



Simplicity for Grants

**Make it simple.
Make it possible.**

MANUAL PRÁCTICO DE
INTELIGENCIA ARTIFICIAL
ALFABETIZACIÓN Y GOBERNANZA CON CRITERIO

Fichas SIMPLICITY de referencia rápida
Tu guía integral para dominar la IA con criterio, responsabilidad y eficacia

2026 | Primera edición

ÍNDICE DE FICHAS

FICHA 0	Cómo funciona la IA generativa	Fundamentos, predicción y naturaleza estocástica
FICHA 1	Creación de Prompts efectivos (Método 5C)	El arte de comunicarse con la IA
FICHA 2	Parametrización y protección de datos	Privacidad, semáforo de datos y anonimización
FICHA 3	Pensamiento crítico y verificación	La Triada de Verificación y señales de alerta
FICHA 4	Gobernanza y ética de la IA	EU AI Act, categorías de riesgo y cumplimiento
FICHA 5	Sesgos en la IA	Detección, tipos y estrategias de mitigación
FICHA 6	IA en la toma de decisiones	Cuándo confiar, cuándo dudar, cuándo supervisar
FICHA 7	Creación de agentes de IA	Del chat al agente autónomo en 4 pasos
FICHA 8	Biblioteca personal de Prompts	Tu multiplicador de productividad
FICHA 9	IA Agéntica y orquestación	El siguiente nivel de automatización
FICHA 10	Estrategia de implementación de IA	Hoja de ruta para organizaciones
ANEXO 1	Tipos y herramientas de IA	Tradicional, generativa y agéntica
ANEXO 2	Glosario de términos clave	55+ conceptos explicados con claridad

Basado en el ebook IA CON CRITERIO que puedes comprar en AMAZON: <https://a.co/d/0inPo3ws> y en el [Libro Blanco de la IA](#).



FICHA

0

CÓMO FUNCIONA LA IA GENERATIVA

Fundamentos que todo usuario debe conocer

Concepto clave: Cuando lo verosímil se disfraza de verdad

Un LLM (Modelo de Lenguaje Grande) no piensa como una persona. No tiene experiencias, ni intención, ni un modelo del mundo con el que contrastar lo que dice. Su habilidad consiste en producir la siguiente palabra de la forma más coherente posible con tu petición y con los patrones que aprendió.

Eso es potentísimo para redactar, estructurar, sintetizar y proponer... pero no es lo mismo que razonar.

Los 4 principios que debes interiorizar

Principio	Qué significa	Implicación práctica
A) No razona: predice	La IA converge hacia lo que suena plausible, no hacia lo correcto. Es como un GPS que sugiere rutas probables, no un juez que determina la verdad.	Si le pides certezas, te devolverá fluidez. La verdad exige fuente, contraste y contexto.
B) Habla como si supiera	Las personas detectamos que no sabemos y frenamos. La IA, por diseño, tiende a responder siempre con tono seguro, incluso cuando improvisa.	Exige que marque incertidumbres: [VERIFICAR], [ESTIMACIÓN]. Si no explicita supuestos, no es fiable.
C) Correlación no es causalidad	Los modelos aprenden qué cosas aparecen juntas, no qué causa qué. Confundirlo en decisiones complejas es peligroso.	Pregunta siempre: ¿Qué parte es evidencia y qué parte es hipótesis? ¿Qué datos probarían causalidad?
D) El riesgo es delegar el juicio	Un texto ordenado y seguro genera sensación de control. Pero el control real está en decidir qué aceptar, qué rechazar y qué verificar.	La IA asiste y amplifica, pero no sustituye el criterio humano, especialmente con valores y consecuencias.

Naturaleza estocástica: por qué cada respuesta es diferente

Los modelos de IA generativa no son deterministas. **Generan texto mediante un proceso de muestreo probabilístico** donde predicen el siguiente token (palabra o subpalabra) basándose en distribuciones de probabilidades. La temperatura controla la creatividad: a mayor temperatura, más variación. A temperatura 0, las respuestas se vuelven casi idénticas.

Regla de oro de la Ficha 0

En la era de la IA generativa, pensar bien será más valioso que responder rápido.

FICHA

1

CREACIÓN DE PROMPTS EFECTIVOS

El Método 5C para instrucciones precisas

Regla de oro

Un buen prompt es una orden clara a un colaborador experto: qué quieres, para quién, con qué criterios y en qué formato. Si no lo pedirías así a un humano, el modelo tampoco lo adivina.

El Método de las 5C

Un marco estructurado para escribir instrucciones que generan resultados precisos, útiles y directamente aplicables. La diferencia entre una respuesta aproximada y un resultado útil está en la especificidad.

C	Elemento	Descripción	Ejemplo
1	CONTEXTO	Dónde estás y qué necesitas. Proporciona el escenario completo.	Soy gestor de proyectos I+D en universidad, necesito informe para comité en 2 semanas.
2	CRITERIO	Cómo evaluar si la respuesta sirve. Define estándar de calidad.	Debe permitir decidir inversión, incluir evidencia cuantificable y casos comparables.
3	CONTENIDO	Qué debe incluir específicamente. Lista exhaustiva de elementos.	Definición (1 párrafo), 3 casos éxito con métricas, KPIs, riesgos, recomendación clara.
4	CONDICIONANTES	Límites de forma y formato: tono, longitud, estructura.	Máx 3 páginas, tono profesional-accesible, resumen en 5 claves, desarrollo, recomendación.
5	COMPROBACIÓN	Qué debe verificar. Pide que explicité supuestos.	Indica criterios asumidos, marca estimaciones como [ESTIMACIÓN], señala datos no verificados.

Prompt mágico (copiar/pegar)

Actúa como [ROL] y crea [RESULTADO] para [AUDIENCIA/OBJETIVO]. Formato: [TABLA/LISTA/TEXTO]. Tono: [FORMAL/TÉCNICO/CLARO]. Idioma: [X]. Longitud: [X]. Criterios: incluye [...] y evita [...]. Si te falta información, pregunta antes de asumir.

Antídoto contra alucinaciones

Añade siempre al final de tu prompt: **"Si no encuentras el dato en el documento aportado o no estás seguro, di que no aparece y pide el dato necesario. No inventes."**

7 errores comunes y soluciones

Error	Solución
Asumir que la IA conoce tu contexto	Indica audiencia, objetivo, limitaciones y ejemplos en cada conversación nueva.
Usar adjetivos vagos (profesional, creativo)	Define con ejemplos concretos lo que significa cada adjetivo para ti.
Conformarse con la primera respuesta	El valor está en iterar. Pide mejoras, ajustes y versiones alternativas.
No pedir que marque incertidumbres	Solicita siempre etiquetas: [VERIFICAR], [ESTIMACIÓN], [SUPUESTO].
Pedir todo a la vez	Divide en pasos y valida el paso 1 antes de seguir al siguiente.
No especificar formato	Exige estructura: tabla, viñetas, máximo de palabras o páginas.
No controlar supuestos	Añade: Si falta información, pregunta; no inventes.

Ejemplo:

Actúa como un consultor senior en estrategia y análisis de negocio, especializado en ayudar a equipos directivos y comités de decisión a tomar decisiones informadas basadas en evidencia clara, estructurada y accionable.

Tu objetivo es generar un informe analítico de alta calidad que permita a un comité de decisión evaluar con rigor una iniciativa estratégica y decidir si debe aprobarse, modificarse o descartarse.

Tarea:

Elabora el contenido solicitado siguiendo estrictamente el Método de las 5C para garantizar claridad, precisión y utilidad práctica del resultado.

Paso 1 – CONTEXTO

Soy responsable de innovación en una empresa mediana del sector tecnológico en España. Estoy preparando un documento para un comité ejecutivo que debe decidir, en un plazo de dos semanas, si invertir o no en el lanzamiento de un nuevo producto digital dirigido al mercado B2B.

Paso 2 – CRITERIO

La respuesta será considerada de alta calidad si:

- Permite tomar una decisión clara de inversión (sí / no / con condiciones).
- Se apoya en datos cuantificables, aunque algunos sean estimaciones claramente marcadas.
- Incluye comparaciones con al menos tres casos similares del mercado.
- Presenta riesgos y beneficios de forma equilibrada y explícita.

Paso 3 – CONTENIDO

Incluye obligatoriamente:

- Definición clara del producto y su propuesta de valor (1 párrafo).
- Análisis del mercado objetivo y tamaño estimado.
- Tres casos comparables de éxito o fracaso con métricas clave.
- KPIs relevantes para evaluar el desempeño en los primeros 12 meses.
- Principales riesgos (técnicos, financieros y de mercado).
- Recomendación final clara y justificada.

Paso 4 – CONDICIONANTES

- Longitud máxima: 1.200 palabras.
- Tono: profesional, claro y accesible para perfiles no técnicos.
- Estructura: resumen ejecutivo en 5 viñetas, desarrollo por secciones, recomendación final.
- Formato: texto estructurado con encabezados claros y listas cuando sea útil.
- Editado en Word garamond 12.
- En los datos relevantes siempre pon la fuente que lo genera y trata siempre que sean oficiales.

Paso 5 – COMPROBACIÓN

Antes de finalizar:

- Haz explícitos todos los supuestos utilizados.
- Marca cualquier dato no confirmado como [ESTIMACIÓN].
- Señala claramente qué información requeriría validación adicional.

No des un dato por correcto sin verificación de la fuente.

FICHA 2

PARAMETRIZACIÓN Y PROTECCIÓN DE DATOS

Tu primera línea de defensa

Por qué es crítico

El 90% de los problemas de privacidad con IA no vienen de mala intención, sino de descuido. Una configuración de 5 minutos puede prevenir meses de problemas legales, reputacionales o de cumplimiento normativo (RGPD, secretos comerciales, datos personales).

Checklist de configuración segura (5 minutos)

1. Revisar política de privacidad (buscar: data usage, training, retention)
2. Desactivar uso de datos para entrenamiento (Settings > Privacy > Disable training)
3. Controlar historial de conversaciones (desactivar guardado si datos sensibles)
4. Configurar retención de datos (eliminar tras X días si es posible)
5. Revisar permisos de terceros (plugins e integraciones con acceso)
6. Activar autenticación de dos factores (2FA)
7. Usar cuentas separadas (trabajo vs personal, nunca mezclar contextos)

El semáforo de datos

VERDE: PÚBLICA	Datos publicados en web, normativa oficial, notas de prensa, dominio público. USO LIBRE en IA.
AMARILLO: INTERNA	Procesos internos (sin nombres), datos agregados, borradores. USO CON ANONIMIZACIÓN.
ROJO: CONFIDENCIAL	DNI, emails personales, contratos, nóminas, datos de salud, secretos comerciales, claves. NUNCA usar en IA pública.

Técnicas de anonimización rápida

Tipo de dato	Transformación
Nombres propios	[PERSONA_A], [EMPRESA_1], [SOCIO_X]
Ubicaciones	[CIUDAD_EUROPEA], [OFICINA_REGIONAL]
Fechas exactas	[Q3_2024], [HACE_2_MESES]
Cifras sensibles	[PRESUPUESTO_ALTO], [>500K]
Nombres de proyectos	[PROYECTO_A], [INICIATIVA_X]

FICHA 3

PENSAMIENTO CRÍTICO Y VERIFICACIÓN

La Triada que separa información de desinformación

Por qué es no negociable

La IA puede inventar información con total confianza. Sin verificación, un dato falso puede propagarse en informes, propuestas y decisiones estratégicas. La verificación no es opcional cuando las consecuencias son reales.

Los 3 filtros de la Triada

Filtro	Qué verifica	Preguntas clave
1. FUENTE	¿De dónde sale la información?	¿Es fuente oficial o creíble? ¿Puede verificarse de forma independiente? ¿Hay varias fuentes que lo confirmen?
2. COHERENCIA	¿Tiene sentido internamente?	¿Hay contradicciones internas? ¿Las cifras cuadran entre sí? ¿La lógica argumentativa es sólida?
3. EXPERTO	¿Qué dirían los especialistas?	¿Pasa la prueba del olfato profesional? ¿Qué opinaría un experto del sector? ¿Suena razonable para quien conoce el tema?

Señales de alerta: cuándo desconfiar inmediatamente

Señal de alerta	Por qué desconfiar	Qué hacer
Cifras exactas sin fuente	73.4% de empresas reportan... ¿Quién lo midió? ¿Cuándo?	Pedir la fuente original. Si no la da, marcar como no verificado.
Afirmaciones categóricas	Siempre, Nunca, Todo... La realidad siempre tiene matices.	Pedir excepciones y condiciones específicas.
Citas sin atribución	Como dijo Einstein... Sin verificar si realmente lo dijo.	Buscar la cita original en fuentes primarias.
Fechas recientes específicas	Eventos de últimos meses pueden estar fuera de los datos de entrenamiento.	Verificar con búsqueda web actualizada.
Contradicciones internas	Dice una cosa y luego se contradice en la misma respuesta.	Señalar la contradicción y pedir aclaración.

Herramientas complementarias de verificación

Herramienta	Para qué sirve
Perplexity / Bing Chat	IA con búsqueda web integrada y citas automáticas de fuentes.
Google Scholar	Verificar estudios, papers académicos y publicaciones científicas.
Snopes / FactCheck.org	Verificar afirmaciones virales y desinformación.
Wayback Machine	Ver versiones históricas de páginas web y verificar cambios.
Búsqueda inversa de imágenes	Verificar autenticidad y origen de imágenes.

FICHA
4

GOBERNANZA Y ÉTICA DE LA IA
El marco europeo que debes conocer

Por qué es imprescindible

La UE ha aprobado el Reglamento de Inteligencia Artificial (AI Act, Regulación 2024/1689), el primer marco legal integral del mundo para la IA. Afecta a cualquier organización que desarrolle, despliegue o use sistemas de IA en Europa. Conocerlo no es opcional: es obligación legal.

El sistema de riesgo del AI Act europeo

El AI Act clasifica los sistemas de IA en cuatro niveles de riesgo, cada uno con obligaciones diferentes:

Nivel de riesgo	Descripción	Ejemplos	Obligaciones clave
INACEPTABLE (Prohibido)	Sistemas que vulneran derechos fundamentales. Prohibidos en la UE.	Puntuación social, manipulación subliminal, vigilancia biométrica masiva en tiempo real.	Prohibición total. No pueden desarrollarse, desplegarse ni usarse.
ALTO	Sistemas con impacto significativo en derechos o seguridad de personas.	Selección de personal con IA, scoring crediticio, diagnóstico médico, gestión de infraestructuras críticas.	Evaluación de conformidad, gestión de riesgos, transparencia, supervisión humana, registro en base de datos de la UE.
LIMITADO	Sistemas que interactúan con personas o generan contenido.	Chatbots, deepfakes, generación de texto/imagen por IA.	Obligaciones de transparencia: informar al usuario de que interactúa con IA o que el contenido es generado por IA.
MÍNIMO	Sistemas sin riesgos significativos.	Filtros de spam, IA en videojuegos, sistemas de recomendación no invasivos.	Sin obligaciones específicas. Se recomienda seguir códigos de buenas prácticas.

Calendario de aplicación del AI Act EU

Fecha	Qué entra en vigor
Febrero 2025	Prohibiciones de prácticas de riesgo inaceptable.
Agosto 2025	Obligaciones para modelos de propósito general (GPAI) y gobernanza.
Agosto 2026	Obligaciones completas para sistemas de alto riesgo (Anexo III).
Agosto 2027	Obligaciones para sistemas de alto riesgo integrados en productos regulados.

Los 7 principios de la IA confiable (HLEG)

El Grupo de Expertos de Alto Nivel de la UE definió 7 requisitos para una IA de confianza:

Principio	Qué significa en la práctica
1. Acción y supervisión humana	El humano mantiene control. La IA asiste, no sustituye el juicio en decisiones con impacto.
2. Robustez técnica y seguridad	El sistema funciona de forma fiable, resiste ataques y tiene planes de contingencia.
3. Privacidad y gobernanza de datos	Protección de datos personales, consentimiento, minimización y cumplimiento RGPD.
4. Transparencia	Explicar cómo funciona el sistema, sus limitaciones y cuándo se usa IA.
5. Diversidad, no discriminación, equidad	Evitar sesgos injustos y garantizar accesibilidad para todos los usuarios.
6. Bienestar social y ambiental	Considerar el impacto en sociedad, medio ambiente y sostenibilidad.
7. Rendición de cuentas	Mecanismos de auditoría, trazabilidad y responsabilidad ante errores o daños.

Checklist rápido de gobernanza para tu organización

Autoevaluación: ¿Cumples lo básico?

1. ¿Sabes qué sistemas de IA usas y en qué categoría de riesgo caen?
 2. ¿Tienes documentada la finalidad de cada uso de IA?
 3. ¿Hay supervisión humana efectiva en decisiones con impacto en personas?
 4. ¿Informas a los usuarios cuando interactúan con IA?
 5. ¿Evalúas periódicamente sesgos y rendimiento de tus sistemas de IA?
 6. ¿Tienes un responsable designado de gobernanza de IA?
 7. ¿Cumples con el RGPD en el tratamiento de datos por IA?
- Si respondiste NO a más de 3 preguntas, necesitas un plan de gobernanza urgente.

FICHA**5****SESGOS EN LA IA***Detectar, entender y mitigar***Por qué importa**

Los sesgos en la IA no son errores aleatorios: son patrones sistemáticos que reproducen y amplifican desigualdades existentes en los datos y en la sociedad. Una IA sesgada puede discriminar en contratación, crédito, salud o justicia sin que nadie lo note.

Los 6 tipos principales de sesgo en IA

Tipo de sesgo	Qué es	Ejemplo real
Sesgo de datos (histórico)	Los datos de entrenamiento reflejan desigualdades pasadas.	IA de contratación penalizaba CVs con la palabra mujer porque históricamente se contrataban más hombres.
Sesgo de representación	Grupos infrarrepresentados en los datos de entrenamiento.	Sistemas de reconocimiento facial con menor precisión en personas de piel oscura.
Sesgo de confirmación	El modelo refuerza creencias preexistentes del usuario.	Burbuja informativa: la IA muestra solo lo que cree que quieres oír.
Sesgo de automatización	Confiar ciegamente en la IA porque es tecnología.	Aceptar un diagnóstico médico de IA sin segunda opinión humana.
Sesgo de anclaje	La primera información recibida condiciona el análisis.	Si el prompt incluye una cifra, la IA tiende a anclar sus estimaciones cerca de ella.
Sesgo de supervivencia	Solo se analizan los casos exitosos, ignorando los fracasos.	IA entrenada solo con empresas exitosas no aprende patrones de fracaso.

Estrategias de mitigación prácticas

Estrategia	Cómo aplicarla	Nivel de dificultad
Auditoría de datos	Revisar los datos de entrenamiento buscando desequilibrios en representación por género, edad, origen, etc.	Medio
Pruebas de equidad	Ejecutar el mismo proceso con perfiles diversos y comparar resultados (A/B testing de sesgo).	Bajo
Diversidad en equipos	Equipos diversos detectan sesgos que equipos homogéneos no ven. Incluir perspectivas variadas.	Bajo
Transparencia de modelos	Exigir explicabilidad: por qué el modelo tomó esta decisión, qué variables pesaron más.	Alto

Supervisión continua	Monitorizar resultados en producción para detectar sesgos emergentes (drift de sesgo).	Medio
Feedback de usuarios	Crear canales para que los afectados reporten resultados injustos o discriminatorios.	Bajo

FICHA

6

IA EN LA TOMA DE DECISIONES

Cuándo confiar, cuándo dudar, cuándo supervisar

Principio rector

La IA es una herramienta de apoyo a la decisión, no un sustituto del juicio humano. Cuanto mayor sea el impacto de la decisión (dinero, personas, reputación, derechos), mayor debe ser la supervisión humana.

La matriz de decisión: cuándo delegar a la IA

Nivel de delegación	Tipo de tarea	Ejemplos	Supervisión requerida
AUTOMATIZAR	Rutinaria, bajo riesgo, alto volumen	Clasificar emails, generar borradores, resumir documentos, transcribir	Revisión por muestreo (10-20%)
ASISTIR	Análisis o propuesta, riesgo medio	Análisis de datos, comparativas de mercado, propuestas de texto	Revisión humana antes de usar el resultado
CONSULTAR	Decisión compleja, alto impacto	Estrategia de inversión, evaluación de candidatos, diagnóstico médico	IA como insumo más entre varios. Decisión siempre humana
NO DELEGAR	Valores, ética, relaciones personales	Despidos, comunicación de crisis, decisiones legales vinculantes	La IA puede informar, pero la decisión y la responsabilidad son 100% humanas

El protocolo DECIDE con IA

Paso	Acción	Pregunta guía
D - Definir	Define el problema con claridad antes de consultar la IA.	¿Qué estoy intentando resolver exactamente?
E - Explorar	Usa la IA para explorar opciones y perspectivas que no habías considerado.	¿Qué alternativas existen que no he visto?
C - Contrastar	Compara las respuestas de la IA con otras fuentes y tu conocimiento experto.	¿Coincide con lo que sé? ¿Tiene lagunas?
I - Identificar	Identifica supuestos, sesgos y limitaciones en las respuestas.	¿Qué está asumiendo la IA? ¿Qué podría faltar?
D - Deliberar	Delibera con personas clave antes de tomar la decisión final.	¿Qué opinan quienes serán afectados?

E - Ejecutar	Ejecuta la decisión documentando el proceso y las fuentes.	¿Puedo justificar esta decisión si me preguntan?
--------------	--	--

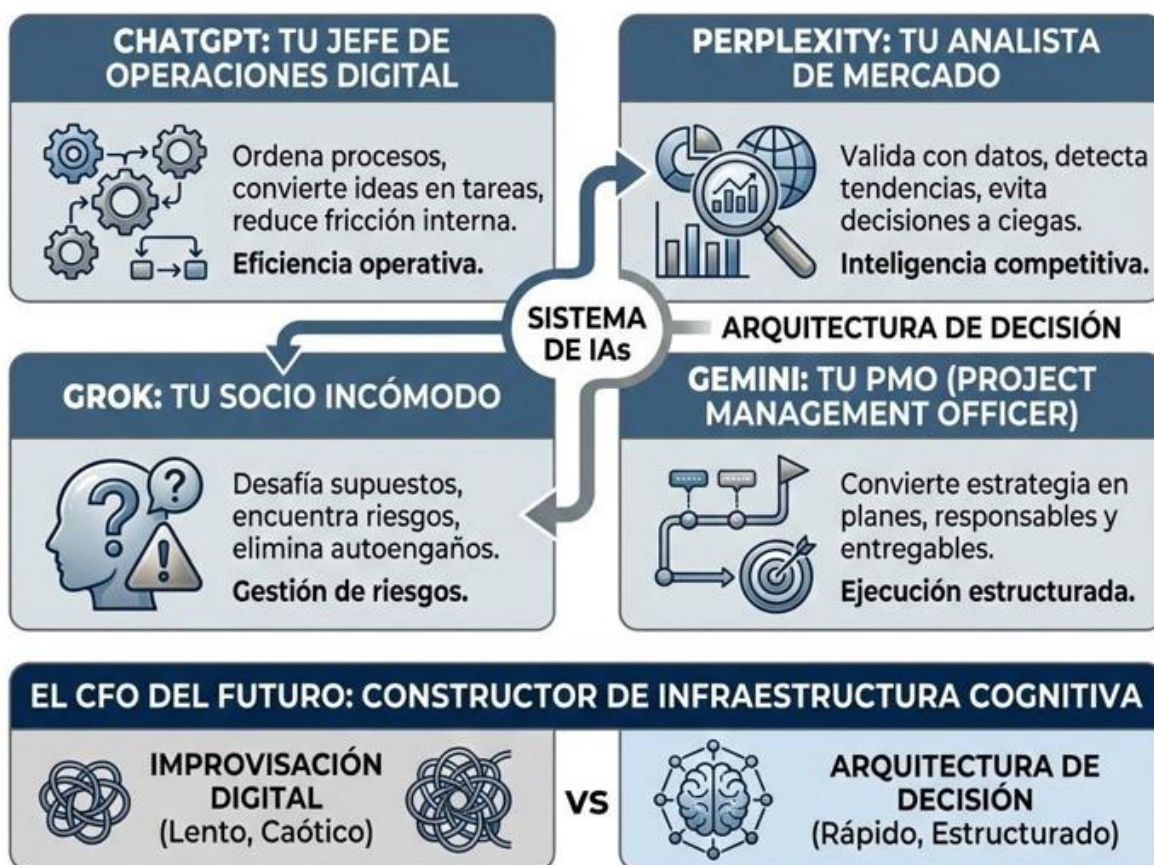
Pregunta clave antes de cada decisión asistida por IA

"Si esta decisión sale mal, ¿podré explicar por qué la tomé y qué información usé?"

Si la respuesta es no, necesitas más supervisión humana.

IA NO ES UNA HERRAMIENTA. ES UNA ARQUITECTURA DE GESTIÓN.

UNA VISIÓN CFO: DE APLICACIONES A ROLES ESTRATÉGICOS.



El nuevo cuello de botella no es el dinero. Es la capacidad de PENSAR, DECIDIR y EJECUTAR más rápido que el mercado.

INTERIM
MANAGER
CONSULTING

FICHA
7

CREACIÓN DE AGENTES DE IA
Del chat básico al asistente autónomo

Qué es un agente de IA

Un sistema que realiza tareas de forma autónoma o semi-autónoma, tomando decisiones basadas en objetivos y ejecutando múltiples pasos secuencialmente. A diferencia del chat básico, los agentes usan herramientas, mantienen contexto y toman decisiones intermedias.

Chat IA vs Agente de IA

Característica	Chat IA normal	Agente de IA
Autonomía	Responde a cada pregunta individual	Descompone y ejecuta objetivos complejos
Herramientas	Solo genera texto	Usa APIs, bases de datos, navegadores, código
Memoria	Solo la conversación actual	Mantiene contexto entre sesiones
Decisiones	Requiere dirección constante del usuario	Toma decisiones intermedias autónomamente
Iteración	Una respuesta por prompt	Intenta múltiples enfoques si algo falla

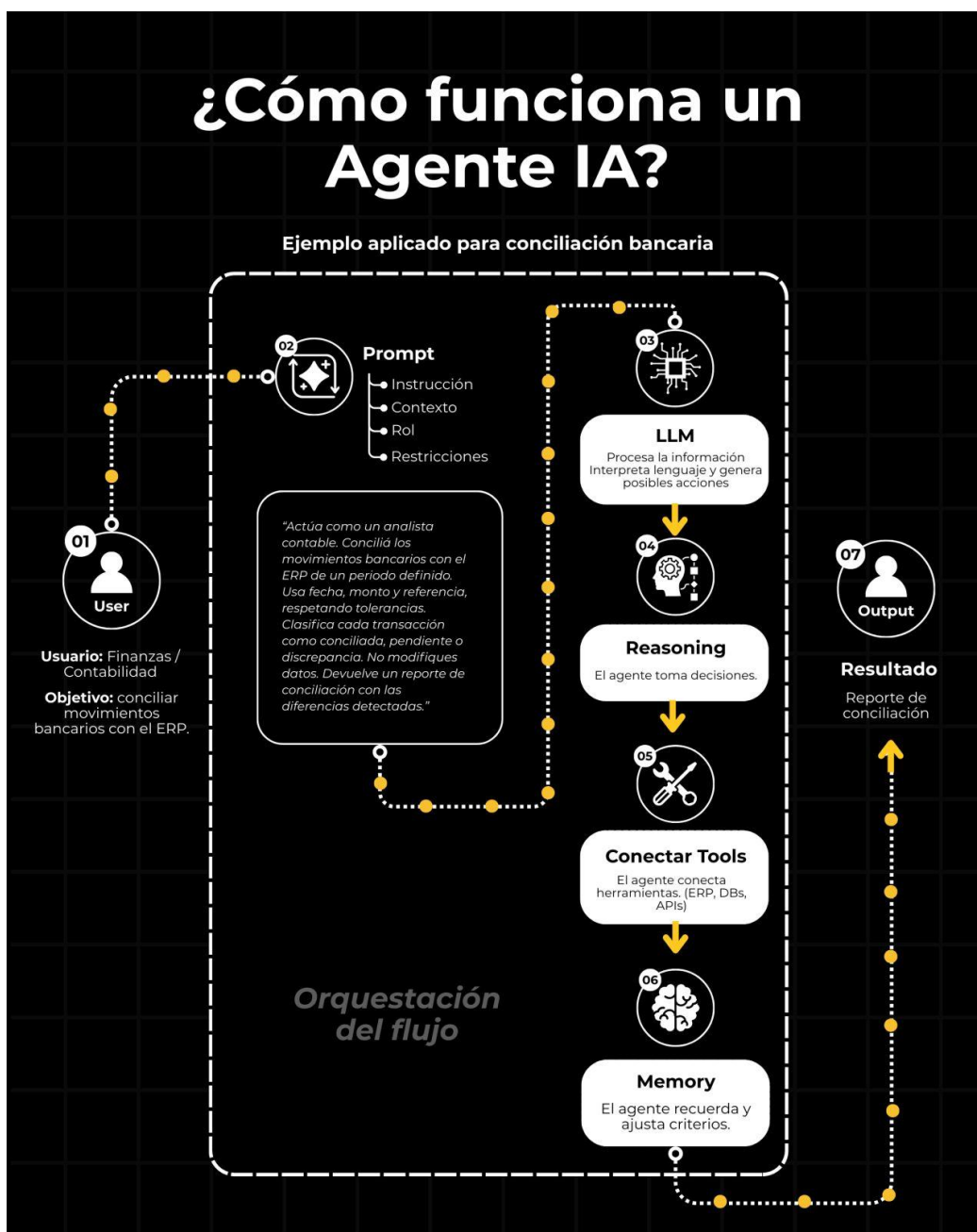
Cómo crear tu primer agente (4 pasos)

Paso	Acción	Ejemplo práctico
1. OBJETIVO	Define un objetivo claro y medible. NO: Ayudarme con proyectos. SÍ: un objetivo específico.	Monitorizar 5 fuentes de convocatorias cada lunes y enviarme resumen de oportunidades relevantes.
2. SUBTAREAS	Descompón en pasos secuenciales concretos.	Acceder a URLs > Extraer datos > Filtrar > Estructurar en tabla > Enviar resumen por email.
3. HERRAMIENTAS	Identifica qué herramientas necesita el agente.	Navegador web + Extractor de texto + Filtrado + Generador de tablas + Cliente email.
4. PLATAFORMA	Elige según tu nivel técnico.	Principiante: GPTs/Zapier. Intermedio: Claude Projects/n8n. Avanzado: LangChain/LangGraph.

3 Consideraciones críticas

- 1. Supervisión humana necesaria:** Los agentes pueden errar en cadenas largas. Define checkpoints de confirmación y revisa logs.
- 2. Gestión de costes:** Los agentes consumen tokens rápidamente (múltiples llamadas). Establece límites mensuales.

3. Seguridad y privacidad: Nunca des acceso a datos ROJOS. Usa credenciales con permisos mínimos. Aplica el semáforo de datos.



FICHA

8

BIBLIOTECA PERSONAL DE PROMPTS

Tu multiplicador de productividad $\times 10$

El dato que convence

Escribir un buen prompt desde cero: 5-10 minutos. Usar uno de tu biblioteca: 30 segundos. Con 15-20 prompts bien contruidos cubres el 80% de tus necesidades habituales.

Sistema de 3 pasos para construir tu biblioteca

Paso	Acción	Ejemplo
1. CAPTURA	Cuando un prompt funcione bien, cópialo inmediatamente.	Usaste un prompt para resumir informe y el resultado fue excelente. Guardar.
2. ETIQUETA	Asigna categoría clara y descriptiva.	[RESUMEN] [EMAIL] [ANÁLISIS] [PLANIFICACIÓN] [VERIFICACIÓN]
3. REFINA	Revisa mensualmente y mejora los que más usas.	Prompt de email usado 20 veces: optimizar para reducir edición manual.

10 Prompts esenciales para empezar

N.	Categoría	Descripción
1	Resumen	Resumir documento largo extrayendo lo esencial en formato ejecutivo.
2	Planificación	Convertir ideas desestructuradas en plan de acción con fases y plazos.
3	Redacción	Mejorar claridad y legibilidad de texto técnico para audiencia no experta.
4	Reuniones	Preparar agenda, puntos clave y preguntas anticipadas para reunión.
5	Comunicación	Responder a objeción o crítica de forma constructiva y profesional.
6	Presentación	Estructurar presentación con narrativa clara y mensajes clave.
7	Análisis	Analizar datos/información y extraer insights accionables.
8	Creatividad	Generar alternativas creativas a un problema definido.
9	FAQ	Anticipar preguntas frecuentes y preparar respuestas claras.
10	Revisión	Revisar documento antes de enviar (calidad, tono, coherencia, errores).

Dónde guardar tu biblioteca

Opción	Ventaja principal
Google Docs / Word	Fácil búsqueda y edición. Ideal para empezar.
Notion / Obsidian	Etiquetas, enlaces internos y búsqueda avanzada.
Hoja de cálculo	Columnas: Categoría Nombre Prompt Frecuencia de uso.
Dentro de la IA	Custom Instructions (ChatGPT), Projects (Claude).

FICHA
9

IA AGÉNTICA Y ORQUESTACIÓN
El siguiente nivel de automatización inteligente

Definición

Es la IA que toma decisiones y ejecuta acciones de forma autónoma para cumplir objetivos complejos. No solo genera respuestas: completa tareas completas sin supervisión constante, usando herramientas, planificando pasos y verificando resultados.

Las 3 capas de IA que se combinan

Capa	Qué hace	Riesgo diferencial	Control mínimo
IA Tradicional (predictiva)	Predice, clasifica y detecta anomalías con datos numéricos.	Descalibración si cambia el contexto (drift).	Métricas + validación + monitorización continua.
IA Generativa (lenguaje)	Genera texto, código, imágenes y estructura ideas.	Plausible con tono de certeza cuando faltan datos.	RAG + fuentes + supuestos explícitos + verificación.
IA Agéntica (acción)	Planifica, usa herramientas y ejecuta tareas completas.	Actúa sobre sistemas reales: un error puede escalar.	Human-in-the-loop + permisos mínimos + trazabilidad.

Mercado y oportunidad en Europa

Indicador	Dato 2024	Proyección 2030
Mercado global AI Agents	5.400 M USD	50.300 M USD (CAGR 45,8%)
Cuota europea	24,5% del mercado global	11.490 M USD
Inversión VC en Agentic AI (UE)	5% de toda la inversión VC	Tendencia creciente

Criterio de oro para IA agéntica

Cuanto más cerca estés de decisiones con valores, dinero, reputación o personas, más necesitas evidencia, trazabilidad y supervisión humana.

FICHA

10

ESTRATEGIA DE IMPLEMENTACIÓN DE IA

Hoja de ruta para organizaciones

Principio rector

La implementación de IA no es un proyecto tecnológico: es un proyecto de transformación organizativa. La tecnología es el 20%; las personas, los procesos y la cultura son el 80%.

Hoja de ruta en 5 fases

Fase	Duración	Objetivos clave	Entregables
1. DIAGNÓSTICO	2-4 semanas	Mapear procesos, identificar oportunidades y evaluar madurez digital.	Mapa de procesos, análisis de madurez, listado de casos de uso.
2. PILOTO	4-8 semanas	Seleccionar 1-2 casos de uso, probar con equipo reducido, medir resultados.	POC funcional, métricas de impacto, lecciones aprendidas.
3. ESCALADO	2-3 meses	Ampliar a más equipos y procesos, definir gobernanza y protocolos.	Política de uso de IA, formación, infraestructura.
4. INTEGRACIÓN	3-6 meses	Incorporar IA en flujos de trabajo habituales, automatizar procesos.	Workflows automatizados, KPIs de rendimiento, presupuesto ajustado.
5. OPTIMIZACIÓN	Continuo	Monitorizar, iterar, medir ROI y adaptar a nuevas capacidades.	Dashboard de seguimiento, informes trimestrales, plan de mejora.

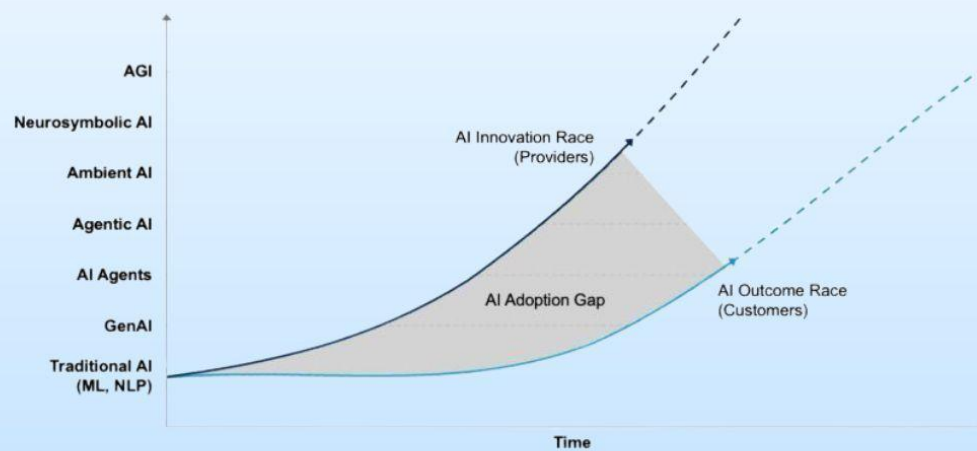
Los 5 errores fatales en implementación de IA

Error	Consecuencia	Cómo evitarlo
Empezar por la tecnología, no por el problema	Soluciones sofisticadas para problemas que no existen.	Siempre empieza preguntando: ¿Qué problema concreto resolvemos?
No involucrar a los usuarios finales	Resistencia, desconfianza y baja adopción.	Co-diseña con los equipos que usarán la herramienta desde el día 1.
Ignorar la calidad de los datos	Basura entra, basura sale. La IA amplifica los problemas de datos.	Audita y limpia datos antes de alimentar cualquier sistema.
No definir métricas de éxito	Imposible saber si funciona o si vale la inversión.	Define KPIs claros antes de empezar: tiempo ahorrado, errores reducidos, etc.
Olvidar la gobernanza	Riesgos legales, éticos y reputacionales sin control.	Desde el día 1: política de uso, semáforo de datos, supervisión humana.

Impacto cuantitativo esperado

Indicador	Rango típico
Reducción de tiempo en procesamiento de datos	70-80%
Reducción de tiempo en comunicaciones	30-50%
Aumento de proyectos gestionables por profesional	25-40%
Mejora en capacidad de respuesta a clientes	35-45%
ROI promedio en implementaciones exitosas	3x-5x
Resolución de consultas sin intervención humana	60-80%
Reducción de errores tras validación con IA	35-70%

The AI Adoption Gap



Source: Gartner
837663

Gartner®

ANEXO 1

TIPOS Y HERRAMIENTAS DE IA

Mapa completo del ecosistema

Las 3 capas explicadas en 10 segundos cada una

Capa	Explicación rápida	Para qué sirve	Dónde falla
IA Tradicional	Mira el pasado para estimar el futuro o clasificar el presente.	Forecast, scoring, segmentación, detección de fraude.	Cuando cambian las condiciones del entorno (drift).
IA Generativa	Escribe y organiza muy bien, pero necesita fuentes y revisión.	Borradores, síntesis, emails, informes, código, documentación.	Puede rellenar con confianza cuando faltan datos.
IA Agéntica	No solo responde: se organiza, usa herramientas y ejecuta.	Flujos completos: búsqueda + análisis + redacción + envío.	Actúa sobre sistemas reales: un error puede escalar.

Niveles de inteligencia artificial

Para desarrollar un criterio sólido en IA, es fundamental entender sus niveles de desarrollo y capacidad. No todas las IA son iguales: algunas son herramientas especializadas, mientras que otras (aún hipotéticas) podrían transformar la sociedad.

Clasificamos la IA en tres niveles principales según su alcance y autonomía:

1. **IA Estrecha (ANI, Artificial Narrow Intelligence):** Representa la mayoría de los sistemas de IA que usamos hoy. Estos están diseñados para tareas específicas y limitadas, como reconocer imágenes en un teléfono, traducir textos o recomendar productos en una tienda online. Son eficientes en lo que hacen, pero no "entienden" el mundo ni pueden aplicar su conocimiento a áreas diferentes. Ejemplo: un asistente como ChatGPT genera texto basado en patrones, pero no razona como un humano. La ANI es la base de la mayoría de las aplicaciones prácticas en este libro, y requiere supervisión humana para evitar errores.
2. **IA General (AGI, Artificial General Intelligence):** Se refiere a sistemas que podrían realizar cualquier tarea intelectual humana con el mismo nivel de comprensión, adaptabilidad y aprendizaje. La AGI no existe aún (en 2026, estamos lejos de ella o cerca depende el experto que nos hable), pero se especula que podría llegar en las próximas décadas. Imagina una IA que no solo genera informes, sino que también diseña estrategias empresariales, resuelve problemas científicos o maneja emergencias con juicio humano. Prepararnos para la AGI implica fortalecer la gobernanza ahora, como se discute en el capítulo 12.

3. **IA Superinteligente (ASI, Artificial Superintelligence):** Va un paso más allá de la AGI, superando la inteligencia humana en todos los aspectos, incluyendo creatividad, velocidad de procesamiento y resolución de problemas complejos. Podría innovar en campos como la medicina o la energía de formas impredecibles. Sin embargo, trae riesgos existenciales, como desalineación con valores humanos o pérdida de control. La ASI es especulativa y no inminente, pero subraya la necesidad de ética y regulaciones globales para mitigar impactos negativos.

Entender estos niveles te ayuda a contextualizar el uso actual de la IA (principalmente ANI) y a anticipar futuros desafíos. En la práctica, enfócate en gobernar lo que ya tienes, pero mantén un ojo en la evolución hacia AGI y ASI para adaptar tu estrategia.

Esta clasificación conceptual es ampliamente utilizada en literatura de IA (Bostrom, Russell, etc.), no definida en marcos regulatorios oficiales.

Herramientas principales por categoría

Categoría	Herramientas destacadas	Caso de uso
LLMs generalistas	ChatGPT (OpenAI), Claude (Anthropic), Gemini (Google), Grok (xAI), Perplexity	Redacción, análisis, síntesis, programación, soporte general.
Búsqueda con IA	Perplexity, Bing Chat, Google AI Overviews	Investigación con fuentes citadas y verificables.
Imágenes y vídeo	DALL-E, Midjourney, Runway, HeyGen	Generación de contenido visual, avatares, vídeos.
Presentaciones	Gamma, Slidesgo, Copilot + PowerPoint	Creación rápida de presentaciones profesionales.
Productividad	Notion AI, ClickUp, Miro AI	Gestión de proyectos, mapas mentales, organización.
Automatización	Zapier AI, Make (Integromat), n8n	Automatización de flujos de trabajo entre aplicaciones.
Documentos y datos	Humata AI, HyNote, Napkin AI	Análisis de documentos, transcripción, infografías.
Traducción	DeepL, Google Translate (avanzado)	Traducción profesional con contexto.

ANEXO
2

GLOSARIO DE TÉRMINOS CLAVE

55+ conceptos explicados con claridad

Referencia rápida de los conceptos más importantes de IA, ordenados temáticamente para facilitar la consulta.

Conceptos fundamentales

Término	Definición clara	Para qué sirve
Inteligencia Artificial (IA)	Campo que crea sistemas capaces de realizar tareas que requieren inteligencia: percepción, lenguaje, decisión.	Marco general que incluye ML, DL, NLP, etc.
Machine Learning (ML)	Subcampo de IA: modelos que aprenden patrones desde datos para predecir o decidir.	Predicción, clasificación, segmentación.
Deep Learning (DL)	ML basado en redes neuronales profundas (muchas capas).	Cuando hay muchos datos y alta complejidad.
LLM (Large Language Model)	Modelo de lenguaje grande que genera y razona con texto.	Chatbots, resumen, redacción, asistentes.
Token	Unidad de texto que procesa un LLM (fragmentos de palabras).	Determina coste y límites de contexto.
Context window	Máximo de tokens que el modelo puede leer a la vez.	Determina cuánto puede recordar por interacción.
Prompt	Instrucción de entrada que guía al modelo.	Control de output, estilo y formato.
Prompt engineering	Diseño sistemático de prompts para resultados fiables.	Mejorar calidad sin necesidad de reentrenar.
Hallucination	Respuesta inventada o no sustentada por datos.	Riesgo clave en LLM sin controles.
RAG	Generación aumentada con recuperación: busca info en documentos y luego genera.	Respuestas con base documental verificable.

Conceptos de gobernanza y ética

Término	Definición clara	Para qué sirve
AI Act EU	Reglamento europeo de IA (2024/1689). Primer marco legal integral del mundo.	Marco obligatorio para uso de IA en Europa.
RGPD	Reglamento General de Protección de Datos de la UE.	Protección de datos personales procesados por IA.

Sesgo (bias)	Errores sistemáticos por datos, etiquetas o diseño.	Riesgo ético y legal a monitorizar.
Explainability (XAI)	Métodos para explicar decisiones del modelo.	Regulación, confianza y diagnóstico.
Human-in-the-loop (HITL)	Supervisión humana en puntos críticos del proceso.	Requisito del AI Act para sistemas de alto riesgo.
Guardrails	Controles para limitar riesgos de seguridad, formato o contenido.	Producción responsable de IA.
Grounding	Anclar respuestas en evidencias verificables.	Fiabilidad y auditoría de resultados.
Drift (deriva)	Cambio del patrón de datos o comportamiento con el tiempo.	Mantenimiento continuo en producción.
Datos sintéticos	Datos generados artificialmente para entrenar o probar.	Privacidad y aumento de datos escasos.
Federated Learning	Entrenar sin centralizar datos sensibles.	Salud, banca, escenarios de privacidad.

Conceptos técnicos esenciales

Término	Definición clara
Entrenamiento	Proceso de ajustar parámetros del modelo minimizando una función de pérdida.
Inferencia	Uso del modelo ya entrenado para producir predicciones en tiempo real.
Fine-tuning	Reentrenar un modelo preentrenado en un dominio o tarea específica.
Transfer Learning	Reusar conocimiento de un modelo entrenado en otra tarea diferente.
Embeddings	Representaciones numéricas densas de texto/imagen que captan significado semántico.
Vector DB	Base de datos para búsqueda por similitud semántica (soporta RAG).
Temperatura	Parámetro que controla la creatividad/aleatoriedad de las respuestas.
Overfitting	El modelo memoriza los datos de entrenamiento y no generaliza a datos nuevos.
Pipeline	Cadena automatizada: limpieza de datos, entrenamiento, evaluación, despliegue.
MLOps	Prácticas para desplegar, monitorizar y gobernar modelos de ML en producción.
Agente (AI Agent)	Sistema que planifica acciones, usa herramientas y ejecuta pasos hacia un objetivo.
Tool use	Capacidad del modelo de llamar herramientas externas (web, BD, cálculos).

¿Por dónde empezar?

Si nunca has usado IA seriamente: Comienza con Ficha 1 (Método 5C) y practica 5 prompts.

Si ya usas IA pero sin protección: Ve directo a Ficha 2 (Parametrización) y configura HOY.

Si necesitas cumplir normativa: Ficha 4 (Gobernanza) es tu punto de partida obligatorio.

Si te preocupan los sesgos: Ficha 5 (Sesgos) te da las herramientas para detectarlos.

Si quieres automatizar flujos: Fichas 7 y 9 (Agentes y IA Agéntica) son tu camino.

Si trabajas con información crítica: Memoriza Ficha 3 (Triada) y aplícala siempre.

El futuro del trabajo no es humanos vs IA.

Es humanos CON IA CON CRITERIO.

Estas fichas te dan el criterio. El resto depende de ti.

LEYENDA SOBRE EL USO DE INTELIGENCIA ARTIFICIAL

La información contenida en este documento ha sido elaborada mediante la herramienta Simplicity for Grants© (SfG), una plataforma de inteligencia artificial, analítica avanzada y blockchain especializada en la identificación y análisis de ayudas públicas, subvenciones, programas europeos y licitaciones globales.

El contenido ha sido posteriormente revisado y validado por profesionales senior de IMC, en cumplimiento del principio Human-in-the-loop (HITL) de la Artificial Intelligence Act (Regulación (UE) 2024/1689).

Interim Manager Consulting, propietario de la plataforma SfG, tiene un compromiso firme con la ética y el cumplimiento normativo europeo y está adherido al Código de Buenas Prácticas para IA de Propósito General (GPAI) de la UE.

Contacto: holat@interimconsulting.es

www.interimconsulting.es

Documento generado por Simplicity for Grants© con metodología EstrategicaMente© propiedad intelectual de IMC

LEYENDA DE PROTECCIÓN Y USO RESTRINGIDO

Este documento es propiedad intelectual de su titular. Queda terminantemente prohibida su reproducción, distribución, comunicación pública, difusión, cesión o transformación, total o parcial, por cualquier medio, sin autorización expresa y por escrito del titular.

A efectos de integridad, trazabilidad y prueba de existencia, el documento ha sido registrado en blockchain. Hash / Identificador de registro: QmaNo9WXjxtH25UxApVGyG'T8Cv8rjoxo86XmLWhugykj4n/1767114499810-v2.json.

El uso no autorizado podrá dar lugar a las acciones legales que correspondan.