

Glosario de IA para Directivos — AI4PDG

61 términos esenciales para la alta dirección Versión 2.0 · Junio 2026 ·

emovere.agency/IA4PDG

A

Agente de IA (AI Agent) Sistema de IA capaz de planificar y ejecutar tareas de forma autónoma, tomando decisiones intermedias sin intervención humana en cada paso. A diferencia de un chatbot, un agente puede usar herramientas (buscar en internet, ejecutar código, enviar emails) para completar objetivos complejos. *Ejemplo directivo: Manus es un agente que puede investigar un mercado, redactar un informe y enviarlo por email sin que el directivo intervenga en cada paso.* Ver también: LLM · Human-in-the-Loop · Manus

AI Act (Reglamento Europeo de IA) ⚠ LEGAL URGENTE — Regulación de la Unión Europea que clasifica los sistemas de IA según su nivel de riesgo (inaceptable, alto, limitado, mínimo) y establece obligaciones para empresas que los desarrollen o usen. Las prohibiciones absolutas están vigentes desde febrero 2025; las obligaciones para sistemas de alto riesgo entran en vigor en agosto 2026. *Ejemplo directivo: Un sistema que evalúa candidatos a empleo automáticamente es “alto riesgo” bajo el AI Act y requiere supervisión humana obligatoria, auditorías y documentación técnica.* Ver también: Sistema de IA de Alto Riesgo · Explicabilidad · RGPD

AI Red Teaming Práctica de simular ataques adversariales contra un sistema de IA para identificar vulnerabilidades, sesgos o comportamientos no deseados antes del despliegue. Equipos especializados intentan “romper” el sistema de forma controlada. *Ejemplo directivo: Antes de desplegar un chatbot de atención al cliente, el equipo de red teaming intenta hacerle decir cosas inapropiadas o revelar datos confidenciales.* Ver también: Prompt Injection · AI Safety · Responsible AI

AI Safety (Seguridad en IA) Campo de investigación y práctica centrado en garantizar que los sistemas de IA se comporten de forma segura, predecible y alineada con los valores humanos. Incluye tanto la seguridad técnica (evitar fallos) como la seguridad ética (evitar daños). *Ejemplo directivo: Un sistema de IA que gestiona inventario debe tener mecanismos de seguridad que impidan ordenar cantidades absurdas aunque el modelo lo “decida”.* Ver también: Responsible AI · Human-in-the-Loop · Alignment

Alignment (Alineación) Capacidad de un sistema de IA para actuar de acuerdo con los valores, intenciones y objetivos del usuario o la organización. Un modelo “desalineado” puede optimizar para un objetivo técnico mientras produce resultados no deseados. *Ejemplo directivo: Un modelo optimizado para “maximizar engagement” puede recomendar contenido polarizador*

aunque el objetivo real sea informar correctamente. Ver también: [AI Safety](#) · [Responsible AI](#) · [Human-in-the-Loop](#)

Alucinación (Hallucination) Fenómeno por el que un modelo de IA genera información falsa pero presentada con aparente confianza. No es un “error” en el sentido tradicional: el modelo no sabe que está equivocado. *Ejemplo directivo: ChatGPT puede citar un artículo académico con título, autor y año perfectamente plausibles que simplemente no existe.* Ver también: [Grounding](#) · [RAG](#) · [Output](#)

API (Application Programming Interface) Interfaz que permite a dos sistemas de software comunicarse entre sí. En el contexto de IA, las APIs permiten integrar modelos de lenguaje en aplicaciones empresariales sin necesidad de desarrollar la IA desde cero. *Ejemplo directivo: La API de OpenAI permite a una empresa integrar ChatGPT en su CRM para generar respuestas automáticas a clientes.* Ver también: [Automatización](#) · [Make](#) · [Zapier](#)

Automatización (Automation) Ejecución de tareas repetitivas por sistemas de software sin intervención humana continua. La IA generativa amplía la automatización a tareas que antes requerían juicio humano (redacción, análisis, clasificación). *Ejemplo directivo: Automatizar el triaje de emails de clientes para clasificarlos por urgencia y asignarlos al departamento correcto.* Ver también: [API](#) · [Make](#) · [SOP](#)

B

Benchmark Prueba estandarizada para medir y comparar el rendimiento de modelos de IA en tareas específicas. Permite evaluar objetivamente qué modelo es mejor para un caso de uso concreto. *Ejemplo directivo: Antes de elegir entre ChatGPT y Claude para análisis financiero, se puede ejecutar un benchmark con 50 casos reales de la empresa.* Ver también: [Fine-tuning](#) · [Output](#)

Brecha de seguridad (Data Breach) ⚠️ LEGAL URGENTE — Acceso no autorizado, pérdida o divulgación de datos personales. En el contexto de IA, puede ocurrir cuando se introducen datos confidenciales en herramientas de IA que los almacenan o usan para entrenamiento. *Ejemplo directivo: Subir un contrato confidencial a ChatGPT sin verificar la política de privacidad puede constituir una brecha de seguridad notificable a la AEPD en 72 horas.* Ver también: [RGPD](#) · [DPA](#) · [Shadow AI](#)

C

Chain-of-Thought (CoT / Cadena de Razonamiento) Técnica de prompting que instruye al modelo a “pensar en voz alta” paso a paso antes de dar una respuesta final. Mejora significativamente la precisión en tareas de razonamiento complejo. *Ejemplo directivo: Añadir*

“Razona paso a paso antes de responder” a un prompt de análisis financiero reduce los errores de cálculo en un 40-60%. Ver también: Prompt · Few-Shot Prompting · Zero-Shot Prompting

Chatbot Programa que simula conversación humana mediante texto o voz. Los chatbots modernos basados en LLMs son mucho más capaces que los chatbots de árbol de decisión tradicionales, pero siguen siendo herramientas de conversación, no agentes autónomos.

Ejemplo directivo: Un chatbot de RRHH puede responder preguntas frecuentes sobre vacaciones, pero no puede aprobar una solicitud de baja. Ver también: LLM · Agente de IA · Output

Claude Modelo de lenguaje grande desarrollado por Anthropic, diseñado con énfasis en seguridad y alineación. Destacado en análisis de documentos largos (hasta 200.000 tokens de contexto), redacción matizada y razonamiento ético. *Ejemplo directivo: Claude es especialmente útil para analizar contratos largos, informes anuales completos o transcripciones de reuniones extensas. Ver también: LLM · Responsable AI · Contexto*

Contexto (Context Window) Cantidad máxima de texto (medida en tokens) que un modelo puede procesar en una sola interacción. Determina cuánta información puede “recordar” el modelo durante una conversación. *Ejemplo directivo: GPT-4 tiene un contexto de ~128.000 tokens (~96.000 palabras), suficiente para analizar un informe anual completo de una empresa mediana. Ver también: LLM · Token · RAG*

Copilot (Microsoft) Suite de herramientas de IA de Microsoft integradas en el ecosistema Office 365 (Word, Excel, PowerPoint, Teams, Outlook). Permite usar IA directamente en las aplicaciones de productividad empresarial más comunes. *Ejemplo directivo: Copilot en Excel puede analizar una hoja de cálculo con 10.000 filas y generar un resumen ejecutivo en lenguaje natural. Ver también: LLM · Prompt · Fine-tuning*

Cumplimiento normativo (Compliance IA) ⚠️ LEGAL URGENTE — Conjunto de obligaciones legales y regulatorias que las organizaciones deben cumplir al desarrollar o usar sistemas de IA. Incluye el AI Act europeo, el RGPD, normativas sectoriales (DORA para finanzas, MDR para sanidad) y políticas internas. *Ejemplo directivo: Una empresa que usa IA para scoring de crédito debe documentar el modelo, garantizar la explicabilidad de las decisiones y permitir que los afectados soliciten revisión humana. Ver también: AI Act · RGPD · DPA*

D

Demo / Piloto / Producción Las tres fases de madurez de un sistema de IA: Demo (prueba de concepto, datos ficticios), Piloto (prueba con datos reales en entorno controlado) y Producción (despliegue completo con usuarios reales). Cada fase tiene requisitos de gobernanza distintos. *Ejemplo directivo: Un chatbot de atención al cliente debe pasar por piloto con 100 clientes antes de desplegarse a toda la base. Ver también: Human-in-the-Loop · SOP*

Digital Twin (Gemelo Digital) Representación virtual de un proceso, producto o sistema real que se actualiza con datos en tiempo real. La IA puede usar gemelos digitales para simular escenarios y optimizar decisiones sin riesgo. *Ejemplo directivo: Un gemelo digital de la cadena de suministro permite simular el impacto de un proveedor fallido antes de que ocurra.* Ver también: Automatización · Inferencia

DPA (Data Processing Agreement / Acuerdo de Tratamiento de Datos) ⚠️ LEGAL

URGENTE — Contrato obligatorio bajo el RGPD entre una empresa y cualquier proveedor que procese datos personales en su nombre. Es obligatorio antes de usar herramientas de IA con datos de clientes o empleados. *Ejemplo directivo: Antes de usar ChatGPT Enterprise con datos de clientes, la empresa debe firmar un DPA con OpenAI que especifique cómo se usan y protegen esos datos.* Ver también: RGPD · Brecha de seguridad · Cumplimiento normativo

E

Embeddings (Vectores semánticos) Representación matemática del significado de un texto como un vector numérico. Permiten a los sistemas de IA comparar la “similitud semántica” entre textos, lo que es la base de la búsqueda semántica y los sistemas RAG. *Ejemplo directivo: Los embeddings permiten buscar “problemas con proveedores” en un repositorio de emails y encontrar mensajes que hablan de “retrasos en entregas” aunque no usen esas palabras exactas.* Ver también: RAG · Base de datos vectorial · NotebookLM

Explicabilidad (XAI — Explainable AI) ⚠️ LEGAL URGENTE — Capacidad de un sistema de IA para explicar sus decisiones en términos comprensibles para humanos. El AI Act exige explicabilidad para sistemas de alto riesgo; el RGPD garantiza el derecho a explicación en decisiones automatizadas. *Ejemplo directivo: Si un sistema de IA rechaza una solicitud de crédito, el cliente tiene derecho a saber qué factores influyeron en la decisión.* Ver también: AI Act · RGPD · Sistema de IA de Alto Riesgo

F

Few-Shot Prompting Técnica de prompting que incluye 2-5 ejemplos del resultado esperado dentro del prompt para guiar al modelo. Más efectivo que las instrucciones abstractas para tareas con formato específico. *Ejemplo directivo: Para generar informes ejecutivos consistentes, incluir 2-3 ejemplos del formato deseado en el prompt reduce significativamente la variabilidad del output.* Ver también: Prompt · Zero-Shot Prompting · Chain-of-Thought

Fine-tuning (Ajuste fino) Proceso de entrenar adicionalmente un modelo de IA preentrenado con datos específicos de una organización para mejorar su rendimiento en tareas concretas. Más costoso que el prompting pero produce resultados más consistentes. *Ejemplo directivo: Una empresa de seguros puede hacer fine-tuning de un LLM con sus propias pólizas y*

resoluciones para que el modelo responda con precisión sobre sus productos específicos. Ver también: LLM · Embeddings · RAG

G

Gemini (Google) Familia de modelos de lenguaje grande de Google DeepMind, integrados en el ecosistema Google (Workspace, Search, Cloud). Destacado en tareas multimodales (texto, imagen, audio, vídeo) y en integración con servicios de Google. *Ejemplo directivo: Gemini en Google Workspace puede analizar una hoja de cálculo de Google Sheets y generar visualizaciones automáticamente.* Ver también: LLM · Multimodal · Copilot

Grounding (Fundamentación en fuentes) Técnica para reducir alucinaciones vinculando las respuestas del modelo a fuentes verificables específicas. Un modelo “grounded” cita sus fuentes y solo afirma lo que puede respaldar con documentación real. *Ejemplo directivo: NotebookLM usa grounding: solo responde con información de los documentos que el usuario ha subido, citando la fuente exacta.* Ver también: RAG · Alucinación · NotebookLM

H

Human-in-the-Loop (HITL) Modelo de supervisión en el que un humano revisa y aprueba las decisiones de la IA antes de que tengan efecto real. Obligatorio para decisiones de alto impacto según el AI Act. *Ejemplo directivo: Un sistema de IA puede generar 100 respuestas a clientes, pero un agente humano las revisa y aprueba antes de enviarlas.* Ver también: AI Act · Agente de IA · Automatización

I

IA Generativa (Generative AI) Categoría de sistemas de IA capaces de generar contenido nuevo (texto, imágenes, audio, código, vídeo) a partir de instrucciones en lenguaje natural. Se diferencia de la IA “discriminativa” que solo clasifica o predice. *Ejemplo directivo: ChatGPT, Claude, Midjourney y Suno son ejemplos de IA generativa; un sistema de detección de fraude es IA discriminativa.* Ver también: LLM · Multimodal · Prompt

Inferencia (Inference) Proceso de usar un modelo de IA entrenado para generar predicciones o respuestas a partir de nuevas entradas. Es lo que ocurre cuando el usuario envía un prompt y recibe una respuesta. *Ejemplo directivo: Cada vez que un directivo hace una pregunta a ChatGPT, el sistema ejecuta una inferencia: aplica el modelo entrenado a ese input específico.* Ver también: LLM · Token · Output

Iteración de prompt (Prompt Iteration) Proceso de refinar y mejorar un prompt de forma sistemática para obtener mejores resultados. Es una habilidad fundamental en el uso efectivo de herramientas de IA generativa. *Ejemplo directivo: Un prompt que genera un informe mediocre puede mejorarse en 3-4 iteraciones añadiendo contexto, formato esperado y ejemplos.* Ver también: [Prompt](#) · [Few-Shot Prompting](#) · [Chain-of-Thought](#)

L

LLM (Large Language Model / Modelo de Lenguaje Grande) Tipo de modelo de IA entrenado con enormes cantidades de texto para predecir y generar lenguaje. Es la tecnología subyacente de ChatGPT, Claude, Gemini y Copilot. *Ejemplo directivo: Un LLM no “entiende” el texto como un humano; predice estadísticamente qué palabras deben seguir a las anteriores, lo que produce respuestas coherentes pero no necesariamente correctas.* Ver también: [IA Generativa](#) · [Fine-tuning](#) · [Alucinación](#)

M

Make (Automatización) Plataforma de automatización visual (antes Integromat) que permite conectar aplicaciones y crear flujos de trabajo automatizados sin programación. Especialmente potente para integrar herramientas de IA con sistemas empresariales. *Ejemplo directivo: Make puede conectar un formulario de contacto web con ChatGPT para generar una respuesta personalizada y enviarla automáticamente por email.* Ver también: [Automatización](#) · [API](#) · [Zapier](#)

Manus Agente de IA autónomo desarrollado por Manus AI, capaz de ejecutar tareas complejas de múltiples pasos usando herramientas (navegador web, código, archivos). A diferencia de los chatbots, Manus puede completar proyectos completos de forma autónoma. *Ejemplo directivo: Manus puede recibir la instrucción “analiza los 5 principales competidores y prepara un informe comparativo” y ejecutarla completamente sin intervención humana.* Ver también: [Agente de IA](#) · [Human-in-the-Loop](#) · [Automatización](#)

Model Card (Ficha del Modelo) Documento técnico que describe las características, capacidades, limitaciones y consideraciones éticas de un modelo de IA. El AI Act exige documentación equivalente para sistemas de alto riesgo. *Ejemplo directivo: Antes de desplegar un modelo de IA para RRHH, la Model Card debe especificar en qué datos fue entrenado, qué sesgos conocidos tiene y para qué casos de uso NO debe usarse.* Ver también: [AI Act](#) · [Explicabilidad](#) · [Responsable AI](#)

Modelo Fundacional (Foundation Model) Modelo de IA de gran escala entrenado con datos masivos que puede adaptarse a múltiples tareas. GPT-4, Claude 3 y Gemini Ultra son modelos fundacionales. Las empresas los usan como base para construir aplicaciones específicas. *Ejemplo directivo: En lugar de entrenar un modelo desde cero (coste: millones de euros), una*

empresa puede usar un modelo fundacional como base y adaptarlo con fine-tuning. Ver también: LLM · Fine-tuning · IA Generativa

Multimodal Capacidad de un sistema de IA para procesar y generar múltiples tipos de datos: texto, imágenes, audio, vídeo y código. Los modelos multimodales pueden, por ejemplo, analizar una imagen y generar texto descriptivo. *Ejemplo directivo: GPT-4V puede analizar un gráfico financiero (imagen) y generar un análisis textual de las tendencias que muestra. Ver también: LLM · IA Generativa · Output*

N

NotebookLM Herramienta de IA de Google que permite crear “cuadernos” con documentos propios y hacer preguntas sobre ellos con grounding garantizado. Especialmente útil para análisis de documentos corporativos y preparación de reuniones. *Ejemplo directivo: Un directivo puede subir 20 informes de mercado a NotebookLM y hacer preguntas como “¿cuáles son las principales amenazas competitivas según estos informes?” Ver también: RAG · Grounding · Embeddings*

O

Output (Salida) Resultado generado por un sistema de IA en respuesta a un input. La calidad del output depende de la calidad del prompt, el modelo usado y el contexto proporcionado. *Ejemplo directivo: El output de un LLM puede ser texto, código, una tabla, un resumen o una lista; el directivo debe especificar el formato deseado en el prompt. Ver también: Prompt · Inferencia · Alucinación*

Overfitting (Sobreajuste) Problema de entrenamiento en el que un modelo aprende demasiado bien los datos de entrenamiento y pierde capacidad de generalizar a nuevos datos. En la práctica, produce modelos que funcionan bien en pruebas pero mal en producción. *Ejemplo directivo: Un modelo de predicción de ventas con overfitting puede tener 99% de precisión en datos históricos pero fallar completamente con datos nuevos. Ver también: Fine-tuning · Benchmark · Inferencia*

P

Perplexity Motor de búsqueda con IA que combina búsqueda web en tiempo real con síntesis generativa. A diferencia de ChatGPT, siempre cita sus fuentes y trabaja con información actualizada. *Ejemplo directivo: Para investigar rápidamente un competidor o un mercado,*

Perplexity es más fiable que ChatGPT porque basa sus respuestas en fuentes web verificables y actuales. Ver también: Grounding · RAG · Output

Prompt Instrucción o conjunto de instrucciones que el usuario proporciona a un sistema de IA para obtener un resultado específico. La calidad del prompt determina en gran medida la calidad del output. *Ejemplo directivo: “Resume este informe en 5 puntos clave para el Consejo de Administración, en tono formal y sin tecnicismos” es un prompt mucho más efectivo que “resume esto”. Ver también: Few-Shot Prompting · Chain-of-Thought · Iteración de prompt*

Prompt Injection (Inyección de prompt) ⚠️ LEGAL URGENTE — Ataque en el que instrucciones maliciosas ocultas en datos externos (emails, documentos, páginas web) manipulan el comportamiento de un agente de IA para ejecutar acciones no autorizadas. *Ejemplo directivo: Un email malicioso puede contener instrucciones ocultas que, al ser procesadas por un agente de IA con acceso al email corporativo, hagan que el agente reenvíe información confidencial. Ver también: AI Safety · Agente de IA · Brecha de seguridad*

Prompt Injection Detection Técnicas y herramientas para detectar y neutralizar intentos de prompt injection antes de que afecten al comportamiento del sistema de IA. Incluye filtros de entrada, sandboxing y supervisión de outputs. *Ejemplo directivo: Los sistemas de IA empresariales con acceso a datos sensibles deben implementar detección de prompt injection como medida de seguridad básica. Ver también: Prompt Injection · AI Safety · Human-in-the-Loop*

R

RAG (Retrieval-Augmented Generation) Técnica que combina un modelo de lenguaje con una base de conocimiento externa. En lugar de depender solo de lo que aprendió durante el entrenamiento, el modelo busca información relevante en documentos específicos antes de responder. *Ejemplo directivo: Un sistema RAG sobre la documentación interna de la empresa puede responder preguntas sobre procedimientos, políticas y contratos con precisión y sin alucinaciones. Ver también: Embeddings · Grounding · NotebookLM*

Responsible AI (IA Responsable) Marco de principios y prácticas para desarrollar y usar IA de forma ética, justa, transparente y segura. Incluye consideraciones de sesgo, privacidad, explicabilidad, impacto social y gobernanza. *Ejemplo directivo: Una empresa con política de Responsible AI documenta qué decisiones puede tomar la IA de forma autónoma, cuáles requieren supervisión humana y cómo se auditan los resultados. Ver también: AI Safety · Alignment · Explicabilidad*

RGPD (Reglamento General de Protección de Datos) ⚠️ LEGAL URGENTE — Regulación europea que protege los datos personales de los ciudadanos de la UE. En el contexto de IA, establece restricciones sobre qué datos pueden usarse para entrenar modelos, qué decisiones automatizadas están permitidas y qué derechos tienen los afectados. *Ejemplo directivo: Usar*

datos de empleados para entrenar un modelo de evaluación de rendimiento sin su consentimiento explícito viola el RGPD. Ver también: DPA · AI Act · Brecha de seguridad

S

Sesgo algorítmico (Algorithmic Bias) ⚠️ LEGAL URGENTE — Tendencia de un sistema de IA a producir resultados sistemáticamente injustos para ciertos grupos (por género, edad, etnia, etc.) debido a sesgos en los datos de entrenamiento o en el diseño del modelo. *Ejemplo directivo: Un modelo de selección de personal entrenado con datos históricos de una empresa con poca diversidad puede discriminar sistemáticamente a candidatas mujeres.* Ver también: Responsable AI · Explicabilidad · AI Act

Shadow AI (IA en la sombra) ⚠️ LEGAL URGENTE — Uso de herramientas de IA por parte de empleados sin conocimiento ni aprobación del departamento de IT o la dirección. Crea riesgos de seguridad, privacidad y cumplimiento normativo. *Ejemplo directivo: Un empleado que usa ChatGPT gratuito para procesar datos de clientes sin que IT lo sepa está creando un riesgo de brecha de seguridad y posible incumplimiento del RGPD.* Ver también: DPA · Brecha de seguridad · Cumplimiento normativo

Sistema de IA de Alto Riesgo ⚠️ LEGAL URGENTE — Categoría del AI Act que incluye sistemas de IA usados en áreas críticas: empleo (selección, evaluación), crédito, educación, servicios esenciales, aplicación de la ley, gestión de infraestructuras críticas. Requieren documentación técnica, supervisión humana y registro en base de datos europea. *Ejemplo directivo: Un sistema que decide automáticamente qué candidatos pasan a entrevista es de alto riesgo bajo el AI Act, independientemente del tamaño de la empresa.* Ver también: AI Act · Explicabilidad · Human-in-the-Loop

SOP (Standard Operating Procedure / Procedimiento Operativo Estándar) Documento que describe paso a paso cómo ejecutar un proceso de forma consistente. En el contexto de IA, los SOPs son fundamentales para definir cuándo y cómo usar herramientas de IA, quién supervisa los outputs y cómo se escalan los problemas. *Ejemplo directivo: Un SOP de uso de IA para comunicaciones externas debe especificar qué herramientas están aprobadas, qué revisión humana es obligatoria y cómo se documenta el uso.* Ver también: Automatización · Human-in-the-Loop · Cumplimiento normativo

Synthetic Data (Datos Sintéticos) Datos generados artificialmente por modelos de IA que imitan las características estadísticas de datos reales sin contener información personal real. Útil para entrenamiento de modelos cuando los datos reales son escasos o sensibles. *Ejemplo directivo: Una empresa puede generar datos sintéticos de transacciones bancarias para entrenar un modelo de detección de fraude sin exponer datos reales de clientes.* Ver también: Fine-tuning · RGPD · DPA

T

Temperatura del modelo (Temperature) Parámetro que controla la aleatoriedad de las respuestas de un LLM. Temperatura baja (0-0.3): respuestas más deterministas y consistentes. Temperatura alta (0.7-1): respuestas más creativas y variadas. *Ejemplo directivo: Para análisis financiero o legal, usar temperatura baja. Para brainstorming creativo o generación de ideas, temperatura alta.* Ver también: LLM · Prompt · Output

Token Unidad básica de texto que procesa un LLM. Aproximadamente 1 token = 0.75 palabras en español. Los modelos tienen un límite máximo de tokens por interacción (contexto) y cobran por tokens procesados. *Ejemplo directivo: Un informe de 10 páginas (~5.000 palabras) equivale a ~6.700 tokens. GPT-4 cobra ~0.01€ por 1.000 tokens de input.* Ver también: LLM · Contexto · Inferencia

V

Vector / Base de datos vectorial Representación matemática de texto como punto en un espacio de alta dimensión. Las bases de datos vectoriales permiten búsqueda semántica eficiente: encontrar documentos similares en significado, no solo en palabras exactas. *Ejemplo directivo: Una base de datos vectorial de todos los contratos de la empresa permite buscar “cláusulas de penalización por retraso” y encontrar todos los contratos relevantes aunque usen terminología diferente.* Ver también: Embeddings · RAG · NotebookLM

Z

Zapier Plataforma de automatización que conecta más de 6.000 aplicaciones empresariales sin necesidad de programación. Similar a Make pero con mayor catálogo de integraciones y más orientado a usuarios no técnicos. *Ejemplo directivo: Zapier puede conectar automáticamente las respuestas de un formulario de leads con ChatGPT para generar un email personalizado de bienvenida.* Ver también: Automatización · API · Make

Zero-Shot Prompting Técnica de prompting en la que se pide al modelo que realice una tarea sin proporcionar ejemplos previos, solo con instrucciones. Funciona bien para tareas simples; para tareas complejas, Few-Shot o Chain-of-Thought son más efectivos. *Ejemplo directivo: “Clasifica este email como urgente, normal o informativo” es un prompt zero-shot efectivo para tareas de clasificación simple.* Ver también: Prompt · Few-Shot Prompting · Chain-of-Thought
