cargas fatoriais não rotacionadas dos dois primeiros fatores extraídos.

λ_1	λ_2	λ_3
2,3	0,7	0,5

variável	fator 1	fator 2
Z_1	0,80	0,10
Z_2	0,75	0,30
Z_3	0,60	0,55
Z_4	0,10	0,85

Com base nessas informações, julgue os próximos itens.

- 93** Com base na regra de Kaiser (ou critério de Kaiser), é correto afirmar que apenas o primeiro componente principal deveria ser retido na ACP.
- 94** O quarto autovalor (λ_4) deve ser igual a zero.
- 95** Conjuntamente, as duas primeiras componentes principais explicam 75% da variância total.
- 96** A variável Z_2 possui maior associação com o fator 2 que com o fator 1.
- 97** A comunalidade da variável Z_4 é inferior a 0,7.

Um estudo está sendo realizado para a classificação de registros administrativos quanto à procedência de inconsistências. A base de dados foi coletada com informações quantitativas padronizadas de 120 registros, já rotulados como pertencentes a três grupos: A (correções técnicas), B (fraude documental) e C (erro de digitação). Inicialmente, foi aplicada uma análise discriminante linear (ADL) para que se pudesse verificar a separação entre os grupos. Em seguida, foi realizada uma análise de agrupamento (*cluster*) do tipo hierárquica, com distância euclidiana ao quadrado e método de Ward, sem se considerarem os rótulos originais.

Julgue os seguintes itens, relacionados à situação hipotética precedente.

- 98** Na análise de agrupamento hierárquico pelo método de Ward, os grupos são formados com base na maximização da variância intragrupo.
- 99** A ADL permite que se obtenham, no máximo, dois eixos discriminantes para o problema descrito.

unidade da Federação	tráfico de entorpecentes	lavagem de dinheiro	corrupção	total
estado A	32	20	15	67
estado B	25	18	20	63
estado C	28	15	12	55
estado D	40	10	5	55
estado E	10	25	30	65

Com base nos dados hipotéticos precedentes, relativos a crimes federais investigados, ao longo de 2024, em cinco unidades da Federação, julgue os itens seguintes.

- 100** O desvio padrão dos totais por unidade da Federação é inferior a 5.
- 101** A amplitude entre os totais de crimes investigados por unidade da Federação é 14.
- 102** A moda entre os tipos de crime apresentados é o tráfico de entorpecentes.
- 103** A unidade da Federação com a maior frequência relativa de corrupção é o estado E.
- 104** A média de casos por unidade da Federação é inferior a 60.

Espaço livre

Para certo conjunto de dados que mensuram uma variável quantitativa discreta com N medições, foram determinados a média m , o desvio padrão e , para uma dada proporção p , o p -quantil $q(p)$.

Com base nessa situação hipotética, julgue os itens seguintes, considerando que há acesso aos dados originais e a nenhuma outra informação acerca do conjunto.

- 105** Se $p > 0,9$ e se $q(p) < \frac{m}{2}$, há assimetria nos dados.
- 106** Mesmo sem acesso aos dados originais, é correto afirmar que a proporção de observações não inferiores a $q(p)$ é igual a $1 - p$.
- 107** Ainda que o conjunto de dados tenha valores extremos, a média e o desvio padrão são medidas adequadas para representar o conjunto, uma vez que são afetados apenas suavemente por esses extremos.
- 108** Mesmo sem acesso aos dados originais, é correto afirmar que, se p_1 indicar uma proporção diferente de p , o quantil $q(p_1)$ será diferente de $q(p)$.
- 109** A média e o desvio padrão fornecem informações suficientes para uma adequada percepção da eventual assimetria da distribuição dos dados.

Julgue os próximos itens, a respeito de Python.

- 110** Considere o código a seguir, desenvolvido em Python.

```
lista = [1, 2, 3, 4]
resultado = 0
x = 1
while x < len(lista) - 1:
    resultado += lista[x]
    x += 1
print(f"O Resultado é: {resultado}")
```

Ao ser executado, esse código apresentará o seguinte resultado.

O Resultado é: 6

- 111** Considere o código a seguir, desenvolvido em Python.

```
lista = [1, 2, 3, 4]
resultado = 1
for x in lista:
    if x % 2 == 0:
        resultado *= x
print(f"O Resultado é: {resultado}")
```

Ao ser executado, esse código apresentará o seguinte resultado.

O Resultado é: 8

- 112** Considere o código a seguir, desenvolvido em Python.

```
def calculo(lista):
    for n in lista:
        for x in range(2, n):
            if n % x == 0:
                n = n // x
                break
        else:
            print(f"O Resultado é: {n}")
lista = [2, 4, 6]
calculo(lista)
```

Ao ser executado, esse código apresentará o seguinte resultado.

O Resultado é: 2

- 113** Considere o código a seguir, desenvolvido em Python.

```
prisao = {
    "p1": {"inquerito": "INQ123",
            "anos_prisao": 5},
    "p2": {"inquerito": "INQ567",
            "anos_prisao": 3},
    "p3": {"inquerito": "INQ789",
            "anos_prisao": 6}
}
del prisao["p1"]
for pessoa, dados in prisao.items():
    dados["anos_prisao"] += 2
resultado = sum(dados["anos_prisao"] for
                dados in prisao.values())
print(f"O Resultado é: {resultado}")
```

Ao ser executado, esse código apresentará o seguinte resultado.

O Resultado é: 13

- 114** Considere o código a seguir, desenvolvido em Python.

```
conjunto = ((2,3), (4,5), (6,7))
x = 0
while x < len(conjunto):
    numeros = conjunto[x]
    resultado = (numeros[0] + numeros[1]) // 2
    x += 1
print(f"O Resultado é: {resultado}")
```

Ao ser executado, esse código apresentará o seguinte resultado.

O Resultado é: 6

- 115** Considere o código a seguir, desenvolvido em Python.

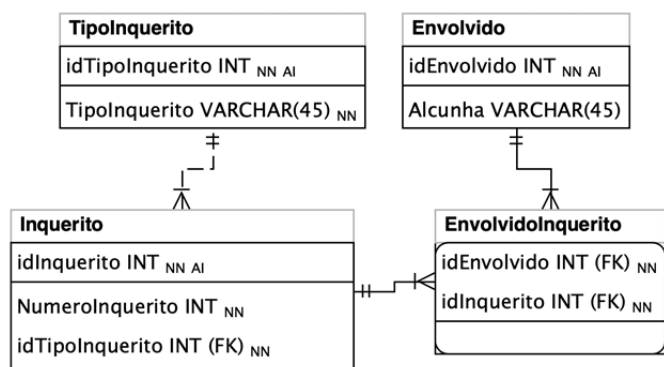
```
conjunto = ((2,7), (4,9), (6,4), (4,6))
resultado = 0
x = 0
while x < len(conjunto):
    numeros = conjunto[x]
    y = numeros[1]
    if y == 4:
        pass
    else:
        resultado += y
    x += 1
print(f"O Resultado é: {resultado}")
```

Ao ser executado, esse código apresentará o seguinte resultado.

O Resultado é: 22

Espaço livre

A seguir, são apresentados um modelo de banco de dados, consistindo de quatro tabelas, bem como os comandos SQL utilizados para a criação e a inserção de dados nessas tabelas.



criação das tabelas

```

CREATE TABLE Envolvido (
    idEnvolvido INT NOT NULL,
    Alcunha VARCHAR(45) NULL DEFAULT NULL,
    PRIMARY KEY (idEnvolvido));

CREATE TABLE TipoInquerito (
    idTipoInquerito INT NOT NULL,
    TipoInquerito VARCHAR(45) NOT NULL,
    PRIMARY KEY (idTipoInquerito));

CREATE TABLE Inquerito (
    idInquerito INT NOT NULL,
    NumeroInquerito INT NOT NULL,
    idTipoInquerito INT NOT NULL,
    PRIMARY KEY (idInquerito),
    FOREIGN KEY (idTipoInquerito)
    REFERENCES TipoInquerito (idTipoInquerito));

CREATE TABLE EnvolvidoInquerito (
    idEnvolvido INT NOT NULL,
    idInquerito INT NOT NULL,
    PRIMARY KEY (idEnvolvido, idInquerito),
    FOREIGN KEY (idEnvolvido)
    REFERENCES Envolvido (idEnvolvido),
    FOREIGN KEY (idInquerito)
    REFERENCES Inquerito (idInquerito));
  
```

inserção de dados nas tabelas

```

INSERT INTO Envolvido (idEnvolvido, Alcunha)
VALUES (10, 'Gargamel');
INSERT INTO Envolvido (idEnvolvido, Alcunha)
VALUES (20, 'Vingador');
INSERT INTO Envolvido (idEnvolvido, Alcunha)
VALUES (30, 'Esqueleto');

INSERT INTO TipoInquerito (idTipoInquerito,
TipoInquerito) VALUES (11, 'Inquérito Policial');
INSERT INTO TipoInquerito (idTipoInquerito,
TipoInquerito) VALUES (12, 'Inquérito Civil');
INSERT INTO TipoInquerito (idTipoInquerito,
TipoInquerito) VALUES (13, 'Notícia-Crime');
  
```

```

INSERT INTO Inquerito (idInquerito,
NumeroInquerito, idTipoInquerito) VALUES (1,111, 13);
INSERT INTO Inquerito (idInquerito,
NumeroInquerito, idTipoInquerito) VALUES (2,121, 12);
INSERT INTO Inquerito (idInquerito,
NumeroInquerito, idTipoInquerito) VALUES (3,131, 13);
INSERT INTO Inquerito (idInquerito,
NumeroInquerito, idTipoInquerito) VALUES (4,444, 11);
INSERT INTO Inquerito (idInquerito,
NumeroInquerito, idTipoInquerito) VALUES (5,555, 13);

INSERT INTO EnvolvidoInquerito (idEnvolvido,
idInquerito) VALUES (10, 1);
INSERT INTO EnvolvidoInquerito (idEnvolvido,
idInquerito) VALUES (20, 1);
INSERT INTO EnvolvidoInquerito (idEnvolvido,
idInquerito) VALUES (20, 4);
INSERT INTO EnvolvidoInquerito (idEnvolvido,
idInquerito) VALUES (30, 1);
INSERT INTO EnvolvidoInquerito (idEnvolvido,
idInquerito) VALUES (30, 3);
INSERT INTO EnvolvidoInquerito (idEnvolvido,
idInquerito) VALUES (30, 5);
  
```

Considerando as informações precedentes, julgue os próximos itens, relativos à SQL.

116 Considere o comando SQL a seguir.

```

SELECT TI.TipoInquerito,
COUNT(I.idTipoInquerito) AS quantidade
FROM Envolvido AS E
INNER JOIN EnvolvidoInquerito AS EI ON
E.idEnvolvido = EI.idEnvolvido
INNER JOIN Inquerito AS I ON EI.idInquerito =
I.idInquerito
INNER JOIN TipoInquerito AS TI ON
I.idTipoInquerito = TI.idTipoInquerito
GROUP BY TI.TipoInquerito;
  
```

Ao ser executado, esse comando apresentará o resultado seguinte.

TipoInquerito	quantidade
Inquérito Policial	1
Inquérito Civil	1
Notícia-Crime	3

117 Considere o comando SQL que se segue.

```

SELECT DISTINCT E.Alcunha,
COUNT(I.idInquerito)
AS Quantidade
FROM Envolvido AS E
INNER JOIN EnvolvidoInquerito AS EI ON
E.idEnvolvido = EI.idEnvolvido
INNER JOIN Inquerito AS I ON EI.idInquerito =
I.idInquerito
GROUP BY E.Alcunha;
  
```

Esse comando, ao ser executado, apresentará o resultado a seguir.

Alcunha	Quantidade
Gargamel	1
Vingador	1
Esqueleto	1

118 Considere o comando SQL a seguir.

```
SELECT DISTINCT I.NumeroInquerito
FROM Envolvido AS E
INNER JOIN EnvolvidoInquerito AS EI ON
E.idEnvolvido = EI.idEnvolvido
INNER JOIN Inquerito AS I ON EI.idInquerito =
I.idInquerito
GROUP BY I.NumeroInquerito
HAVING COUNT(I.idInquerito) > 1;
```

Ao ser executado, esse comando apresentará o resultado seguinte.

```
NumeroInquerito
111
```

119 Considere o comando SQL que se segue.

```
SELECT      DISTINCT      COUNT(E.Alcunha)      as
Quantidade
FROM Envolvido AS E
INNER JOIN EnvolvidoInquerito AS EI ON
E.idEnvolvido = EI.idEnvolvido
WHERE ALCUNHA LIKE '%E%';
```

Ao ser executado, esse comando apresentará o resultado a seguir.

```
Quantidade
1
```

120 Considere o comando SQL a seguir.

```
SELECT DISTINCT EI.idEnvolvido,
I.idTipoInquerito
FROM Envolvido AS E
INNER JOIN EnvolvidoInquerito AS EI ON
E.idEnvolvido = EI.idEnvolvido
INNER JOIN Inquerito AS I ON EI.idInquerito =
I.idInquerito
WHERE NumeroInquerito = 111;
```

Ao ser executado, esse comando apresentará o resultado seguinte.

```
idEnvolvido      idTipoInquerito
30                13
```

Espaço livre