

Exemplo de Exame

Respostas Comentadas

Set A
Versão 1.0

ISTQB® Testing with Generative AI Syllabus Specialist Level

Compatível com a versão 1.0 do Syllabus

International Software Testing Qualifications Board



Aviso de direitos autorais

Copyright © International Software Testing Qualifications Board (doravante ISTQB®).

ISTQB® é uma marca registrada do International Software Testing Qualifications Board.

Todos os direitos reservados.

Os autores, por meio deste documento, transferem os direitos autorais para o ISTQB®. Os autores (como atuais detentores dos direitos autorais) e o ISTQB® (como futuro detentor dos direitos autorais) concordaram com as seguintes condições de uso:

Extratos deste documento, para uso não comercial, podem ser copiados desde que a fonte seja citada.

Qualquer Provedor de Treinamento Credenciado pode usar este exemplo de exame em seu curso de treinamento se os autores e o ISTQB® forem reconhecidos como a fonte e os proprietários dos direitos autorais do exemplo de exame e desde que qualquer anúncio de tal curso de treinamento seja feito somente após o Credenciamento oficial dos materiais de treinamento ter sido recebido de um Conselho de Membros reconhecido pelo ISTQB®.

Qualquer indivíduo ou grupo de indivíduos pode usar este exemplo de exame em artigos e livros, desde que os autores e o ISTQB® sejam reconhecidos como a fonte e os proprietários dos direitos autorais do exemplo de exame.

É proibido qualquer outro uso deste exemplo de exame sem antes obter a aprovação por escrito do ISTQB®.

Qualquer Conselho Membro reconhecido pelo ISTQB® pode traduzir este exemplo de exame desde que reproduza o Aviso de Direitos Autorais acima mencionado na versão traduzida do exemplo de exame.

Responsabilidade pelo documento

O *ISTQB® Examination Working Group* é responsável por este documento.

Este documento é mantido por uma equipe central do ISTQB®, composta pelo *Syllabus Working Group* e pelo *Exam Working Group*.

Agradecimentos

Este documento foi produzido por uma equipe central do ISTQB® : Abbas Ahmad, Alessandro Collino, Bruno Legeard.

A equipe principal agradece à equipe de revisão do Grupo de Trabalho de Exames, ao Grupo de Trabalho do Programa e aos Conselhos de Membros por suas sugestões e contribuições

Histórico

Versão	Data	Comentários
V1.0	30	Conjunto de exames de amostra CT-GenAI A 1.0 Versão alfa
V1.0	17/04/2025	Conjunto de exames de amostra CT-GenAI A 1.0 Versão Beta
V1.0	10/06/2025	Conjunto de exames de amostra CT-GenAI A 1.0 Lançamento

Índice

Aviso de direitos autorais.....	2
Responsabilidade pelo documento	3
Agradecimentos	4
Histórico	5
Índice	6
Introdução	7
Objetivo deste documento.....	7
Instruções	7
Gabarito	8
Respostas Comentadas	9
Apêndice - Gabarito das Questões Complementares	23
Respostas Adicionais Comentadas	24

Introdução

Objetivo deste documento

Os exemplos de perguntas e respostas e as justificativas associadas neste exemplo de exame foram criados por uma equipe de especialistas no assunto e redatores experientes em perguntas com o objetivo de:

- Auxiliar os Conselhos de Membros e os Provedores de Exames do ISTQB® em suas atividades de elaboração de perguntas
- Fornecer aos provedores de treinamento e candidatos a exames exemplos de perguntas de exames

Essas perguntas não podem ser usadas como estão em nenhum exame oficial.

Observe que os exames reais podem incluir uma grande variedade de perguntas, e este exemplo de exame **não tem** a intenção de incluir exemplos de todos os tipos, estilos ou durações possíveis de perguntas. Este exemplo de exame pode ser mais difícil ou menos difícil do que qualquer exame oficial.

Instruções

Neste documento, você encontrará:

- Tabela de respostas, incluindo para cada resposta correta:
 - Nível K, objetivo de aprendizado e valor de pontos
- Conjuntos de respostas, incluindo para todas as perguntas:
 - Resposta correta
 - Justificativa para cada opção de resposta
 - Nível K, objetivo de aprendizado e valor da pontuação
- Conjuntos de respostas adicionais, inclusive para todas as perguntas [não se aplica a todos os exemplos de exame]:
 - Resposta correta
 - Justificativa para cada opção de resposta
 - Nível K, objetivo de aprendizagem e valor em pontos
- As perguntas estão contidas em um documento separado

Gabarito

(Q) Questão – (RC) Resposta correta – (LO) Objetivo de Aprendizagem – (K) Nível K – (P) Pontos

Q	RC	LO	K	P
1	d	GenAI-1.1.1	K1	1
2	c	GenAI-1.1.2	K2	1
3	b	GenAI-1.1.2	K2	1
4	d	GenAI-1.1.3	K2	1
5	b	GenAI-1.1.4	K2	1
6	a, e	GenAI-1.2.1	K2	1
7	d	GenAI-1.2.2	K2	1
8	b	GenAI-2.1.1	K2	1
9	c	GenAI-2.1.1	K2	1
10	b	GenAI-2.1.2	K2	1
11	a	GenAI-2.1.3	K2	1
12	a	GenAI-2.2.1	K3	2
13	b	GenAI-2.2.2	K3	2
14	c	GenAI-2.2.3	K3	2
15	c	GenAI-2.2.4	K3	2
16	a	GenAI-2.2.5	K3	2
17	a	GenAI-2.3.1	K2	1
18	a	GenAI-2.3.2	K2	1
19	c	GenAI-3.1.1	K1	1
20	d	GenAI-3.1.2	K3	2

Q	RC	LO	K	P
21	b	GenAI-3.1.3	K2	1
22	b	GenAI-3.1.4	K1	1
23	d	GenAI-3.2.1	K2	1
24	d	GenAI-3.2.2	K2	1
25	a	GenAI-3.2.2	K2	1
26	b	GenAI-3.2.3	K2	1
27	c	GenAI-3.3.1	K2	1
28	b, d	GenAI-3.4.1	K1	1
29	a	GenAI-4.1.1	K2	1
30	a	GenAI-4.1.2	K2	1
31	c	GenAI-4.1.3	K2	1
32	b	GenAI-4.2.1	K2	1
33	b	GenAI-4.2.2	K2	1
34	d	GenAI-5.1.1	K1	1
35	b	GenAI-5.1.2	K2	1
36	b	GenAI-5.1.3	K2	1
37	a	GenAI-5.1.4	K1	1
38	c	GenAI-5.2.1	K2	1
39	c	GenAI-5.2.2	K1	1
40	a	GenAI-5.2.3	K1	1

Respostas Comentadas

(Q) Questão – (RC) Resposta correta – (LO) Objetivo de Aprendizagem – (K) Nível K – (P) Pontos

Q	RC	Explicação/Fundamentação	LO	K	P
1	d	Considerando: • A AI simbólica usa um sistema baseado em regras para imitar a tomada de decisões humanas, representando o conhecimento usando símbolos e regras lógicas (1B). • O Machine Learning clássico usa uma abordagem orientada por dados que requer preparação de dados, seleção de recursos e treinamento de modelos (2D). • O Deep Learning usa redes neurais (estruturas de aprendizado de máquina) para aprender automaticamente recursos a partir de dados (3A). • AAI generativa usa técnicas de deep learning para criar dados, aprendendo e imitando padrões a partir de seus dados de treinamento (4C). Portanto: a) Não está correto. b) Não está correto. c) Não está correto. d) Está correto.	GenAI-1.1.1	K1	1
2	c	a) Não está correto porque as janelas de contexto não controlam a sequência temporal, mas limitam o escopo do texto que pode ser considerado simultaneamente. b) Não está correto porque as janelas de contexto afetam o escopo dentro da entrada atual, não os recursos de referência entre documentos. c) Está correto. Quando o texto excede o tamanho da janela de contexto, o modelo não pode considerar simultaneamente todas as partes do documento. À medida que o modelo processa novos tokens, ele deve efetivamente “esquecer” ou descartar os tokens que estão fora dos limites da janela de contexto. d) Não está correto porque as janelas de contexto não determinam as abordagens de análise, mas limitam a quantidade de texto que pode ser processada simultaneamente.	GenAI-1.1.2	K2	1

Q	RC	Explicação/Fundamentação	LO	K	P
3	b	a) Não está correto. Isso descreve incorporações, não tokenização. b) Está correto. A tokenização envolve a divisão do texto em unidades menores (tokens) que representam os blocos de construção das tarefas de geração de linguagem natural e permitem que os LLMs compreendam e gerem texto. Veja também a definição de “Tokenização” no programa (consulte “Apêndice D – Termos específicos de AI generativa”). c) Não está correto. Isso descreve a função geral dos LLMs, não a tokenização. d) Não está correto. Isso se refere à forma como os LLMs geram texto, não à tokenização.	GenAI-1.1.2	K2	1
4	d	Considerando: i. Não está correto. Embora os LLMs básicos possam gerar casos de teste, eles não são inerentemente “excelentes” nessa tarefa sem uma entrada estruturada. Gerar casos de teste sem uma entrada estruturada representa incorretamente suas capacidades. ii. Não está correto. Criar scripts de teste baseados em modelos requer a execução de instruções explícitas, que é a função dos LLMs ajustados para instruções, não dos LLMs de raciocínio. Os LLMs de raciocínio são especializados em inferência lógica e resolução de problemas, não na adesão rígida a modelos. iii. Não está correto. Os LLMs ajustados para instruções são projetados para seguir prompts estruturados, não para tomar decisões autônomas. Priorizar testes com base no feedback requer raciocínio, o que está além do seu escopo. iv. Correto. Os LLMs de raciocínio são explicitamente projetados para sintetizar várias fontes de dados e realizar inferência lógica, resolução de problemas e tomada de decisões. Detectar tendências e priorizar esforços de teste está alinhado com sua função na análise contextual e com a descrição fornecida. v. Está correto. Os LLMs ajustados às instruções são especificamente treinados para seguir instruções, incluindo a adesão aos formatos, estilos e regras de sintaxe solicitados. Gerar casos de teste em conformidade com a sintaxe da linguagem Gherkin é uma tarefa adequada às suas capacidades. Portanto: a) Não está correto. b) Não está correto. c) Não está correto. d) Está correto.	GenAI-1.1.3	K2	1
5	b	a) Não está correto. Os modelos de linguagem visual são um subconjunto dos LLMs multimodais, e não o contrário. b) Está correto. Os modelos de linguagem visual integram especificamente dados visuais e textuais, tornando-os um subconjunto dos LLMs multimodais. c) Não está correto. Os modelos de linguagem visual estão intimamente relacionados aos LLMs multimodais e se concentram em dados visuais e textuais. d) Não está correto. LLMs multimodais e modelos de linguagem visual têm escopos distintos e não são intercambiáveis.	GenAI-1.1.4	K2	1

Q	RC	Explicação/Fundamentação	LO	K	P
6	a, e	a) Está correto. Os LLMs podem analisar e esclarecer requisitos identificando ambiguidades e inconsistências. b) Não está correto. Gerar código de aplicativo completo não é uma capacidade essencial dos LLMs em tarefas de teste. c) Não está correto. Os LLMs podem dar suporte à automação de testes, sugerindo melhorias nos scripts de teste e identificando padrões de design, mas não executam diretamente scripts de teste nem automatizam totalmente scripts de teste sem supervisão humana. d) Não está correto. Os LLMs não podem realizar testes exploratórios manuais porque se trata de um processo intuitivo e adaptativo que requer criatividade, experiência e tomada de decisões humanas. e) Correto. Os LLMs podem gerar dados de teste diversificados, incluindo combinações e limites.	GenAI-1.2.1	K2	1
7	d	a) Não está correto. Os chatbots com AI são mais adequados para interações ad hoc, não para tarefas de teste específicas. b) Não está correto. Os chatbots com AI e os aplicativos de teste com LLM têm finalidades distintas e não são idênticos em termos de funcionalidade ou configurabilidade. c) Não está correto. Os aplicativos de teste com LLM se concentram na integração em processos de teste, não em prompts de conversação. Os chatbots com AI não requerem integração em ferramentas de teste e em processos de teste porque sua função principal é facilitar interações ad hoc, em vez de realizar tarefas de teste específicas. d) Correto. Os chatbots com AI fornecem interfaces conversacionais para tarefas ad hoc, enquanto os aplicativos de teste com LLM oferecem soluções personalizadas para necessidades específicas.	GenAI-1.2.2	K2	1
8	b	a) Não está correto. O contexto fornece informações básicas sobre o ambiente de teste ou o sistema que está sendo testado, não dados específicos (de entrada) a serem analisados. Embora mencione testes de desempenho, ele lista fontes de dados reais em vez de fornecer informações básicas. b) Está correto. A descrição lista fontes de dados específicas (relatórios de teste, logs de monitoramento, benchmarks de desempenho) que serão processadas pelo LLM para realizar a tarefa de análise. c) Não está correto. As restrições descrevem limitações ou considerações especiais sobre como a tarefa deve ser executada. A descrição lista fontes de dados reais, em vez de restrições à abordagem de análise. d) Não está correto. O formato de saída especifica a estrutura e as características esperadas da resposta do LLM. A descrição lista fontes de dados reais, sem qualquer referência à forma como os resultados devem ser apresentados.	GenAI-2.1.1	K2	1
9	c	a) Não está correto. A linha fornecida não descreve qual tarefa deve ser realizada, mas sim como apresentar os resultados. As instruções dizem ao LLM o que fazer, enquanto esta linha diz ao LLM como formatar a saída. b) Não está correto. A linha fornecida não está impondo restrições ou limitações ao processo de análise, mas especificando o formato de apresentação desejado. As restrições podem incluir aspectos como “excluir questões cosméticas”, enquanto esta linha trata da formatação da saída. c) Está correto. A linha fornecida corresponde à definição do programa de estudos sobre especificações de formato de saída que orientam como o LLM deve formatar a saída. d) Não está correto. A linha fornecida não fornece informações básicas, como o contexto da especificação dos requisitos. Em vez disso, trata-se da formatação da saída.	GenAI-2.1.1	K2	1

Q	RC	Explicação/Fundamentação	LO	K	P
10	b	a) Não está correto. A cadeia de prompts é, na verdade, definido pela decomposição de uma tarefa em prompts sequenciais, e não pelo “fornecimento de exemplos”. O prompting com poucos disparos é a técnica que fornece exemplos, e não aquela que “divide tarefas em subtarefas”. O meta prompting depende do LLM para revisar os prompts por si mesmo, e não do testador “refinar os prompts manualmente”. b) Está correto. O prompting com poucos disparos fornece orientação com exemplos, a cadeia de prompts divide as tarefas em etapas intermediárias (vários prompts) e o meta prompting usa AI para refinar seus próprios prompts de forma iterativa. c) Não está correto. O objetivo principal do meta prompting é o refinamento automático dos prompts pelo LLM, não “dividir as tarefas em etapas”. A cadeia de prompt é caracterizada pela decomposição passo a passo, não pelo “uso de exemplos”, como afirmado. O prompting de poucos disparos é o método que fornece exemplos ao modelo e não se concentra na “otimização manual dos prompts”. d) Não está correto. A cadeia de prompt é definida pela decomposição passo a passo, não pela “orientação sem exemplos”. O meta prompting depende do LLM para gerar/refinar prompts, não apenas da formulação definida pelo testador.	GenAI-2.1.2	K2	1
11	a	a) Está correto. O prompt do sistema permanece constante durante toda a sessão de interação e estabelece a estrutura fundamental para como o LLM deve responder. b) Não está correto. Isso descreve a função de um prompt do usuário, não de um prompt do sistema. c) Não está correto. O prompt do sistema não se ajusta dinamicamente; ele permanece constante. d) Não está correto. O prompt do sistema está oculto e não inclui entradas visíveis do usuário.	GenAI-2.1.3	K2	1

Q	RC	Explicação/Fundamentação	LO	K	P
12	a	Considerando: i. Está correto porque garante que o processo comece com a geração de condições de teste a partir dos requisitos, o que está alinhado com a descrição do syllabus que afirma que a análise de teste com GenAI envolve a geração de condições de teste com base na base de teste, por exemplo, nos requisitos. ii. Está correto porque esclarece a necessidade de critérios de aceite (condições de teste) para melhorar os resultados gerados pelo LLM e segue a orientação do programa de estudos de que os LLMs podem priorizar as condições de teste com base no nível de risco quando fornecidos com o contexto adequado. iii. Está correto porque orienta o LLM a realizar uma análise de cobertura, abordando todos os aspectos dos requisitos, o que corresponde à afirmação do programa de estudos de que os LLMs podem realizar análises de cobertura para determinar se todos os aspectos da base de teste estão cobertos. iv. Não está correto porque confiar em informações mínimas sem orientação não está alinhado com a técnica de encadeamento de prompts nem garante a eficácia. Essa opção tenta fazer tudo em uma única etapa, em vez de dividir o processo. v. Não está correto porque a identificação de defeitos não é fundamental para o objetivo do teste de gerar condições de teste priorizadas e identificar lacunas de cobertura. O cenário afirma especificamente que os requisitos são estáveis e já foram cuidadosamente revisados quanto a defeitos (que normalmente incluem ambiguidades e inconsistências). Portanto: a) Está correto. b) Não está correto. c) Não está correto. d) Não está correto.	GenAI-2.2.1	K3	2
13	b	a) Não está correto. Embora especifique o uso de exemplos predefinidos, não exige explicitamente o uso da sintaxe “Given-When-Then”, que é crucial para casos de teste no estilo Gherkin. Também indica dependência de práticas recomendadas vagas na definição das restrições e não garante especificamente o alinhamento com o critério de aceite. b) Está correto. Abrangente e aproveita exemplos predefinidos para orientar o LLM. c) Não está correto. Carece de ênfase no uso de exemplos predefinidos e indica dependência de práticas recomendadas vagas na definição das instruções. d) Não está correto. Concentra-se em casos extremos, mas negligencia a cobertura abrangente e o uso de exemplos para orientação.	GenAI-2.2.2	K3	2

Q	RC	Explicação/Fundamentação	LO	K	P
14	c	a) Não está correto. Aborda o agrupamento e a verificação cruzada, mas omite outras etapas críticas, como a separação dos resultados dos testes. b) Não está correto. Melhora a clareza das funções, mas não expande as instruções nem aborda etapas estruturadas. c) Está correto. Expande as instruções para incluir todas as etapas críticas da análise estruturada. Separar os resultados esperados dos resultados reais ajuda a identificar incompatibilidades de forma eficaz, o agrupamento de problemas facilita uma melhor priorização e reduz a redundância, e destacar as discrepâncias ajuda a concentrar a atenção nas descobertas mais importantes. d) Não está correto. Introduce restrições irrelevantes, desalinhando o prompt com os requisitos da tarefa.	GenAI-2.2.3	K3	2
15	c	a) Não está correto. Esclarece a função, mas não adiciona nenhuma instrução concreta que melhore a precisão, a aplicabilidade ou a interpretabilidade das métricas fornecidas. b) Não está correto. Adiciona uma tarefa de análise de risco que pode distrair o modelo das métricas principais e ainda não oferece nenhum mecanismo para tornar os resultados claros para as partes interessadas. c) Está correto. Expandir o formato de saída com um resumo em linguagem simples que interpreta as métricas e descreve as próximas etapas apoia diretamente a compreensão e a capacidade de ação das partes interessadas (consulte também o programa na seção 2.2.4: “Visualização e relatórios aprimorados das métricas de teste”). d) Não está correto. Apenas reafirma uma restrição existente e não fornece orientações específicas sobre como o LLM deve alcançar a interpretabilidade no nível das partes interessadas.	GenAI-2.2.4	K3	2
16	a	a) Está correto. O prompting com poucos disparos é ideal neste cenário, pois permite fornecer alguns exemplos (os casos de teste existentes com entradas e resultados esperados) para orientar o LLM. Ao ilustrar como essas regras de transformação são aplicadas para gerar novos casos de teste, o prompting com poucos disparos pode ajudar o LLM a compreender e replicar o processo para gerar casos de teste adicionais. b) Não está correto. Embora a cadeia de prompt possa ser usada, ele pode complicar desnecessariamente essa tarefa muito simples. c) Não está correto. Embora o meta-prompting possa ser usado, ele não é tão direto e específico quanto o prompting de poucos disparos para lidar com essa tarefa muito simples. d) Não está correto. O prompt de zero disparos não seria eficaz, pois não aproveitaria os casos de teste existentes como exemplos.	GenAI-2.2.5	K3	2
17	a	a) Está correto. A diversidade garante uma cobertura abrangente de casos extremos, e a taxa de sucesso na execução dos testes avalia a confiabilidade dos scripts de teste da API (consulte a descrição e o exemplo da métrica “diversidade” fornecida pelo programa na tabela da seção 2.3.1). b) Não está correto. Embora a precisão e a integridade sejam importantes, basear-se principalmente na eficiência de tempo não avalia totalmente a cobertura ou a confiabilidade da execução do teste. c) Não está correto. A precisão e a adequação contextual são valiosas, mas não abordam a diversidade nem avaliam exaustivamente o sucesso da execução dos testes para scripts de teste de API. d) Não está correto. A relevância e a adequação contextual são cruciais, mas sem considerar métricas como a taxa de sucesso da execução do teste ou a precisão, aspectos críticos da avaliação são perdidos.	GenAI-2.3.1	K2	1

Q	RC	Explicação/Fundamentação	LO	K	P
18	a	a) Está correto. A análise de saída examina as saídas do LLM em busca de imprecisões ou inconsistências. Ao classificar os resultados esperados incorretos que contradizem os requisitos, ela revela por que o prompt induziu o LLM ao erro e fornece insights concretos para o refinamento do prompt. b) Não está correto. O teste A/B de prompts é mais adequado para comparar diferentes versões de prompts do que para diagnosticar a causa de resultados esperados incorretos. c) Não está correto. Embora ajustar o comprimento e a especificidade das prompts possa melhorar a qualidade das respostas ao adicionar ou reduzir o contexto, isso não aborda diretamente a causa dos resultados esperados incorretos. d) Não está correto. Embora o uso das informações coletadas dos testadores sobre a utilidade e a clareza da saída gerada possa ajudar a refinar os prompts para melhor atender às necessidades de teste do mundo real, isso não aborda diretamente a causa dos resultados esperados incorretos.	GenAI-2.3.2	K2	1
19	c	a) Não está correto. Isso descreve erros de raciocínio, não alucinações. b) Não está correto. Isso descreve preconceitos na saída da IA, não alucinações. c) Está correto. Alucinações ocorrem quando o LLM gera resultados factualmente incorretos ou irrelevantes para uma determinada tarefa. Consulte o programa na seção 3.1.1. d) Não está correto. Isso descreve vieses devido à sub-representação nos dados de treinamento, não alucinações.	GenAI-3.1.1	K1	1
20	d	a) Não está correto. O gerenciamento do carrinho é explicitamente mencionado no briefing do projeto, tornando este caso de teste relevante. b) Não está correto. A aplicação do código de desconto é explicitamente mencionada no briefing do projeto, tornando este caso de teste relevante. c) Não está correto. Os e-mails de confirmação de pedido são mencionados explicitamente no briefing do projeto, portanto, este caso de teste é válido. d) Está correto. O gerenciamento da lista de desejos não é mencionado no briefing do projeto, tornando este o caso de teste mais provável de ser uma alucinação.	GenAI-3.1.2	K3	2
21	b	a) Não está correto. O esforço para ajustar os LLMs para tarefas de teste está associado a um treinamento adicional desses modelos em um conjunto de dados específico, permitindo que eles aprendam o conhecimento específico do domínio e as nuances necessárias para realizar essas tarefas de teste. O uso de formatos de dados de entrada claros e estruturados não afeta esse esforço. b) Está correto. Quando os dados de entrada são apresentados de forma clara e estruturada, os possíveis mal-entendidos são minimizados. A ambiguidade surge frequentemente de dados pouco claros e mal estruturados. c) Não está correto. A relevância do contexto depende mais do fornecimento das informações contextuais corretas do que apenas do uso de formatos de dados de entrada claros e estruturados. d) Não está correto. O uso de formatos de dados de entrada claros e estruturados não aumenta a criatividade dos LLMs na geração de resultados. Em particular, o uso desses formatos não incentiva respostas inovadoras, pois exige que os LLMs gerem respostas que sigam os formatos especificados.	GenAI-3.1.3	K2	1

Q	RC	Explicação/Fundamentação	LO	K	P
22	b	a) Não está correto. A taxa de aprendizagem é uma configuração (escolhida antes do início do treinamento) que controla o processo de aprendizagem das redes neurais. Essa configuração determina o quanto os pesos do modelo devem ser alterados durante o treinamento. Não está relacionada à variabilidade da inferência. Aumentá-la pode levar a uma convergência mais rápida durante o treinamento, mas não afeta a variabilidade das saídas durante a inferência. b) Está correto. A temperatura é um parâmetro que controla a aleatoriedade da saída, atuando na distribuição de probabilidade durante a inferência. Diminuir a temperatura reduz a aleatoriedade, levando a saídas mais consistentes. c) e d) não estão corretas. Embora definir uma semente aleatória possa melhorar a reprodutibilidade, isso por si só não restringe a distribuição de probabilidade durante a inferência. Além disso, o valor ser alto ou baixo é irrelevante.	GenAI-3.1.4	K1	1
23	d	a) Não está correto. Trata-se de uma questão de privacidade, uma vez que a “exposição involuntária de dados” está relacionada com os resultados da AI generativa (GenAI) que revelam acidentalmente informações confidenciais. b) Não está correto. Trata-se de uma questão de privacidade, pois envolve a falta de controle sobre como os dados confidenciais são armazenados ou processados. c) Não está correto. Trata-se de uma questão de privacidade relacionada ao não cumprimento de regulamentos como o Regulamento Geral sobre a Proteção de Dados (RGPD), o que acarreta riscos legais. d) Está correto. Um LLM não expõe dados confidenciais reais se tiver alucinações ao gerar dados de teste sintéticos, desde que não tenha sido treinado com dados confidenciais reais. Nesse caso, as alucinações seriam puramente sintéticas e baseadas em padrões e estruturas que o modelo aprendeu durante o treinamento. O LLM poderia gerar involuntariamente dados de teste sintéticos que correspondessem a dados confidenciais reais (e isso é definitivamente uma preocupação), mas não estaria expondo dados confidenciais reais. No entanto, essa geração involuntária é muito improvável.	GenAI-3.2.1	K2	1
24	d	a) Não está correto. A geração de código malicioso envolve “manipular um LLM para gerar backdoors (por exemplo, chamadas de comando externo) durante o uso”, não injetar dados falsos em conjuntos de dados de treinamento. b) Não está correto. A exfiltração de dados envolve “enviar solicitações projetadas para extrair dados confidenciais de treinamento”, não injetar dados falsos em conjuntos de dados de treinamento. c) Não está correto. A manipulação de solicitações envolve “introduzir dados que interrompem a saída da IA” durante o uso em tempo de execução, não injetar dados falsos em conjuntos de dados de treinamento. d) Está correto. De acordo com o programa da seção 3.2.2, o envenenamento de dados envolve “manipular dados de treinamento” com o exemplo específico de “fornecer avaliações falsas ao classificar os resultados de um relatório de teste gerado por IA”. A Questão descreve este vetor de ataque: um invasor está injetando resultados de teste (execução) falsificados no conjunto de dados de treinamento, o que manipula diretamente os dados de treinamento para comprometer a capacidade do LLM de recomendar com precisão estratégias ideais de cobertura de teste.	GenAI-3.2.2	K2	1

Q	RC	Explicação/Fundamentação	LO	K	P
25	a	Considerando: - A exfiltração de dados baseia-se no envio de solicitações destinadas a extrair dados confidenciais de treinamento. (1C) - A manipulação de solicitações é baseada na introdução de uma imagem que perturba a saída do LLM. (2D) - O envenenamento de dados baseia-se na manipulação de recomendações por meio de avaliações falsas usadas no treinamento. (3A) - A geração de código malicioso é baseada na manipulação de um LLM para gerar backdoors (por exemplo, chamadas de comandos externos) durante o uso. (4B) Portanto: a) Está correto. b) Não está correto. c) Não está correto. d) Não está correto.	GenAI-3.2.2	K2	1
26	b	a) Não está correto. Embora avaliar resultados comparando vários LLMs seja uma prática útil, neste caso, ela se concentra na qualidade dos resultados em termos de precisão, sem abordar os riscos à privacidade dos dados. b) Está correto. A anonimização é uma estratégia eficaz para mitigar os riscos à privacidade dos dados. c) Não está correto. O acesso irrestrito a dados confidenciais aumenta o risco de violação de dados confidenciais. d) Não está correto. Desativar a criptografia de dados confidenciais enfraquece a segurança e aumenta o risco de roubo de dados confidenciais.	GenAI-3.2.3	K2	1
27	c	a) Não está correto. A geração de imagens consome muito mais energia e CO₂ do que a geração de texto. O consumo de energia e as emissões de CO₂ estão normalmente correlacionados, a menos que seja especificado um fator-chave (como diferenças na fonte de energia). b) Não está correto. Pelo contrário, as pesquisas alimentadas por GenAI consomem consideravelmente mais energia do que as pesquisas tradicionais na web. c) Está correto. A geração de imagens requer significativamente mais recursos computacionais do que a geração de texto, tornando-a muito mais intensiva em energia. d) Não está correto. O impacto cumulativo da geração de texto em milhões de usuários não é insignificante.	GenAI-3.3.1	K2	1

Q	RC	Explicação/Fundamentação	LO	K	P
28	b, d	a) Não está correto. ISO/IEC 25010:2023 é uma norma (não mencionada no syllabus, mas mencionada no programa CTFL) que define um modelo de qualidade do produto relacionado às propriedades de qualidade dos produtos de TIC/software. Ela não aborda o uso de GenAI em testes de software. b) Está correto. ISO/IEC 23053:2022 é uma norma (mencionada na seção 3.4.1 do syllabus) que fornece uma estrutura para qualidade de dados, transparência e tolerância a falhas ao usar GenAI para testes. c) Não está correto. ISO/IEC/IEEE 29119-2:2021 é uma parte (volume) da norma (não mencionada no syllabus, mas mencionada no syllabus CTFL) que trata dos processos de teste. Não aborda o uso de GenAI em testes de software. d) Está correto. ISO/IEC 42001:2023 é uma norma (mencionada na seção 3.4.1 do syllabus) que especifica requisitos para gerenciar sistemas baseados em AI dentro de uma organização, fornecendo as melhores práticas para GenAI em testes de software e promovendo consistência e confiabilidade. e) Não está correto. ISO/IEC/IEEE 29119-3:2021 é uma parte (volume) da norma (não mencionada no syllabus, mas mencionada no syllabus CTFL) que trata da documentação de testes. Não aborda o uso de GenAI em testes de software.	GenAI-3.4.1	K1	1
29	a	a) Está correto. O back-end é responsável por recuperar dados de bancos de dados relacionais e vetoriais, combiná-los com as entradas do usuário e preparar o prompt personalizado para o LLM. b) Não está correto. O front-end serve como interface do usuário para enviar entradas, mas não prepara o prompt para o LLM. c) Não está correto. O componente de autenticação garante o acesso seguro, mas não está relacionado à recuperação de dados ou à preparação do prompt. d) Não está correto. O componente de pós-processamento refina a saída gerada pelo LLM, mas não lida com a preparação do prompt.	GenAI-4.1.1	K2	1
30	a	a) Está correto. O RAG recupera automaticamente as especificações relevantes e trechos de casos de teste históricos e os envia para o LLM, que então gera casos de teste precisos e contextualmente adequados. b) Não está correto. O RAG é mais adequado quando solicitado a recuperar informações específicas do banco de dados vetorial, e não todo o conjunto de dados. c) Não está correto. Revisar e refinar manualmente a consulta adiciona esforço desnecessário, o que contradiz o design do RAG de automatizar a recuperação e usar os dados recuperados dinamicamente na geração de respostas do LLM. d) Não está correto. Confiar nos dados internos do LLM ignora a capacidade central do RAG de aprimorar as respostas integrando informações externas e atualizadas.	GenAI-4.1.2	K2	1

Q	RC	Explicação/Fundamentação	LO	K	P
31	c	a) Não está correto. Agentes autônomos e semi-autônomos podem ser implementados como sistemas de agente único e/ou multiagente para melhorar a eficiência e a qualidade nos processos de teste. No entanto, as principais melhorias que esses agentes trazem referem-se ao aumento da eficiência e da qualidade nos processos de teste devido à sua capacidade de equilibrar a autonomia com a supervisão humana. b) Não está correto. Embora os agentes autônomos e semi-autônomos incorporem verificações para garantir a qualidade, essas verificações também são projetadas para serem eficientes. O uso de verificações automatizadas visa, na verdade, simplificar os processos de teste e melhorar a qualidade sem sacrificar a eficiência (e vice-versa), que, por sua vez, também é aprimorada. c) Está correto. Os agentes autônomos aumentam a eficiência executando tarefas de teste dentro dos processos de teste com intervenção humana mínima e usando verificações automatizadas projetadas para serem eficientes. Os agentes semi-autônomos envolvem supervisão humana estratégica, o que garante a qualidade dos resultados dos testes. Esses dois tipos de agentes permitem implementar uma abordagem equilibrada para alcançar eficiência e qualidade, aproveitando sua capacidade de operar com vários níveis de interação humana. d) Não está correto. A eliminação completa da verificação não é realista nem desejável. A verificação continua sendo crucial mesmo quando se utilizam tecnologias avançadas, como agentes autônomos e semiautônomos alimentados por LLM.	GenAI-4.1.3	K2	1
32	b	a) Não está correto. Esta é uma descrição correta do ajuste fino, que melhora o desempenho de um modelo em um domínio específico por meio de treinamento específico para a tarefa. b) Está correto. Esta afirmação está incorreta porque o ajuste fino não substitui o conhecimento geral do modelo, e o sobreajuste continua sendo um desafio potencial. c) Não está correto. Esta é uma descrição correta, pois o ajuste fino envolve o ajuste de parâmetros usando conjuntos de dados direcionados. d) Não está correto. Esta é uma descrição correta dos desafios associados ao ajuste fino, pois conjuntos de dados de alta qualidade são cruciais para uma adaptação eficaz.	GenAI-4.2.1	K2	1
33	b	a) Não está correto. As Operações de Modelos de Linguagem Grande (LLMOps) se concentram em gerenciar LLMs, não em impedir seu uso. b) Está correto. LLMOps gerencia LLMs ao longo de seu ciclo de vida, abordando privacidade, segurança e custo para testes. c) Não está correto. LLMOps se aplica a várias aplicações, não apenas a soluções baseadas em chatbots. d) Não está correto. LLMOps não visa automatizar totalmente todas as tarefas de teste ou eliminar a supervisão humana.	GenAI-4.2.2	K2	1
34	d	a) Não está correto. O uso de ferramentas GenAI não aprovadas não impõe as políticas de dados da organização. b) Não está correto. A AI oculta introduz riscos relacionados a contratos de licenciamento pouco claros, o que pode levar a disputas de propriedade intelectual. c) Não está correto. A AI oculta aumenta, em vez de reduzir, o risco de disputas de propriedade intelectual. d) Está correto. Ferramentas de AI não aprovadas geralmente carecem de medidas de segurança robustas, aumentando o risco de violações de dados ou acesso não autorizado.	GenAI-5.1.1	K1	1

Q	RC	Explicação/Fundamentação	LO	K	P
35	b	a) Não está correto. Embora certificações para LLMs específicos possam ser úteis (e mais disponíveis no futuro), os programas de treinamento devem ser projetados para garantir que os membros da equipe de teste tenham as habilidades técnicas necessárias para usar as ferramentas GenAI de maneira eficaz. b) Está correto. Selecionar LLMs compatíveis com a infraestrutura de teste existente e os requisitos de escalabilidade é um aspecto crítico de uma estratégia de GenAI. Ferramentas de teste e ambientes de teste são elementos fundamentais da infraestrutura de teste. c) Não está correto. A qualidade dos dados e a entrada estruturada e segura são cruciais para obter resultados confiáveis com GenAI. É necessário ter uma quantidade adequada de dados de entrada de alta qualidade (de acordo com os objetivos do teste), não o máximo de dados possível. d) Não está correto. A seção 5.1.2 afirma que uma estratégia de GenAI para testes de software deve coletar métricas específicas da tarefa para medir a eficácia das saídas do LLM. As métricas listadas na seção 2.3.1 (por exemplo, relevância, taxa de sucesso de execução, eficiência de tempo) são adaptadas para tarefas de teste, em vez de serem emprestadas integralmente do aprendizado de máquina supervisionado clássico.	GenAI-5.1.2	K2	1
36	b	a) Não está correto. Um critério de seleção fundamental envolve avaliar o desempenho do modelo nas tarefas de teste em relação aos benchmarks da organização, utilizando métricas como as apresentadas no próprio programa (consulte o syllabus na seção 5.1.3). O critério especificado pode ser relevante se as tarefas de teste envolverem geração de código. No entanto, não é aplicável a outros tipos de tarefas de teste. b) Está correto. Um critério de seleção importante envolve a avaliação de custos recorrentes (consulte o syllabus na seção 5.1.3), e esta resposta refere-se a um exemplo típico de custo recorrente. c) Não está correto. Avaliar o LLM em relação aos benchmarks da comunidade disponíveis publicamente para garantir total compatibilidade com eles não é um dos critérios-chave mencionados no programa (na seção 5.1.3) para selecionar um LLM apropriado para tarefas de teste específicas. d) Não está correto. Um critério de seleção importante envolve a avaliação de custos recorrentes (consulte o syllabus na seção 5.1.3), mas esta resposta refere-se a um exemplo típico de um custo não recorrente.	GenAI-5.1.3	K2	1

Q	RC	Explicação/Fundamentação	LO	K	P
37	a	O programa descreve as seguintes três fases principais: • Descoberta • Iniciação e definição de uso • Utilização e iteração Portanto: a) Está correto: menciona explicitamente todas estas três fases principais. b) Não está correto: não menciona explicitamente nenhuma dessas três fases (apenas menciona alguns objetivos e/ou atividades associados a essas fases). c) Não está correto: não menciona explicitamente nenhuma dessas três fases (apenas menciona alguns objetivos e/ou atividades associados a essas fases). d) Não está correto: não menciona explicitamente nenhuma dessas três fases (apenas menciona alguns objetivos e/ou atividades associados a essas fases).	GenAI-5.1.4	K1	1
38	c	a) Não está correto. Alucinações em LLMs são desafios intrínsecos às tecnologias de AI atuais, e os testadores não podem impedir que alucinações e erros de raciocínio ocorram. Em vez disso, os testadores devem ser capazes de (identificar e) mitigar os riscos de alucinações e erros de raciocínio (e preconceitos) ao testar com GenAI. b) Não está correto. Embora a proficiência em automação de testes seja valiosa, ela não aborda especificamente a integração da GenAI nos processos de teste. c) Está correto. O programa de estudos menciona “avaliar as capacidades do LLM” como uma competência fundamental. Isso implica naturalmente comparar modelos candidatos e selecionar aquele que pode ser adaptado ou ajustado para tarefas de teste específicas. Esse é um exemplo direto do conhecimento/habilidades que os testadores precisam para trabalhar de forma eficaz com LLMs. d) Não está correto. Refere-se às habilidades necessárias para desenvolver os LLMs, como pesquisadores de AI e cientistas de dados, e não aos testadores que usam GenAI para tarefas de teste. Esses testadores se concentram em garantir a qualidade dos dados de teste que usam diretamente ao interagir com LLMs.	GenAI-5.2.1	K2	1
39	c	a) Não está correto. Embora cursos externos com especialistas e prática possam ser benéficos, depender principalmente deles diminui a importância da prática interna e da construção de uma comunidade. Além disso, não é recomendável integrar a AI em todas as tarefas de teste diárias de forma abrupta. b) Não está correto. Embora a experimentação faça parte do processo de aprendizagem, a experimentação independente sem estrutura pode não levar ao desenvolvimento consistente ou eficaz de habilidades. c) Está correto. Concentrar-se no desenvolvimento de habilidades práticas por meio de atividades estruturadas, aprendizagem entre pares e comunidades de compartilhamento de conhecimento é uma abordagem recomendada. d) Não está correto. Depender principalmente de cursos teóricos ministrados por especialistas externos diminui a importância do desenvolvimento de habilidades práticas e know-how por meio do compartilhamento dentro da organização.	GenAI-5.2.2	K1	1

Q	RC	Explicação/Fundamentação	LO	K	P
40	a	a) Está correto. Os testadores evoluem para especialistas em testes assistidos por IA, refinando prompts e verificando os resultados da IA. b) Não está correto. As responsabilidades dos gerentes de teste foram atualizadas para incluir o desenvolvimento de uma estratégia de teste baseada em IA, gerenciamento de risco baseado em AI e supervisão dos processos de teste baseados em IA. Assim, seu foco continua sendo o gerenciamento de testes e não a compreensão do funcionamento interno (tecnicamente complexo) das tecnologias GenAI. c) Não está correto. Os testadores não mudam seu foco para supervisionar os processos de teste baseados em IA. Essa supervisão é de responsabilidade dos gerentes de teste. d) Não está correto. Os gerentes de teste devem equilibrar as capacidades humanas e da AI para alcançar resultados eficientes.	GenAI-5.2.3	K1	1

Apêndice - Gabarito das Questões Complementares

As respostas encontram-se nos comentários.

Respostas Adicionais Comentadas

(Q) Questão – (RC) Resposta correta – (LO) Objetivo de Aprendizagem – (K) Nível K – (P) Pontos

Q	RC	Explicação / Justificativa	LO	K	P												
A1	N/A	O processo RAG pode ser resumido da seguinte forma: 1. Divisão dos documentos em partes para facilitar a recuperação. 2. Limpar e processar partes para garantir a consistência. 3. Codificar partes em incorporações e armazená-las em um banco de dados vetorial. 4. Recuperar partes relevantes durante o processamento da consulta do usuário. 5. Gerar respostas combinando os dados recuperados com o conhecimento do LLM. Portanto, a resposta correta é: <table><tr><th colspan="2">Resposta</th></tr><tr><td>Primeiro</td><td>Divida documentos grandes em fragmentos menores</td></tr><tr><td></td><td>Limpe e processe os pedaços do documento</td></tr><tr><td></td><td>Armazene incorporações em um banco de dados vetorial</td></tr><tr><td></td><td>Recupere partes relevantes com base na semelhança semântica</td></tr><tr><td>Último</td><td>Gere uma resposta usando as partes recuperadas e o LLM</td></tr></table>	Resposta		Primeiro	Divida documentos grandes em fragmentos menores		Limpe e processe os pedaços do documento		Armazene incorporações em um banco de dados vetorial		Recupere partes relevantes com base na semelhança semântica	Último	Gere uma resposta usando as partes recuperadas e o LLM	Gen AI-4.1.2	K2	1
Resposta																	
Primeiro	Divida documentos grandes em fragmentos menores																
	Limpe e processe os pedaços do documento																
	Armazene incorporações em um banco de dados vetorial																
	Recupere partes relevantes com base na semelhança semântica																
Último	Gere uma resposta usando as partes recuperadas e o LLM																

Q	RC	Explicação / Justificativa	LO	K	P
A2	N/A	Considerando: - Cadeia de prompts - 1º item à direita: faz referência direta ao uso iterativo de LLM para tarefas de análise de teste (Seção 2.2.1). - Prompting com poucos exemplos – 2º item à direita: vincula explicitamente a casos de teste gerados por LLM com exemplos (Seção 2.2.2b). - Meta Prompting (C) – 4º item à direita: destaca a colaboração do LLM para o refinamento do prompt (Seção 2.1.2). - Prompting Zero-Shot (D) – 3º item à direita: enfatiza o raciocínio intrínseco do LLM para priorização (Seção 2.1.2). Portanto, a resposta correta é: <pre> graph LR A[Prompt Chaining] --- B[Breaking down test analysis tasks into smaller steps requiring iterative LLM interactions and human verification] C[Few-Shot Prompting] --- D[Generating Gherkin-style test cases from user stories using the LLM with pre-defined examples] E[Meta Prompting] --- F[Prioritizing test cases based on risk by leveraging the LLM's inherent reasoning without examples] G[Zero-Shot Prompting] --- H[Interacting with the LLM to iteratively refine prompts for creating test oracles from ambiguous requirements] </pre>	Gen AI-2.1.2	K2	1

Q	RC	Explicação / Justificativa	LO	K	P
A3	N/A	Considerando: • Descoberta – As atividades incluem treinar equipes de teste em conceitos de GenAI, fornecer acesso a LLMs/SLMs e experimentar casos de uso iniciais para familiarizar os testadores com a GenAI e construir confiança • Iniciação e definição de uso – Foca na identificação e priorização de casos de uso práticos para GenAI em testes de software • Utilização e iteração – Monitoramento contínuo do progresso da GenAI para testes de software e ferramentas relacionadas, bem como medição e gerenciamento da transformação para garantir benefícios sustentáveis e escalabilidade Portanto, a resposta correta é: Activities Identifying and prioritizing practical use cases Training testers on Generative AI Providing testers access to LLMs Experimenting with initial use cases to familiarize testers with Generative AI Managing the evolution of test processes following the full integration of Generative AI Key phases Discovery Initiation and Usage Definition Exploitation and Iteration	Gen AI-5.1.4	K1	1