

Glossário de Termos de Inteligência Artificial

Guia para o público geral - INCRA

Orientações para o uso ético e responsável da IA



Diretoria de Gestão Estratégica - DE
Versão 1.0 - Maio/2026

A

Algoritmo

Sequência lógica e ordenada de passos ou regras que um computador segue para realizar uma tarefa ou resolver um problema. Em sistemas de Inteligência Artificial, os algoritmos são responsáveis por processar informações, identificar padrões e produzir resultados de forma automatizada.

Alto Impacto Adverso a Direitos Fundamentais

Situação em que um sistema de IA provoca consequências graves e de difícil reparação sobre direitos constitucionais das pessoas, como privacidade, igualdade, devido processo legal, liberdade de expressão ou dignidade humana. O reconhecimento desse impacto exige controles reforçados e supervisão especializada.

Alto Risco

Classificação atribuída a sistemas de IA cujo uso pode gerar consequências graves, irreversíveis ou de difícil reparação para pessoas ou grupos. Sistemas nessa categoria exigem avaliação ética rigorosa, supervisão humana ativa e medidas adicionais de transparência e prestação de contas.

Alucinação

Fenômeno em que um sistema de IA generativa produz respostas que parecem coerentes e confiáveis, mas contêm informações incorretas, desatualizadas ou completamente inventadas. O modelo não tem consciência do erro - por isso, toda resposta gerada por IA deve ser verificada antes de uso institucional.

Anonimização

Processo pelo qual dados pessoais são transformados de forma irreversível, tornando tecnicamente inviável a identificação do titular por meios razoáveis e disponíveis - padrão estabelecido pela LGPD (art. 5º, III). Ao contrário da pseudonimização, a anonimização não pode ser desfeita. Exemplos de técnicas: supressão de identificadores, generalização e randomização com perda de vínculo.

Aprendizado de Máquina (Machine Learning)

Área da Inteligência Artificial em que algoritmos aprendem padrões a partir de grandes volumes de dados e exemplos, melhorando seu desempenho em tarefas específicas sem serem explicitamente programados para cada situação. É a base da maioria dos sistemas de IA modernos.

Aprendizado Profundo (Deep Learning)

Subcampo do aprendizado de máquina que utiliza redes neurais artificiais com múltiplas camadas de processamento para aprender representações complexas dos dados. Essa abordagem permite que sistemas de IA reconheçam imagens, compreendam linguagem natural e realizem tarefas criativas com alto grau de precisão.

Auditabilidade

Capacidade de um sistema de IA de ter seu funcionamento, decisões e resultados examinados e verificados por pessoas ou instâncias de controle. Um sistema auditável permite rastrear como chegou a determinado resultado, o que é fundamental para a transparência e a responsabilização no setor público.

Automação

Uso de tecnologia para executar tarefas repetitivas ou padronizadas com mínima ou nenhuma intervenção humana. No contexto da IA, a automação pode incluir triagem de documentos, categorização de dados e geração de roteiros de trabalho. Seu uso deve ser proporcional à complexidade e criticidade da atividade realizada.

Avaliação de Impacto Ético (AIE / Framework AIE)

Processo estruturado para identificar, analisar e mitigar riscos éticos associados ao desenvolvimento ou adoção de sistemas de IA. No âmbito federal, o Ministério da Gestão e Inovação (MGI) disponibiliza o Framework AIE como ferramenta de autoavaliação para órgãos públicos que desenvolvem projetos com IA.

Agente de IA (Agentic AI)

Sistema de IA capaz de planejar e executar sequências de ações de forma autônoma para atingir objetivos definidos pelo usuário, utilizando ferramentas externas como buscas, bancos de dados ou APIs. Difere dos modelos conversacionais por tomar iniciativas e encadear múltiplas operações sem supervisão constante a cada etapa. Requer controles adicionais de governança no setor público.

B

Baixo Risco

Classificação de sistemas de IA com potencial mínimo de causar impactos negativos sobre direitos fundamentais, segurança ou bem-estar das pessoas. Em geral, enquadram-se aqui aplicações internas, experimentais ou administrativas que não envolvem decisões automatizadas sobre indivíduos.

C

Camada de Abstração

Mecanismo técnico que isola os fluxos de trabalho e os processos internos de uma organização das particularidades de um fornecedor específico de IA. Com ela, é possível trocar a plataforma de IA utilizada sem precisar reescrever todos os processos e prompts, preservando a autonomia tecnológica da instituição.

Centralidade da Decisão Humana

Princípio institucional do INCRA que estabelece que a IA deve sempre ocupar uma posição de apoio ao servidor, nunca de substituição. Todo conteúdo produzido por ferramentas de IA deve ser revisado, e a responsabilidade final sobre os resultados permanece com o servidor público. Trata-se de diretriz normativa do órgão - diferentes sistemas de IA admitem graus variados de automação (human-in-the-loop, human-on-the-loop), mas no INCRA a decisão humana é inegociável.

Ciclo de Vida de Sistemas de IA

Conjunto de etapas que um sistema de IA percorre desde sua concepção até sua descontinuação: planejamento, desenvolvimento, implementação, operação, monitoramento e encerramento. A gestão responsável desse ciclo garante que a solução permaneça segura, ética e alinhada às diretrizes institucionais ao longo do tempo.

Confidencialidade

Princípio que impõe a proteção de informações sensíveis contra acessos não autorizados. No uso de IA, a confidencialidade exige que dados institucionais, pessoais ou sigilosos não sejam inseridos em plataformas públicas ou externas à infraestrutura controlada pela organização.

D

Dados Anonimizados

Dados que passaram por processo técnico de eliminação irreversível de qualquer elemento capaz de identificar seu titular, tornando impossível, com meios razoáveis, associá-los a uma pessoa específica. Dados nessa condição não são considerados dados pessoais para fins da LGPD.

Dados Pessoais

Qualquer informação relacionada a uma pessoa natural identificada ou identificável, como nome, CPF, endereço, email-pessoal, localização geográfica ou registros de comportamento. A Lei Geral de Proteção de Dados (LGPD) estabelece regras específicas para o tratamento desse tipo de dado.

Dados Pessoais Sensíveis

Categoria especial de dados pessoais que exige proteção reforçada por envolver aspectos íntimos ou que podem gerar discriminação. Incluem: origem racial ou étnica, convicção religiosa, opinião política, filiação sindical, saúde, vida sexual, dados genéticos e biométricos.

Data Poisoning (Envenenamento de Dados)

Ataque à integridade de um sistema de IA que consiste em inserir dados maliciosos, corrompidos ou deliberadamente distorcidos no conjunto de treinamento do modelo. O objetivo é comprometer o comportamento do sistema - fazendo-o produzir erros específicos, ignorar categorias de informação ou agir de forma prejudicial. É uma das principais ameaças à confiabilidade de modelos de IA em ambientes críticos.

Deepfake

Conteúdo audiovisual - vídeo, áudio ou imagem - criado ou manipulado por IA generativa para simular de forma realista pessoas reais em situações que nunca ocorreram. Frequentemente utilizado para fraude, desinformação e danos à reputação, exige atenção especial na verificação da autenticidade de conteúdos recebidos.

Dependência Tecnológica (Vendor Lock-in)

Situação em que uma organização torna-se excessivamente dependente de um único fornecedor de tecnologia, dificultando ou inviabilizando a migração para outras soluções. No contexto de IA, consolida-se quando fluxos de trabalho, prompts e rotinas são otimizados exclusivamente para o comportamento de uma plataforma específica. A adoção de camadas de abstração é a principal forma de mitigação.

Direitos Fundamentais

Direitos e garantias individuais e coletivos assegurados pela Constituição Federal e por tratados internacionais de direitos humanos, como liberdade de expressão, privacidade, não discriminação, devido processo legal e dignidade humana. O respeito a esses direitos é condição inegociável no desenvolvimento e uso de sistemas de IA no setor público.

E

Estratégia Brasileira de Inteligência Artificial (EBIA)

Política nacional que orienta o desenvolvimento e a adoção ética, responsável e confiável da IA no Brasil, com foco especial no setor público. A EBIA estabelece diretrizes para garantir que a tecnologia seja utilizada em benefício do cidadão, respeitando valores democráticos, a legislação vigente e o interesse público.

Embeddings

Representação numérica de palavras, frases ou documentos em um espaço vetorial de alta dimensão, que captura relações de significado entre conceitos. Permitem que modelos de IA compreendam proximidade semântica entre termos - por exemplo, que 'servidor público' e 'funcionário federal' são conceitualmente próximos. São a base técnica de sistemas de busca semântica, recuperação de documentos e modelos de linguagem.

Explicabilidade

Capacidade de um sistema de IA de apresentar, de forma compreensível, os critérios e o raciocínio que levaram a determinado resultado ou decisão. A explicabilidade é essencial para garantir transparência, permitir auditorias e possibilitar a contestação de decisões automatizadas por parte dos cidadãos.

F

Fine-tuning (Ajuste Fino)

Técnica de especialização de um modelo de IA pré-treinado por meio de treinamento adicional com dados específicos de um domínio ou tarefa. Permite que modelos genéricos passem a responder com maior precisão e aderência ao contexto de uma área de conhecimento ou organização específica.

G

Governança de IA

Conjunto de políticas, processos, responsabilidades e mecanismos de controle que orientam o desenvolvimento, a adoção e o uso de sistemas de IA dentro de uma organização ou governo. Uma boa governança garante que a tecnologia seja utilizada de forma ética, segura, transparente e alinhada ao interesse público.

Guardrails

Mecanismos de controle implementados em sistemas de IA para delimitar o comportamento do modelo - impedindo respostas inadequadas, conteúdo prejudicial, vazamento de informações sensíveis ou saídas fora do escopo permitido. Podem ser aplicados por meio de filtros, instruções de sistema, verificações automatizadas ou camadas de validação. São componentes essenciais da governança de IA em ambientes institucionais.

Grande Modelo de Linguagem (LLM)

Tipo de rede neural artificial treinada em volumes massivos de texto - livros, sites, artigos e outros documentos - para compreender e gerar linguagem humana. Famílias de modelos como GPT (OpenAI), Gemini (Google) e Claude (Anthropic) são construídas sobre essa arquitetura. São a base tecnológica da maioria dos assistentes de IA generativa disponíveis hoje.

I

IA Corporativa / Institucional

Soluções de Inteligência Artificial desenvolvidas internamente, homologadas pela área de TI ou contratadas com cláusulas específicas de privacidade e proteção de dados. Esses ambientes garantem que as informações processadas permaneçam dentro do perímetro de segurança da organização, sem uso para fins de terceiros.

IA Generativa

Categoria de sistemas de Inteligência Artificial capaz de criar conteúdos - textos, imagens, áudios ou vídeos - com base em padrões aprendidos durante o treinamento. Ferramentas como ChatGPT, Gemini e Copilot são exemplos de IA generativa. Seu uso no setor público exige atenção especial à verificação dos resultados e à proteção de dados.

IA Pública

Ferramentas de Inteligência Artificial acessadas em ambiente externo à infraestrutura institucional e sem garantias contratuais específicas de proteção de dados. O risco não depende apenas do custo de acesso, mas de fatores como termos contratuais, retenção de dados, possibilidade de uso para treinamento posterior, jurisdição aplicável e ausência de isolamento do ambiente. Seu uso deve se restringir a tarefas genéricas que não envolvam informações institucionais, sigilosas ou pessoais.

Inferência

Processo pelo qual um modelo de IA, após treinado, aplica o conhecimento adquirido para gerar respostas ou previsões a partir de novas entradas. É a etapa de uso do modelo em produção - distinta do treinamento, que ocorre uma única vez. A velocidade, o custo e o consumo de energia da inferência determinam a viabilidade de implantação de sistemas de IA em escala.

Impessoalidade

Princípio da administração pública que determina que decisões - incluindo as apoiadas por IA - devem ser objetivas, imparciais e fundamentadas em critérios legais e técnicos, sem favoritismos ou discriminações. No uso de IA, a impessoalidade exige que resultados gerados automaticamente sejam verificados quanto a vieses.

Inteligência Artificial (IA)

Campo da ciência da computação dedicado ao desenvolvimento de sistemas capazes de realizar tarefas que, historicamente, exigiam inteligência humana - como reconhecer padrões, compreender linguagem, tomar decisões e gerar conteúdo. No contexto institucional, a IA deve ser compreendida como ferramenta de apoio ao trabalho humano, não como substituta do julgamento técnico ou da responsabilidade administrativa.

J

Janela de Contexto (Context Window)

Quantidade máxima de texto que um modelo de linguagem pode processar de uma vez - incluindo a conversa anterior, instruções do sistema e o conteúdo a ser analisado. Quando o volume de informação ultrapassa esse limite, o modelo perde acesso às partes mais antigas da conversa, o que pode comprometer a coerência e a qualidade das respostas em interações longas.

L

Legalidade

Princípio que determina que o uso de IA deve estar em conformidade com todas as leis e regulamentos vigentes, incluindo a Lei Geral de Proteção de Dados (LGPD), normas de direitos autorais, a Constituição Federal e demais legislações aplicáveis. Nenhuma conveniência tecnológica justifica o descumprimento do ordenamento jurídico.

LGPD (Lei Geral de Proteção de Dados)

Lei nº 13.709/2018 que estabelece as regras para o tratamento de dados pessoais no Brasil - tanto por organizações públicas quanto privadas. Define os direitos dos titulares de dados, as bases legais para o tratamento, as obrigações dos controladores e as sanções em caso de descumprimento. O uso de IA deve observar rigorosamente seus preceitos.

M

Médio Risco

Classificação de sistemas de IA que interagem com o público ou apoiam decisões com impacto moderado sobre indivíduos ou grupos. Incluem aplicações que processam dados pessoais não sensíveis ou fornecem recomendações para decisões humanas. Exigem controles proporcionais à natureza e ao alcance dos dados tratados.

Minimização de Dados

Princípio que orienta o uso apenas das informações estritamente necessárias para alcançar a finalidade pretendida. Ao interagir com sistemas de IA, o servidor deve filtrar previamente os dados fornecidos, excluindo qualquer informação irrelevante ou excessiva, especialmente dados pessoais ou sigilosos.

Multimodalidade

Capacidade de sistemas de IA de processar e gerar conteúdos em múltiplos formatos de forma integrada - como texto, imagem, áudio e vídeo. Modelos multimodais podem, por exemplo, interpretar uma imagem e responder perguntas sobre ela, ou gerar um texto a partir de um áudio. Ampliam as possibilidades de uso institucional e também os riscos associados à criação de conteúdo sintético.

Modelo de Linguagem

Sistema computacional treinado para compreender e produzir texto em linguagem natural. Aprende padrões estatísticos em grandes volumes de texto e os utiliza para gerar respostas, resumos, traduções e outros conteúdos. O comportamento do modelo depende dos dados com que foi treinado e das instruções fornecidas pelo usuário.

P

Persona

Representação fictícia de um perfil de usuário, construída com base em dados reais, usada para orientar o desenvolvimento e a configuração de ferramentas de IA generativa. Permite que o sistema adapte tom, linguagem e foco ao público-alvo para o qual está sendo preparado.

Plataformas Externas (APIs de Terceiros)

Soluções de IA generativa acessadas por meio de interfaces de programação (APIs) fornecidas por terceiros, como ChatGPT, DeepSeek, Gemini ou Claude. Por não integrarem a infraestrutura interna da organização, seu uso exige avaliação criteriosa quanto à sensibilidade das informações a serem compartilhadas.

Prompt / Prompting Qualificado

Instrução ou comando textual fornecido pelo usuário a um modelo de IA generativa para orientar a geração de uma resposta. O prompting qualificado refere-se à prática de formular comandos precisos, contextualizados e detalhados, de forma iterativa, confrontando os resultados com fontes confiáveis e critérios institucionais aplicáveis.

Prompt Injection

Técnica de ataque a sistemas de IA em que instruções maliciosas são inseridas nas entradas do modelo - via documentos, campos de texto ou conteúdos externos - para manipular seu comportamento. O modelo, ao processar o conteúdo comprometido, pode ignorar suas instruções originais, vaziar informações sensíveis ou executar ações não autorizadas. Risco relevante em sistemas que processam documentos de terceiros.

Propriedade Intelectual

Conjunto de direitos legais que protegem criações do intelecto humano - textos, imagens, códigos, músicas e outros conteúdos - contra uso não autorizado. No contexto da IA, é necessário atenção ao fato de que conteúdos gerados por modelos podem reproduzir obras protegidas, gerando risco de violação de direitos autorais.

Pseudonimização

Técnica de proteção de dados que reduz a identificabilidade das informações, mas não a elimina completamente - ou seja, ainda é possível reidentificar o titular mediante uso de dados adicionais controlados. Difere da anonimização por ser reversível. Exemplos: mascaramento parcial de CPF, tokenização e substituição por códigos ou chaves.

R

Racismo Algorítmico

Reprodução de discriminação racial por sistemas algorítmicos ou de IA como consequência de desequilíbrios, sub-representações ou preconceitos presentes nos dados de treinamento. Manifesta-se em scoring de crédito, reconhecimento facial, classificação automatizada, policiamento preditivo e outros contextos - não se limita à IA generativa. Resulta em outputs sistematicamente desfavoráveis a determinados grupos étnicos ou raciais, perpetuando desigualdades históricas.

RAG (Retrieval-Augmented Generation)

Abordagem técnica que combina um modelo de linguagem com uma base de conhecimento externa - como documentos, bases de dados ou sistemas corporativos. Em vez de depender exclusivamente do que foi aprendido durante o treinamento, o modelo consulta fontes atualizadas antes de gerar uma resposta. Reduz alucinações, aumenta a precisão em domínios específicos e permite trabalhar com informações não incluídas no treinamento original.

Rastreabilidade

Capacidade de reconstituir o histórico de uso de uma ferramenta de IA - quais prompts foram utilizados, que parâmetros foram configurados e que conteúdos foram gerados. A rastreabilidade é fundamental para garantir a transparência, a auditabilidade e a integridade dos atos administrativos realizados com apoio de IA.

Red Teaming

Prática de segurança em que uma equipe testa os limites e vulnerabilidades de um sistema de IA de forma intencional e adversarial, simulando ataques, manipulações e usos indevidos. Serve para identificar falhas de segurança, comportamentos inesperados e pontos de exploração antes da implantação em produção. É recomendado como etapa de governança em sistemas de alto risco.

Relatório de Impacto à Proteção de Dados (RIPD)

Documento exigido pela LGPD que descreve os processos de tratamento de dados pessoais com potencial de risco às liberdades e direitos fundamentais dos titulares, acompanhado das medidas adotadas para mitigar esses riscos. É especialmente relevante em projetos que envolvam processamento de dados por sistemas de IA.

Responsabilidade pelo Resultado

Princípio que estabelece que o uso de IA não isenta o servidor da responsabilidade pelos documentos, análises ou manifestações institucionais produzidos com seu auxílio. Erros, omissões e inconsistências não podem ser atribuídos à ferramenta. A validação deve abranger forma e conteúdo técnico, e a responsabilidade administrativa permanece integral.

Revisão Humana

Processo de verificação e validação por parte de um profissional qualificado dos conteúdos gerados por sistemas de IA, assegurando precisão, adequação ética e conformidade com o contexto institucional. A revisão humana não é etapa opcional - é condição indispensável para o uso responsável da IA no setor público.

Risco Excessivo

Nível máximo de risco em sistemas de IA, caracterizado por potencial de dano inaceitável à sociedade, com violação sistemática de princípios legais ou éticos fundamentais. Incluem aplicações que manipulam comportamento humano de forma prejudicial ou operam com autonomia total em contextos críticos sem supervisão adequada. Sistemas nessa categoria não devem ser adotados.

Riscos no Uso de IA

Possibilidades de consequências negativas associadas ao uso de sistemas de IA, classificadas conforme sua gravidade: baixo risco (impacto mínimo), médio risco (impacto moderado sobre indivíduos), alto risco (consequências graves ou irreversíveis) e risco excessivo (dano inaceitável à sociedade). A identificação do nível de risco orienta a intensidade dos controles a serem aplicados.

S

Segurança da Informação

Conjunto de práticas, controles e políticas destinados a proteger dados e sistemas contra acessos não autorizados, vazamentos, alterações indevidas e indisponibilidade. No uso de IA, a segurança da informação exige atenção ao ambiente onde a ferramenta opera, ao tipo de dado inserido e à procedência e confiabilidade da plataforma utilizada.

Serviços Essenciais

Serviços cuja interrupção compromete gravemente a saúde, a segurança, o bem-estar ou o funcionamento da sociedade - como saúde, educação, segurança pública, transporte, energia e saneamento. O uso de IA em serviços essenciais exige controles especialmente rigorosos e supervisão humana constante.

Supervisão Humana Significativa

Capacidade efetiva, informada e qualificada de supervisores humanos de compreender o funcionamento de um sistema de IA, monitorar suas decisões em tempo real, intervir quando necessário e reverter resultados antes que causem impactos adversos. Vai além da mera presença humana no processo - exige preparo e autoridade para agir.

T

Token

Unidade básica de processamento de modelos de linguagem - geralmente correspondente a uma palavra, parte de palavra ou símbolo de pontuação. Modelos de IA convertem texto em sequências de tokens antes de processá-lo. O número total de tokens suportados define a janela de contexto do modelo e determina quanto texto pode ser analisado de uma vez. Também é a unidade de cobrança da maioria das APIs de IA.

Transparência

Princípio que determina clareza sobre o uso de IA nos processos institucionais. Opera em três níveis: (1) transparência institucional - informar internamente quais sistemas de IA são utilizados e como; (2) rastreabilidade interna - registrar prompts, parâmetros e outputs para auditoria; (3) obrigação de informar ao cidadão - aplicável quando houver impacto relevante sobre decisões que afetam direitos. Nem todo uso auxiliar de IA exige “*disclosure*” formal ao público.

Transparência Algorítmica

Princípio que determina que o funcionamento de sistemas algorítmicos seja compreensível e verificável por usuários, especialistas e órgãos reguladores. Não significa revelar o código-fonte de sistemas proprietários, mas garantir que os critérios de decisão possam ser explicados, auditados e compreendidos pelos afetados.

U

Uso Proporcional / Adequação ao Interesse Público

Princípio que orienta a utilização de IA de forma compatível com a complexidade, a criticidade e o impacto da tarefa realizada. Nem toda atividade se beneficia do uso de IA - e seu emprego indiscriminado pode gerar mais riscos que benefícios. Quanto maior o risco envolvido, maior deve ser o controle humano sobre o processo.

V

Viés Algorítmico

Tendência sistemática de um sistema de IA a produzir resultados distorcidos ou discriminatórios, decorrente de desequilíbrios nos dados de treinamento, escolhas de design ou decisões humanas ao longo do desenvolvimento do modelo. Pode perpetuar desigualdades sociais e comprometer a imparcialidade de processos automatizados.

Viés de Automação

Tendência humana de aceitar e confiar acriticamente nos resultados produzidos por sistemas automatizados, ignorando o próprio julgamento e o bom senso. No uso de IA, manifesta-se quando o usuário adota a resposta gerada sem verificação, transferindo inconscientemente a responsabilidade pela decisão à máquina.

Viés de Exclusão

Falha em modelos de IA causada pela sub-representação ou ausência de determinados grupos sociais nos dados de treinamento. Resulta em sistemas que apresentam desempenho inferior ou produzem interpretações equivocadas quando aplicados a esses grupos, gerando discriminação indireta.

Viés de Resultado

Resultado distorcido gerado por um sistema de IA como consequência de tendências presentes nos dados de treinamento ou nos próprios algoritmos do modelo. Pode levar a recomendações ou decisões injustas mesmo quando os dados de entrada parecem neutros.

Viés do Modelo

Conjunto de predisposições embutidas em um modelo de IA durante seu desenvolvimento - oriundas das escolhas de design, dos dados utilizados e das decisões humanas ao longo do processo de treinamento. Esses vieses afetam diretamente a qualidade, a equidade e a confiabilidade dos resultados produzidos.

A elaboração das imagens/ilustrações deste documento contou com o auxílio da ferramenta de inteligência artificial Claude (Opus 4.7), com revisão do servidor Wesley Teixeira Rodrigues de Menezes, 1467990.

Referências

BRASIL. Ministério da Gestão e da Inovação em Serviços Públicos. Secretaria de Governo Digital. **Glossário de Termos Relacionados à IA**. Brasília: MGI, 2025. Disponível em: <https://www.gov.br/governodigital/pt-br/infraestrutura-nacional-de-dados/inteligencia-artificial-1/publicacoes/glossario-de-termos-relacionados-a-ia>. Acesso em: maio 2026.

BRASIL. Ministério da Gestão e da Inovação em Serviços Públicos. Secretaria de Governo Digital. **Framework para a Autoavaliação de Impacto Ético em Inteligência Artificial no Setor Público Federal (AIE)**. Brasília: MGI, 2024. Disponível em: <https://www.gov.br/governodigital/pt-br/infraestrutura-nacional-de-dados/inteligencia-artificial-1/publicacoes/framework-para-a-autoavaliacao-de-impacto-etico-em-inteligencia-artificial-no-setor-publico-federal>. Acesso em: maio 2026.

INCRA - Instituto Nacional de Colonização e Reforma Agrária. **Diretrizes e Análise de Riscos no Uso de Inteligência Artificial**. Brasília: INCRA, 2026. Versão 1.0 – Maio/2026.

INCRA - Instituto Nacional de Colonização e Reforma Agrária. **Guia de Boas Práticas para o Uso de Inteligência Artificial**. Brasília: INCRA, 2026. Versão 1.0 – Maio/2026.

UNESCO. **Recomendação sobre a Ética da Inteligência Artificial**. Brasília: UNESCO, 2022. Disponível em: https://unesdoc.unesco.org/ark:/48223/pf0000381137_por. Acesso em: maio 2026.