

Guia de estudo para o exame de certificação

Generative AI Leader

Índice

Introdução	02
Conceitos básicos da IA generativa	03
Soluções de IA generativa do Google Cloud	05
Técnicas para melhorar a saída do modelo de IA generativa	08
Estratégias de negócios para uma solução de IA generativa bem-sucedida	10
Como criar seu próprio guia de estudo	11

Introdução: Generative AI Leader

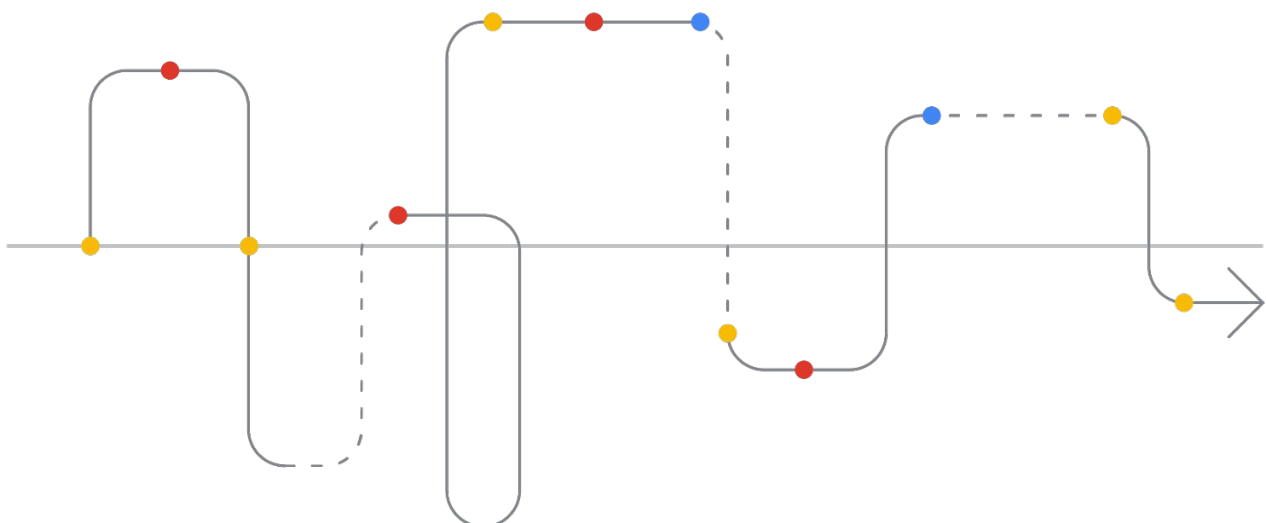
O treinamento e o exame Líder em IA Generativa do Google Cloud são para profissionais que precisam entender como a IA generativa pode ser aplicada em um contexto de negócios, principalmente ao usar ferramentas e serviços do Google Cloud. Essa pessoa consegue identificar casos de uso da IA generativa, participar de discussões com equipes técnicas e não técnicas e influenciar iniciativas de IA generativa.

Foco do exame

O exame avalia seu conhecimento em quatro áreas principais:

- Conceitos básicos da IA generativa (cerca de 30% do exame)
- Soluções de IA generativa do Google Cloud (cerca de 35% do exame)
- Técnicas para melhorar a saída do modelo de IA generativa (cerca de 20% do exame)
- Estratégias de negócios para uma solução de IA generativa bem-sucedida (cerca de 15% do exame)

Este guia de estudo é um ponto de partida, e não uma lista completa de recursos, que ajuda você a se preparar para o exame da certificação Generative AI Leader.



Conceitos básicos da IA generativa

Inteligência artificial (IA): criação de máquinas capazes de realizar tarefas que normalmente exigem inteligência humana, como aprendizado, resolução de problemas e tomada de decisão.

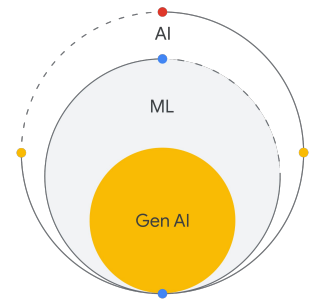
Machine learning (ML): uma subárea da IA em que as máquinas usam dados para aprender a realizar tarefas específicas.

IA generativa: um dos usos do ML voltado à criação de conteúdo novo.

Aprendizado profundo: um subconjunto de ML que usa redes neurais artificiais com múltiplas camadas para extrair padrões complexos dos dados.

Modelos de fundação: modelos avançados de ML treinados com grandes quantidades de dados não rotulados para que desenvolvam uma compreensão ampla do mundo.

Modelos de linguagem grandes (LLMs): um tipo de modelo de fundação criado para entender e gerar linguagem humana.



A **IA generativa** é um tipo de IA que pode criar novos conteúdos e ideias. Os aplicativos de IA generativa podem ser multimodais, o que permite processar e gerar simultaneamente diferentes tipos de dados, como texto, imagens e código. A IA generativa pode:



Criar

Gerar conteúdo novo



Resumir

Condensar informações em resumos concisos



Descobrir

Encontrar informações na hora certa



Automatizar

Automatizar tarefas que antes eram manuais

Modelos de fundação: modelos de IA grandes treinados em conjuntos de dados enormes, permitindo que eles sejam adaptados a muitas tarefas. Eles são a base da IA generativa.

Principais recursos dos modelos de fundação:

- **Treinados** com dados diversos
- **Flexíveis** para uma grande variedade de casos de uso
- **Adaptáveis** a domínios especializados com treinamento adicional e direcionado

Comandos: são a forma de interagir com os modelos de fundação e orientá-los. A ideia é fornecer a eles instruções ou entradas para gerar as saídas desejadas.

Peça ao Gemini



Engenharia de comando: a arte e a ciência de criar entradas eficazes, conhecidas como comandos, para modelos de IA generativa. O objetivo é maximizar o valor desses modelos e adaptar as respostas às demandas dos usuários.

Dados rotulados: informações que têm tags associadas, como nome, tipo ou número.

Dados não rotulados: informações brutas e não processadas que não receberam tags e não têm significado quando estão sozinhas, como fotos desorganizadas ou fluxos de gravações de áudio.

O ML tem três abordagens principais de aprendizagem:



Aprendizagem supervisionada: treina modelos com dados rotulados para prever saídas para novas entradas.



Aprendizagem não supervisionada: usa dados não rotulados para encontrar agrupamentos e padrões naturais.



Aprendizagem por reforço: é baseada na interação e no feedback para maximizar as recompensas e minimizar as penalidades.

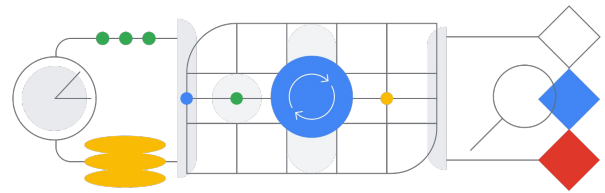
Conceitos básicos da IA generativa

Dados são informações que podem vir de várias formas: números, datas, descrições de texto e até imagens ou sons.

- ◆ **Dados estruturados:** dados organizados e fáceis de pesquisar, geralmente armazenados em bancos de dados relacionais.
- ◆ **Dados não estruturados:** informações que não têm uma estrutura predefinida e exigem técnicas de análise sofisticadas.
- ◆ **Dados de qualidade:** informações precisas, completas, consistentes e relevantes.
- ◆ **Dados acessíveis:** para treinamento de modelo, as informações precisam estar disponíveis, utilizáveis e no formato adequado.

Ciclo de vida de ML

- **Ingestão e preparação de dados:** o processo de coleta, limpeza e transformação de dados brutos em um formato utilizável para análise ou treinamento de modelos.
- **Treinamento de modelos:** o processo de criação de um modelo de ML usando dados.
- **Implantação de modelos:** o processo de disponibilização de um modelo treinado para uso.
- **Gerenciamento de modelos:** o processo de gerenciamento e manutenção dos modelos ao longo do tempo.



Cenário da IA generativa

- **Aplicativo com tecnologia de IA generativa:** a parte da IA generativa voltada ao usuário. É a camada que permite que ele interaja e use os recursos da IA.
- **Agente:** um software que aprende a atingir uma meta com base nas entradas e ferramentas disponíveis.
- **Plataforma:** essa camada oferece APIs, recursos de gerenciamento de dados e ferramentas de implantação de modelos. Ela preenche a lacuna entre modelos e agentes, além de simplificar as complexidades do gerenciamento da infraestrutura.
- **Modelo:** um algoritmo complexo treinado com enormes quantidades de dados. Ele aprende padrões e relações nos dados, o que permite gerar novos conteúdos, traduzir idiomas, responder perguntas e muito mais.
- **Infraestrutura:** essa camada fornece os principais recursos de computação necessários para a IA generativa. Isso inclui o hardware físico (como servidores, GPUs e TPUs) e o software necessários para armazenar e executar modelos de IA e dados de treinamento.



Gemini: compatível com compreensão multimodal, IA de conversação avançada, criação de conteúdo e respostas a perguntas.



Gemma: oferece aos desenvolvedores uma solução fácil de usar e personalizável para implantações locais e aplicativos especializados de IA.



Imagen: é um modelo de difusão para conversão de texto em imagem que gera imagens de alta qualidade com base em descrições textuais.



Veo: gera conteúdo de vídeo com base em descrições de texto ou imagens estáticas.

Recursos para aprender mais

[Uma introdução à IA generativa para executivos com pouco tempo](#)

[Modelos de linguagem grandes \(LLMs\) com IA do Google](#)

[Casos de uso reais da IA generativa das principais organizações do mundo](#)

Soluções de IA generativa do Google Cloud

O Google é uma empresa que coloca a IA em primeiro lugar

- As ferramentas de IA generativa estão integradas ao ecossistema do Google.
- O Google garante que você fique por dentro dos avanços mais recentes da IA.
- O Google oferece um ecossistema que prioriza a segurança e a ética.
- O Google Cloud oferece uma base de nível empresarial para você construir suas soluções.
- A abordagem aberta do Google oferece flexibilidade e opções para suas soluções de IA.

Ferramentas do Google para produtividade pessoal

O Gemini é um modelo de IA generativa do Google usado em várias soluções diferentes.



O **app do Gemini** é o chatbot de IA generativa do Google que ajuda em tarefas como escrever, resumir, traduzir e criar imagens. Com o **Gemini Advanced**, as empresas têm acesso a recursos extras e proteções de nível empresarial.



O **Gemini para Google Workspace** integra a IA generativa aos apps do Workspace que você já conhece, permitindo escrever e-mails no Gmail, gerar imagens no app Apresentações, resumir notas no Meet e muito mais.



O **Gemini para Google Cloud** é seu assistente de IA no Google Cloud. Ele pode ajudar você a escrever e depurar código, gerenciar e otimizar aplicativos de nuvem, analisar dados no BigQuery e fortalecer sua postura de segurança.

Além do Gemini, ferramentas como o NotebookLM atendem às demandas dos usuários, como a compreensão de documentos.



O **NotebookLM** permite fazer upload de arquivos e atua como um assistente de pesquisa, já que resume pontos importantes, responde a perguntas e gera ideias com base no material de origem.



A **Vertex AI** é a plataforma de ML unificada do Google Cloud. Ela permite criar, treinar e implantar aplicativos de ML. A Vertex AI dá acesso a modelos de IA generativa (como o Gemini) e permite que você os ajuste de acordo com suas necessidades e, depois, faça a implantação.

Vertex AI para Pesquisa: soluções de pesquisa e recomendações para sua empresa.

Use-a com a API Gemini no **Google AI Studio** ou no **Vertex AI Studio**.

- O Google AI Studio está disponível sem custo financeiro e é destinado à prototipagem rápida de IA.
- O Vertex AI Studio serve para criar e implantar aplicativos de IA prontos para produção em escala.

Pacote de engajamento do cliente: ferramentas para ajudar sua empresa a interagir de forma eficaz com os clientes. Elas podem ser criadas com base na [Central de atendimento como serviço \(CCaaS\)](#) do Google, uma solução de central de atendimento de nível empresarial nativa da nuvem.

- [Agentes de conversação](#): atuam como chatbots eficazes para seus clientes.
- [Agent Assist](#): para ajudar os agentes humanos da sua central de atendimento.
- [Insights de Conversação](#): receba insights sobre todas as suas comunicações com os clientes.

Gemini Enterprise: integre agentes de conversação e pesquisa personalizados aos sites ou painéis internos da sua organização, que poderão acessar e compreender dados de várias fontes internas.

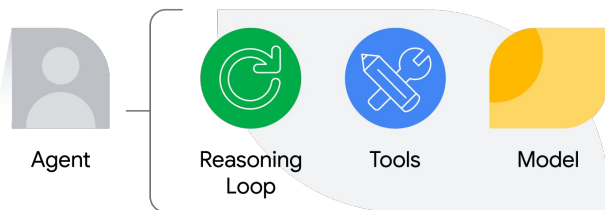
Ferramentas

- **Extensões:** permitem a conexão com serviços externos (via APIs).
- **Funções:** definem ações ou tarefas específicas.
- **Repositórios de dados:** fornecem acesso a informações.
- **Plug-ins:** adicionam novas habilidades e integrações.

Soluções de IA generativa do Google Cloud

Agente: um agente de IA generativa é um aplicativo que tenta atingir uma **meta** ao **observar** o mundo e **agir** com base nas ferramentas que tem à disposição.

Componentes do agente:



Raciocínio de repetição: processo iterativo em que o agente observa, interpreta, pensa e age, geralmente usando engenharia de comando.

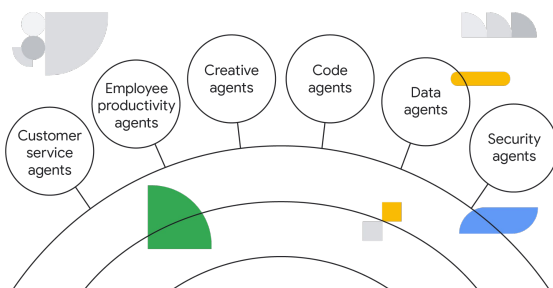
- **Ferramentas:** funcionalidades que permitem ao agente interagir com o ambiente, como acessar e processar dados ou interagir com hardware.
- **Modelo:** é o cérebro do sistema de IA. Ele é composto por vários algoritmos que aprendem padrões com base em dados e conseguem fazer previsões ou gerar conteúdo.

Tipos de agentes

- **Deterministas (tradicionais):** agentes criados com caminhos e ações predefinidos.
- **Generativos:** agentes definidos por linguagem natural usando LLMs para proporcionar uma experiência de conversa mais realista ao seu chatbot.
- **Híbridos:** combinam recursos deterministas e generativos, que os tornam muito eficientes.

Os agentes podem desempenhar várias funções diferentes.

Exemplos:



Plataforma: a base para criar e escalonar iniciativas de IA.



Como a plataforma unificada de machine learning (ML) do Google Cloud, a Vertex AI foi projetada para simplificar todo o fluxo de trabalho de ML. Ela fornece a infraestrutura, as ferramentas e os modelos pré-treinados necessários para criar, implantar e gerenciar suas soluções de ML e IA generativa.

Com as ferramentas de MLOps da Vertex AI, as equipes de IA podem colaborar melhor para monitorar e melhorar os modelos.

Modelo: a Vertex AI oferece opções para trabalhar com modelos de IA no seu projeto.

- **Model Garden:** selecione modelos prontos do Google, de terceiros ou de código aberto.
- **Model Builder:** treine e use seus próprios modelos. Personalize tudo na criação e treinamento de modelos em escala com um framework de ML. Também é possível usar o AutoML para criar e treinar modelos de forma mais fácil sem precisar de conhecimento técnico.

Infraestrutura: oferece os principais recursos de computação necessários para a IA generativa, incluindo hardware (como servidores, GPUs e TPUs) e o software necessário para treinar, armazenar e executar modelos de IA.

IA na borda: você pode executar soluções de IA na infraestrutura (dispositivos ou servidores) mais perto de onde a ação acontece.

O Google oferece ferramentas como o **Lite Runtime (LiteRT)** para ajudar desenvolvedores a implantar modelos de IA em dispositivos de borda.

Gemini Nano é o modelo de IA mais eficiente e compacto do Google, criado especialmente para rodar nos dispositivos.

Soluções de IA generativa do Google Cloud: APIs

API Speech-to-Text

- A API converte fala em texto.
- Ela também transcreve conteúdo de áudio e vídeo.

API Text-to-Speech

- Ela converte texto em fala com som natural.
- A API também cria interfaces de usuário de voz e comunicação personalizada.

API Translation

- A API Translation traduz textos, documentos, sites e arquivos de áudio e vídeo.

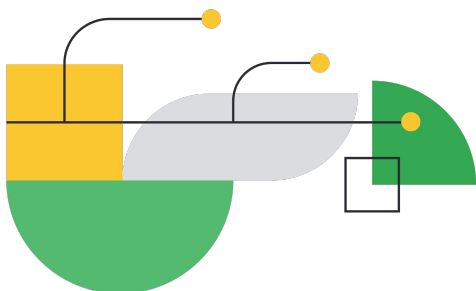
API Document Translation

- Ela traduz documentos formatados mantendo o layout original.

Como criar aplicativos com os agentes

Você pode acessar a API Gemini usando ferramentas como:

- Ferramentas para desenvolvedores do Google Cloud, como o Cloud Run Functions e o Cloud Run
- Ferramentas com pouco e nenhum código, como o Apps Script e o AppSheet



API Document AI

- A API Document AI extrai dados de vários formatos.
- Ela automatiza a captura de dados e o processamento de documentos.
- A API também pode resumir documentos.

API Cloud Vision

- A API analisa o conteúdo de imagem, incluindo tags com base nos objetos e textos detectados.
- Ela também pode identificar rostos e pontos de referência.
- A API é compatível com casos de uso como moderação de conteúdo e pesquisa visual.

API Cloud Video Intelligence

- Permite que desenvolvedores analisem conteúdo de vídeo e extraíam informações relevantes.
- Recomendação de conteúdo, pesquisa de vídeo e análise de mídia.

API Natural Language

- Gera insights com base em textos não estruturados.
- Ajuda a entender o sentimento do texto, classificar o conteúdo e extrair entidades importantes.

Recursos para aprender mais

[Agentes de IA: Parceiros para a inovação empresarial](#)

[Pesquisa com a Vertex AI](#)

[O que são aplicativos de IA?](#)

[IA acessível: comece a usar o Apps Script e o Gemini](#)

Técnicas para melhorar a saída do modelo de IA generativa

Técnicas de criação de comandos



Zero-shot: peça para o modelo concluir uma tarefa sem exemplos anteriores.



One-shot: forneça ao modelo um exemplo para aprender.



Few-shot: forneça ao modelo vários exemplos para aprender.



Papel: atribua uma persona ao modelo para influenciar o estilo, o tom e o foco dele.



Encadeamento de comandos: mantenha uma conversa com a IA.

Raciocínio de repetição: técnicas de engenharia de comando

- **Raciocínio e ação (ReAct, na sigla em inglês):** permite que o LLM raciocine e realize ações com base em uma consulta do usuário.
- **Linha de raciocínio (CoT, na sigla em inglês):** guie um LLM por um processo de resolução de problemas fornecendo exemplos com etapas intermediárias de raciocínio.
- **Metacomandos:** use comandos para orientar o modelo de IA a gerar, modificar ou interpretar outros comandos.

Simplificação dos fluxos de trabalho de comandos

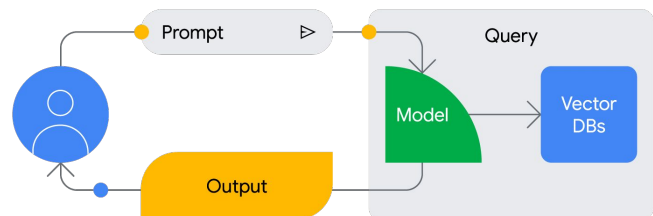
- **Reutilize comandos:** salve comandos como modelos para usar novamente.
- Use o **encadeamento de comandos:** conversas contínuas com o mesmo chatbot para manter o contexto.
- Use as **informações salvas** no Gemini: armazene informações específicas para o modelo usar de forma consistente.
- Os **Gems** são assistentes de IA personalizados no Gemini. Eles fornecem respostas personalizadas de acordo com instruções específicas e simplificam fluxos de trabalho como modelos, comandos e interações guiadas.

Orientação e refinamento de modelos

Embasamento: conecte a saída da IA a fontes de informações verificáveis.

RAG: geração aumentada por recuperação

1. **Recuperação:** o LLM recupera informações relevantes de fontes externas usando ferramentas.
2. **Aumento:** as informações recuperadas são incorporadas ao comando destinado ao LLM.
3. **Geração:** o LLM processa o comando e gera uma resposta.
4. **Iteração (opcional):** o LLM pode iterar o processo de recuperação conforme necessário.



Parâmetros de amostragem

Configurações que influenciam o comportamento do modelo de IA, permitindo resultados personalizados.

- **Contagem de tokens:** representa blocos significativos de texto (como palavras e pontuação).
- **Temperatura:** esse parâmetro controla a "criatividade" ou aleatoriedade das palavras escolhidas pelo modelo durante a geração de texto.
- **Top-P (amostragem de núcleo):** a probabilidade cumulativa dos tokens mais prováveis considerados durante a geração de texto. Essa é outra forma de controlar a aleatoriedade do resultado do modelo.
- **Configurações de segurança:** com elas, é possível remover conteúdo nocivo ou inadequado do resultado do modelo.
- **Tamanho da saída:** determina o tamanho máximo do texto gerado.

Técnicas para melhorar a saída do modelo de IA generativa

Limitações do modelo de fundação



Dependência de dados: o desempenho do modelo de fundação depende muito dos dados. Informações tendenciosas ou incompletas afetam as respostas.



Limite de conhecimento: os modelos de IA são treinados até uma data específica. Isso significa que eles podem não ter conhecimento sobre eventos posteriores a esse ponto.



Viés: os LLMs aprendem com grandes conjuntos de dados, que podem conter vieses. Até mesmo vieses sutis podem ser ampliados nas respostas do modelo.



Imparcialidade: avaliar a imparcialidade dos modelos de IA generativa é um aspecto fundamental do desenvolvimento responsável.



Alucinações: quando os modelos de IA produzem resultados que não são precisos nem baseados em informações reais.



Casos extremos: cenários raros e atípicos podem expor os pontos fracos de um modelo, levando a resultados inesperados.

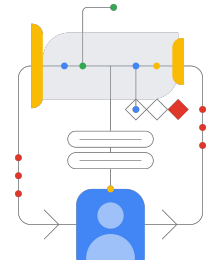
Gerenciar o modelo

O Google Cloud conta com ferramentas para gerenciar todo o ciclo dos modelos de ML. Isso inclui o seguinte:

- **Controle de versões:** monitorar as diferentes versões do modelo com o Vertex AI Model Registry.
- **Monitoramento de desempenho:** revisar as métricas do modelo para conferir o desempenho dele.
- **Monitoramento de degradação:** observar as mudanças na acurácia do modelo ao longo do tempo com o Vertex AI Model Monitoring.
- **Gerenciamento de dados:** usar o Vertex AI Feature Store para gerenciar os atributos de dados usados pelo modelo.
- **Armazenamento:** usar o Vertex AI Model Garden para armazenar e organizar os modelos em um só lugar.
- **Automatização:** usar o Vertex AI Pipelines para automatizar as tarefas de machine learning.

humans in the loop (HITL): um processo em que feedback e entrada humanos são integrados diretamente aos fluxos de trabalho de ML.

- **Moderação de conteúdo:** o HITL garante que o conteúdo gerado pelo usuário seja moderado contextualmente, capturando conteúdo prejudicial que os algoritmos podem não detectar.



Aplicações sensíveis: o HITL oferece supervisão essencial em áreas como saúde e finanças, garantindo a acurácia e reduzindo os riscos de sistemas automatizados.

- **Tomada de decisões de alto risco:** em decisões críticas, o HITL pode ajudar a garantir a acurácia e a responsabilidade com a revisão humana das respostas geradas pelo modelo de ML.
- **Revisão pré-geração:** especialistas humanos revisam e validam as saídas de ML antes da implantação, detectando erros ou vieses para evitar que os usuários sejam afetados.
- **Revisão pós-geração:** a revisão e o feedback humanos contínuos após a implantação ajudam os modelos de ML a melhorar e se adaptar aos contextos e às necessidades dos usuários.

Termos adicionais

- **Janela de contexto:** a quantidade de texto que o modelo pode considerar.
- **Ajuste de detalhes:** uma técnica usada para melhorar o desempenho de modelos pré-treinados ou de fundação com foco em tarefas ou domínios específicos.

Recursos para aprender mais

[O que é a geração aumentada por recuperação \(RAG\)?](#)

[Guia de engenharia de comando para IA](#)

[O que é human-in-the-loop \(HITL\) em IA e ML?](#)

Estratégias de negócios para uma solução de IA generativa bem-sucedida

Antes de iniciar seu projeto de IA generativa, considere estes pontos:

Necessidades:

- **Escala:** como será definida a quantidade de usuários?
- **Personalização:** qual o nível de especialização dessa IA?
- **Interação do usuário:** como os usuários vão interagir?
- **Privacidade:** os dados são sensíveis?
- **Latência:** qual é o tempo de resposta aceitável?
- **Conectividade:** qual a sua conectividade?

Recursos:

- **Pessoas:** você tem acesso a especialistas em IA?
- **Dinheiro:** qual o seu orçamento?
- **Tempo:** quais são os prazos do seu projeto?

Estratégia de IA generativa: combine uma abordagem de cima para baixo (a liderança define a visão e a estratégia) com uma abordagem de baixo para cima (os funcionários identificam aplicações práticas e dão feedback).

- **Foco estratégico:** priorize implementações de IA generativa focadas com valor comercial claro.
- **Exploração:** incentive a experimentação e a colaboração para descobrir aplicações valiosas da IA generativa.
- **IA responsável:** estabeleça diretrizes éticas e garanta o desenvolvimento seguro e responsável da IA.
- **Recursos:** invista em estratégia de dados, aproveite os recursos atuais e desenvolva talentos em IA.
- **Impacto:** avalie o impacto da IA generativa nas metas de negócios e demonstre benefícios tangíveis.
- **Melhoria contínua:** aprimore continuamente as soluções de IA generativa com base em feedback e dados.

IA responsável: como garantir que seus aplicativos de IA sejam usados de forma ética e não causem danos.

A IA responsável precisa ser considerada em todo o ciclo de vida da IA, desde a preparação dos dados e o treinamento do modelo até a implantação e o monitoramento contínuo.

IA segura: como proteger seus aplicativos de IA contra danos.

O **framework de IA segura (SAIF)** ajuda as organizações a gerenciar os riscos dos modelos de IA/ML e garantir a segurança.

A infraestrutura com segurança incorporada ao design do Google Cloud ajuda a proteger todo o ciclo de vida de IA/ML. Várias ferramentas auxiliam na segurança de dados, modelos e aplicativos.

- Identity and Access Management (IAM) para controlar o acesso aos recursos.
- Security Command Center para visibilidade da postura de segurança.
- Ferramentas de monitoramento de cargas de trabalho para ajudar a criar e manter sistemas de IA seguros.

O que considerar na hora de escolher o modelo ideal para seu caso de uso

- **Modalidade:** escolha um modelo de IA generativa com tipos de dados de entrada e saída (modalidade) alinhados às demandas do seu aplicativo, seja texto, imagens, áudio ou vídeo.
- **Janela de contexto:** talvez seja necessário equilibrar a capacidade de um modelo de IA de gerar respostas coerentes e relevantes com o aumento dos custos computacionais.
- **Desempenho:** a acurácia, velocidade e eficiência de um modelo são fatores críticos. Pense nas vantagens e desvantagens de desempenho e custo.
- **Disponibilidade e confiabilidade:** escolha um modelo que esteja sempre disponível e tenha um desempenho confiável quando carregado. Considere fatores como garantias de tempo de atividade, redundância e mecanismos de recuperação de desastres.

Para **planejar sua estratégia de IA generativa**, estabeleça uma visão clara, priorize os casos de uso, invista em recursos, gerencie mudanças, meça o valor e defenda as práticas de IA responsável.



Recursos para aprender mais

[Princípios de IA](#)

[Introdução à Vertex Explainable AI](#)

[Framework de IA segura do Google \(SAIF\)](#)

Como criar seu próprio guia de estudo com IA generativa

Siga estas etapas para praticar suas novas habilidades de IA generativa enquanto se prepara para o exame.

Etapa 1: abrir o NotebookLM

1. Navegue até o site do NotebookLM: notebooklm.google.com.
2. Se você nunca o usou antes, provavelmente vai precisar fazer login com sua Conta do Google.

Etapa 2: criar um novo notebook e fazer upload de uma fonte

1. Depois de acessar o NotebookLM, clique em **Criar novo**. Observação: pode ser um sinal de mais ou um botão com o rótulo "Novo".
2. Você vai precisar adicionar documentos de origem. Selecione **Site**.
3. Copie e cole o link deste guia de estudo em PDF no NotebookLM e selecione **Inserir**.

Observação: depois de você colar o link e fazer upload, o NotebookLM vai processar o documento. Isso pode levar alguns instantes, mas ele vai aparecer na sua lista de fontes quando estiver pronto.

Etapa 3: fazer perguntas específicas para criar as seções do seu guia de estudo

Comece a extrair informações importantes fazendo perguntas específicas ao NotebookLM. Pense nas seções comuns que você gostaria de incluir em um guia de estudo:

- **Conceitos principais:** quais são as ideias e os princípios fundamentais?
- **Serviços/produtos:** quais serviços específicos do Google Cloud são cobertos? Como eles funcionam?
- **Casos de uso:** quando e por que usar esses serviços?
- **Práticas recomendadas:** quais são as formas recomendadas de abordar determinadas tarefas?
- **Termos importantes:** quais são o vocabulário e as definições principais?

Etapa 4: refinar e organizar as informações

Não pare nas respostas iniciais. Faça perguntas complementares para esclarecer conceitos ou se aprofundar. À medida que o NotebookLM fornecer respostas, você poderá fazer o seguinte:

1. **Copiar e colar:** selecione as informações relevantes na janela de chat e cole-as em um documento separado (como um documento Google, um arquivo de texto simples ou até mesmo o NotebookLM se quiser criar um resumo nele).
2. **Reformular e resumir:** as respostas do NotebookLM são um ótimo ponto de partida. Reformule as informações com suas próprias palavras para garantir melhor compreensão e retenção. Resuma explicações mais longas em pontos concisos para uma revisão rápida.

Etapa 5: iterar e expandir seu guia de estudo

Sua primeira versão não vai ficar perfeita, mas tudo bem.

1. **Revise:** leia seu guia de estudo inicial e identifique as áreas que precisam de mais detalhes ou esclarecimentos.
2. **Inclua novas informações:** à medida que você aprender mais ou encontrar outros recursos relevantes, integre essas informações ao seu guia de estudo.
3. **Crie um podcast:** o NotebookLM pode gerar um resumo ou uma explicação em áudio das principais seções deste guia de estudo para você ouvir.
4. **Use a seção de perguntas e respostas:** faça perguntas específicas ao NotebookLM sobre conceitos ou serviços mencionados no materiais de estudo em PDF que você enviou e receba respostas concisas e contextualmente relevantes que vão esclarecer suas dúvidas.

Parabéns! Você começou a criar seu próprio guia de estudo personalizado usando o NotebookLM.