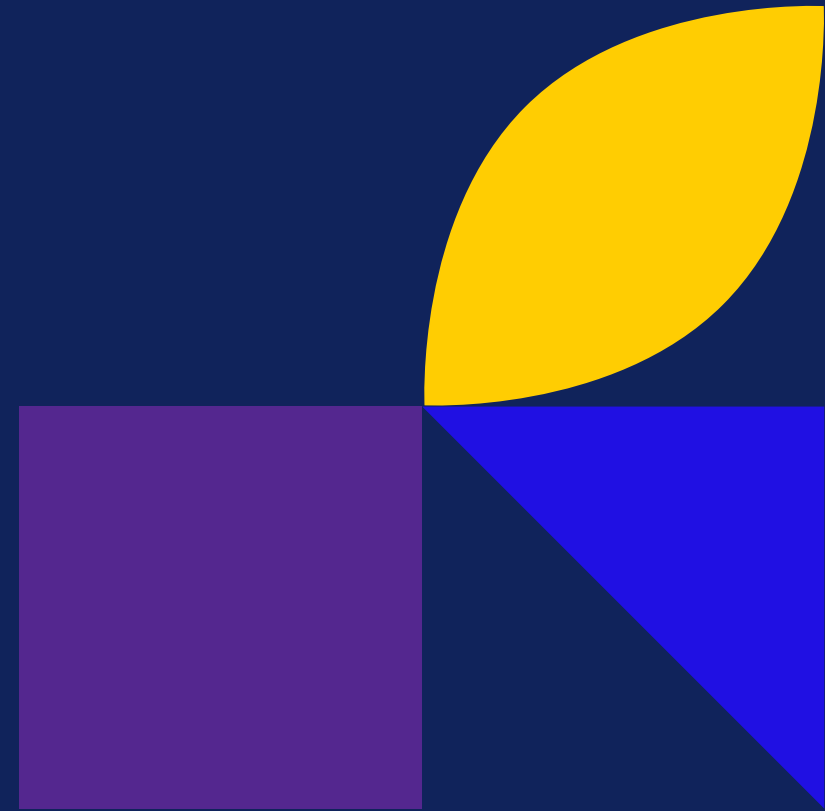


Módulo 1.

Tendencias tecnológicas emergentes: innovación e impacto

Parte 1



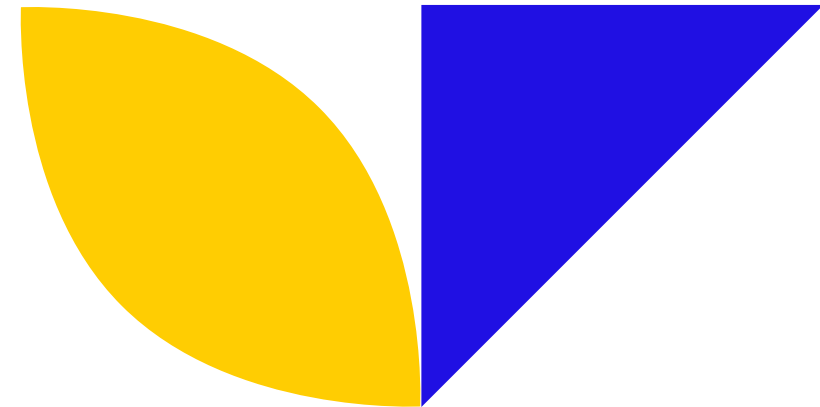
Índice

Introducción y marco general · Tendencias tecnológicas emergentes

Inteligencia Artificial · Estado del arte y disciplinas profesionales

Blockchain y Web3 · Más allá de la especulación

Realidad extendida · RV, RA y dispositivos espaciales



01

Introducción y marco general. Tendencias tecnológicas emergentes

Cómo se identifican, cómo se evalúan,
cómo se incorporan



Introducción y marco general

Definición operativa de tendencia tecnológica emergente

Tres rasgos simultáneos · sin ellos no hay tendencia, hay fenómeno aislado.

Los tres rasgos

- Novedad relativa · usabilidad accesible por primera vez fuera del laboratorio.
- Crecimiento no lineal · adopción, inversión, publicaciones académicas o citas crecen exponencialmente.
- Impacto sistémico potencial · capacidad para reconfigurar industrias, no sólo crear productos.

Falsos positivos comunes

- Producto exitoso de un único actor (ej. una app viral concreta).
- Hype mediático sin base técnica madura (ej. blockchain en 2017 para usos no financieros).
- Tecnología real pero confinada a un nicho (ej. impresión 3D fuera de prototipado).

Marco de referencia académico

- Rotolo, Hicks, Martin (2015) Research Policy · 18 papers revisados, 5 atributos clave.
- Convergencia con Diffusion of Innovations (Rogers, 1962) y curva del Hype (Gartner).

Introducción y marco general

Cuatro fuerzas que producen una tendencia tecnológica

Modelo causal · cuando concurren las cuatro, hay tendencia.

1

Avance técnico habilitador

- Un descubrimiento o ingeniería que rompe el estado del arte.
- Ejemplos: Transformer (2017), AlphaFold (2020), litografía EUV (2017).

2

Reducción drástica de coste

- El coste por unidad de capacidad cae órdenes de magnitud.
- Ej. coste por token GPT-4-class · de 60 USD/1M tokens (2023) a 0,15 USD (2025).

3

Caso de uso visible y medible

- ChatGPT como momento iPhone — interfaz que hace tangible la capacidad.
- Sin caso de uso visible, la tecnología se queda en publicaciones académicas.

4

Capital y atención sostenida

- Inversión privada, despliegue empresarial, regulación naciente.
- AI: inversión privada >100B USD anuales 2024 (Stanford AI Index).

Artificial Intelligence Index Report 2025



Introducción y marco general

Hype Cycle de Gartner

Modelo de atención mediática vs tiempo · útil como diagnóstico, no como predicción.

Cinco fases canónicas

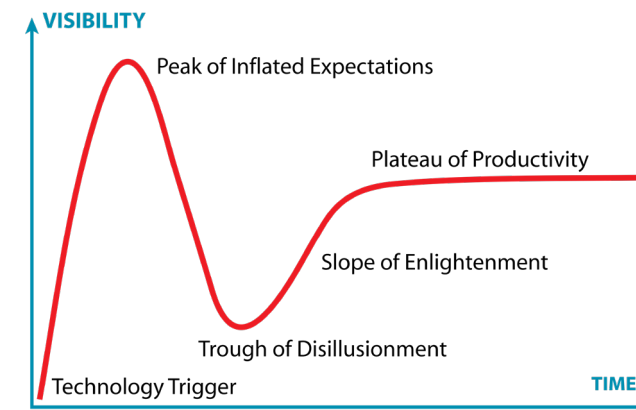
- Disparador de innovación · paper, demo, primera aplicación visible.
- Pico de expectativas infladas · atención masiva, inversión sin discriminar.
- Abismo de desilusión · fracasos públicos visibles, reflujo de capital.
- Rampa de iluminación · casos de uso reales sobreviven · disciplina técnica.
- Meseta de productividad · adopción sostenida, herramientas maduras.

Lectura crítica

- No es predictivo · una tecnología puede saltar fases o quedarse atascada.
- Útil como diagnóstico · ubicar dónde está hoy informa decisiones de inversión.
- Cada Hype Cycle anual sitúa decenas de tecnologías relativas, no absolutas.

Posicionamiento aproximado actual (2025)

- IA generativa · descenso desde el pico hacia el abismo.
- IA agéntica · pico de expectativas infladas.
- Blockchain (uso empresarial) · rampa de iluminación tras el abismo de 2022.
- Computación cuántica · todavía en disparador / pre-pico.



Introducción y marco general

Curva de difusión de Rogers y el chasm de Moore

Modelo de adopción · 5 segmentos · el salto que mata productos.

Cinco categorías de adoptantes (Rogers, 1962)

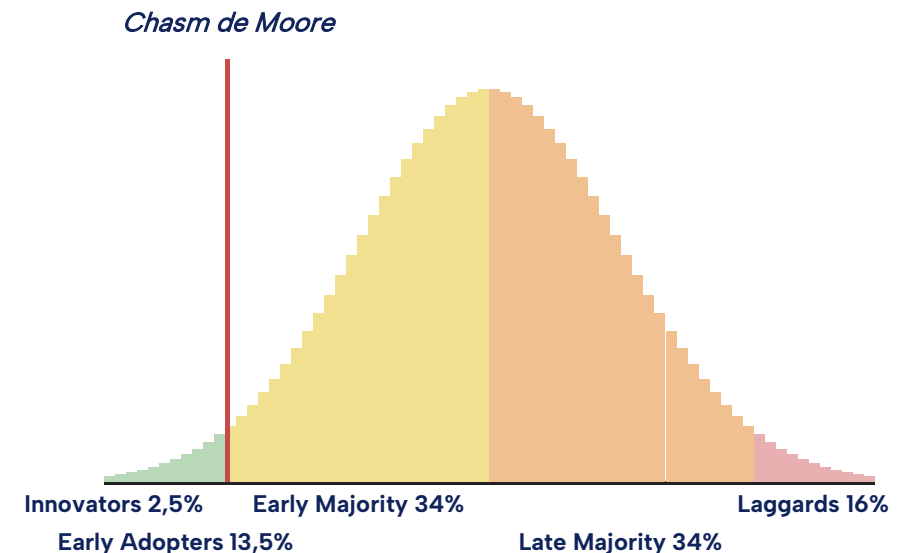
- Innovators · 2,5% · técnicos pioneros, exploradores.
- Early adopters · 13,5% · líderes de opinión, visionarios.
- Early majority · 34% · pragmáticos, exigen referencias.
- Late majority · 34% · escépticos, adoptan tarde por presión.
- Laggards · 16% · tradicionales, sólo si no queda alternativa.

El abismo de Moore (1991)

- Salto crítico entre Early adopters y Early majority (acumulado 16%).
- Cambia el público comprador: visionarios → pragmáticos.
- Donde mueren la mayoría de productos tech con tracción inicial.

Implicación práctica

- Conocer en qué segmento está vuestra empresa frente a una nueva tecnología.
- Adoptar como Early majority requiere referencias, integración estable, soporte.



Introducción y marco general

Las cuatro palancas de adopción · qué acelera y qué frena

Marco operativo para directores tecnológicos.

01

Capacidad técnica acumulada

- Personas con habilidades, código de calidad, datos limpios, infraestructura cloud-native.
- Aceleradora si está · gravísimo cuello de botella si falta.

02

Apetito de riesgo y gobernanza

- Capacidad organizativa de tolerar fallos controlados, marcos de aprobación claros.
- Aceleradora · permite pilotos rápidos · frena · si exige aprobaciones de 6 meses.

03

Alineamiento estratégico

- La tecnología responde a una pregunta de negocio real, no es solución buscando problema.
- Pregunta clave: ¿qué KPI específico mejora?.

04

Capacidad de cambio organizativo

- Procesos que se adaptan · personas que aceptan formación · sponsors ejecutivos.
- El factor más subestimado · suele ser el cuello de botella real.

Introducción y marco general

Marco regulatorio europeo · arquitectura completa

No son normas aisladas · forma un sistema coherente.

Capa horizontal · datos y privacidad

- GDPR (2016/679, vigente 2018) · base · datos personales.
- Data Act (2023/2854, vigente sept 2025) · acceso a datos generados por dispositivos IoT.
- Data Governance Act · intermediación y reutilización de datos públicos.

Capa vertical · sectores específicos

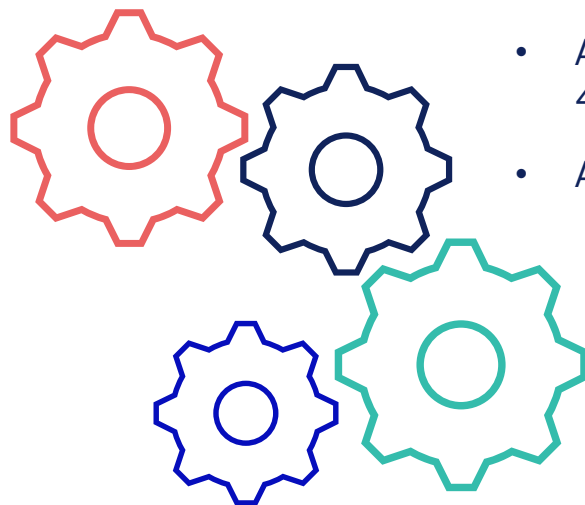
- MiCA (2023/1114) · cripto-activos · vigente 30 dic 2024.
- Regulación contenidos · DSA, DMA · plataformas digitales.

Capa vertical · IA

- AI Act (2024/1689, escalonado 2025–2027) · 4 niveles de riesgo · obligaciones por nivel.
- Aplica a sistemas, no a tecnología abstracta.

Capa vertical · ciberseguridad

- NIS2 (2022/2555, vigente oct 2024) · OES y DSP, sectores críticos.
- DORA (2022/2554, vigente ene 2025) · resiliencia financiera operativa.
- CRA (Cyber Resilience Act, 2024/2847) · seguridad de productos con elementos digitales.



Introducción y marco general

Soberanía digital europea · Chips Act y ecosistema cloud

Por qué el debate es estratégico, no ideológico.

Diagnóstico de partida

- EU produce ~9% de chips mundiales (era 24% en 1990) · objetivo Chips Act: 20% en 2030.
- Hyperscalers (AWS, Azure, GCP) controlan >70% del mercado cloud europeo.
- Modelos frontier de IA dominados por USA y China · ningún europeo en top 5 en 2025.

Respuestas regulatorias y de inversión

- EU Chips Act (2023) · 43.000 M€ públicos · ASML, Intel Magdeburg, TSMC Dresden.
- GAIA-X · cloud federado europeo · adopción modesta · debate sobre eficacia.
- EuroHPC · supercomputación pública · LUMI, Leonardo, MareNostrum 5.
- AI Factories · 7 instalaciones EU para entrenar modelos europeos · 2024-2026.

Dilemas estratégicos

- Soberanía total imposible · ¿qué nivel mínimo asegurar?
- Coste de duplicar capacidad existente vs depender de aliados.
- Aspectos críticos: silicio para defensa, modelos IA en sectores soberanos, infraestructura cloud para sanidad y administración.



Introducción y marco general

Sostenibilidad como límite estructural

No es ESG cosmético · es restricción física real para escalar.



Datos del problema

- Datacenters globales · 1-2% consumo eléctrico mundial (IEA 2024) · proyección 4-6% en 2030.
- Una consulta a GPT-4: ~10× la energía de una búsqueda Google clásica.
- Entrenamiento GPT-4 · ~50.000 MWh · equivalente a 5.000 hogares un año.



Respuestas en marcha

- Microsoft + Three Mile Island · reactivación reactor nuclear para datacenters AI (2024).
- Amazon · 24/7 carbon-free energy · acuerdos PPA con eólica y solar.
- Google DeepMind · TPUs hasta 4× más eficientes que GPUs comparables.
- EU CSRD (2022/2464) · reporting obligatorio de huella climática.



Implicaciones para arquitectos

- Decidir CUÁNDO usar IA generativa · no en todas las consultas.
- Foveated compute · usar modelos pequeños cuando bastan.
- Localización geográfica · mover entrenamiento a regiones con energía baja en carbono.
- Métricas de sostenibilidad como KPIs operativos, no cosmética ESG.

Introducción y marco general

El nuevo perfil del ingeniero informático

Lo que el rol exige hoy frente a hace 5 años.

Habilidades técnicas en alza

- Ingeniería de prompts y de contexto · saber cuándo y cómo usar LLMs en pipelines.
- Diseño de sistemas IA-native · RAG, agentes, evaluación, observabilidad.
- Seguridad reforzada · OWASP Top 10 LLM, supply chain, SBOM.
- Eficiencia y coste · medir y optimizar tokens, GPU horas, latencia.

Habilidades técnicas que pierden centralidad

- Programación de UI rutinaria · cubierta por copilotos.
- Boilerplate de backend CRUD · generado y refactorizado por IA.
- Ciertos roles QA manual · automatizables con agentes.

Habilidades blandas que ganan peso

- Pensamiento crítico para evaluar outputs probabilísticos.
- Comunicación con stakeholders no técnicos sobre limitaciones de IA.
- Capacidad de aprender continuamente · stack cambia cada 12-18 meses.
- Ética y responsabilidad profesional.

Convergencia disciplinar

- Software + ML + cyber + cumplimiento normativo · fronteras desdibujadas.
- Roles emergentes: AI engineer, prompt architect, AI governance officer.

Introducción y marco general

Mapa de los ocho vectores tecnológicos de la acción formativa

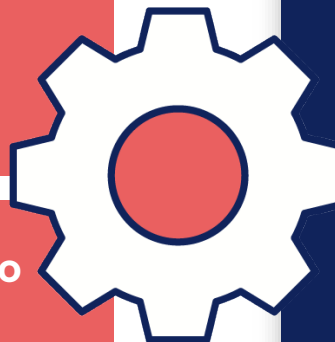
Coherencia de la primera y segunda mitad · 5 horas cada una.

Primera mitad · 5 horas · enfoque profundo

- Unidad 1 · Esta introducción.
- Unidad 2 · Inteligencia Artificial.
- Unidad 3 · Blockchain.
- Unidad 4 · Realidad virtual y aumentada.

Segunda mitad · 5 horas · operacional y crítico

- Unidad 5 · Cloud y Big Data.
- Unidad 6 · IoT.
- Unidad 7 · Ciberseguridad (incluye ciber + IA).
- Unidad 8 · Computación cuántica + cierre.



Hilo conductor · cuatro preguntas en cada bloque

- ¿Qué es y qué no es? (definición operativa).
- ¿Cómo funciona técnicamente? (lo justo para decidir).
- ¿Dónde está hoy? (Hype Cycle + adopción real).
- ¿Qué hago el lunes? (acciones para vuestro rol).

Introducción y marco general

Métodos de evaluación de tendencias para profesionales

Filtros prácticos para distinguir señal de ruido.

Análisis cualitativo

- PESTEL adaptado a tech
· Político · Económico · Social · Tecnológico · Ecológico · Legal.
- SWOT específico · oportunidad para nuestra organización vs competencia.
- Análisis de stakeholders
· ¿quién gana, quién pierde con esta tecnología?

Análisis cuantitativo

- Bibliometría · papers/citaciones por año (Google Scholar, Web of Science).
- Inversión privada · CB Insights, PitchBook, Crunchbase.
- Patentes · Espacenet, USPTO · indicador adelantado de comercialización.
- Adopción · informes Gartner, IDC, Forrester · benchmarks sectoriales.

Filtros de criterio profesional

- Wardley Mapping · ubicar componentes en evolution-stages.
- Test de los "cinco años" · ¿esta tecnología existirá en 5 años? ¿En qué forma?
- Test del "compromiso reversible" · ¿podemos abandonarla si falla?

Fuentes recomendadas

- ACM Communications · IEEE Spectrum · arXiv (revisado).
- MIT Technology Review · Nature/Science (selectos).
- Anthropic Research Blog · Google DeepMind Blog · OpenAI Research.

Introducción y marco general

Convergencia tecnológica como acelerador estructural

El efecto multiplicador entre vectores.

Convergencias en marcha · ejemplos verificados

- IA + biotecnología · AlphaFold 3 · proteínas, ácidos nucleicos · DeepMind 2024.
- IA + robótica · Figure 02, Tesla Optimus · agentes embodied · 2024-2025.
- IA + cuántica · simulación química con LLM + simulador cuántico · IBM, Google.
- Blockchain + IoT · supply chain con sensores (MediLedger, IBM Food Trust).
- Cloud + Big Data + IA · stack productivo de cualquier empresa data-driven.

Por qué importa para vuestra empresa

- Las tendencias no se evalúan aisladas · se evalúan en sistema.
- Una empresa que adopta IA sin Big Data limpio NO obtiene retorno.
- Una empresa con IoT sin ciberseguridad multiplica superficie de ataque.
- Una iniciativa blockchain sin caso B2B real es desperdicio de capital.

Implicación organizativa

Equipos cross-funcionales · no silos por tecnología.

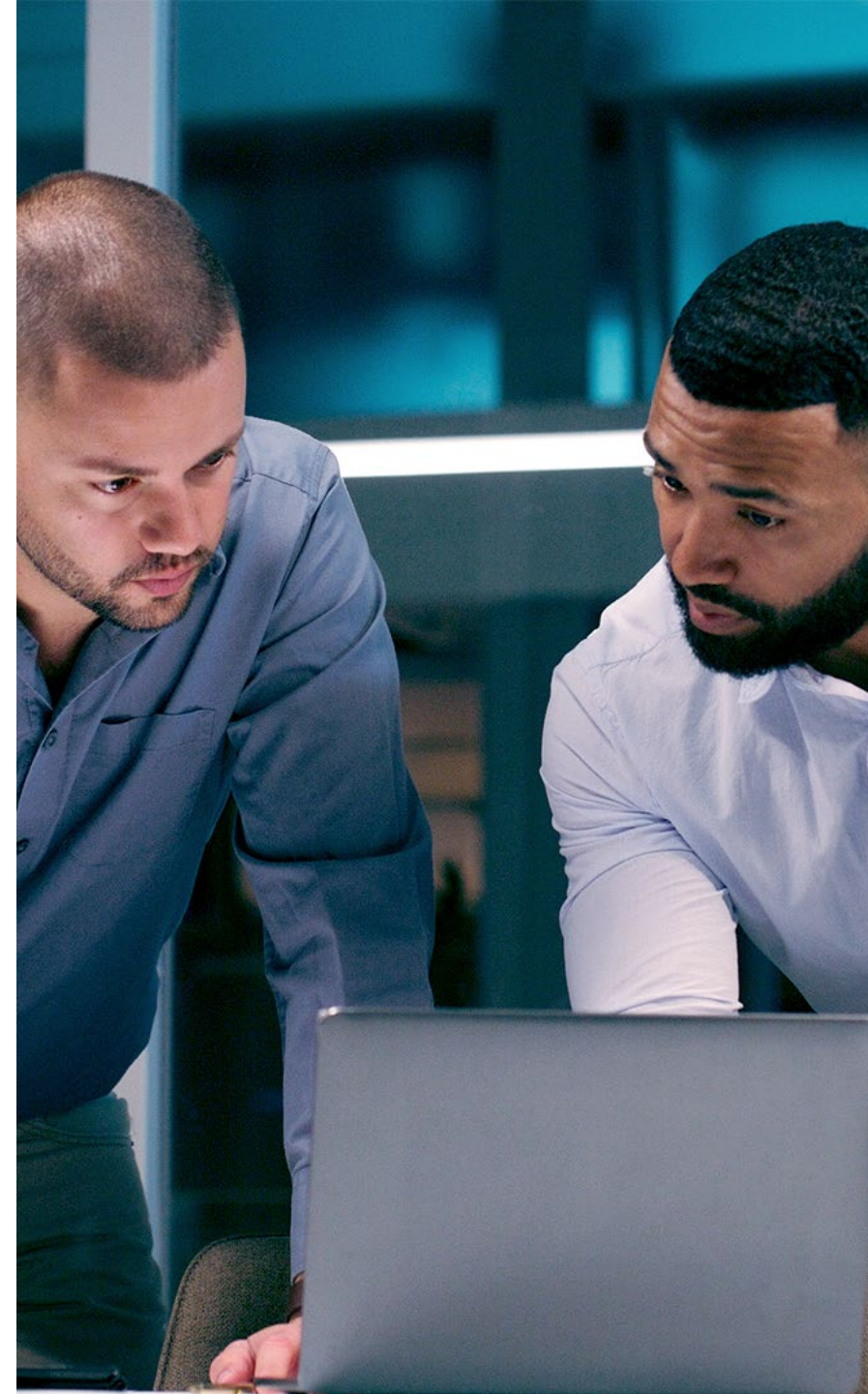
Roadmaps integrados · no roadmap de IA + roadmap de cloud + roadmap de seguridad separados.

Competencias híbridas en plantilla · esto es el cambio profesional más profundo.

Introducción y marco general

Cierre · qué nos llevamos a la unidad 2

- 01** Una tendencia es novedad relativa + crecimiento no lineal + impacto sistémico · sin las tres a la vez es fenómeno aislado.
- 02** Hype Cycle + Curva de Rogers son los dos diagnósticos que merecen estar en cualquier comité tecnológico.
- 03** El marco europeo (GDPR, AI Act, NIS2, DORA, MiCA, CRA) es denso pero coherente · navegar el sistema da certidumbre operativa.
- 04** Soberanía digital y sostenibilidad son restricciones físicas y geopolíticas reales · no slogans.
- 05** El nuevo perfil profesional combina software + ML + cyber + cumplimiento · convergencia disciplinar.
- 06** La convergencia entre vectores es lo que multiplica · evaluar tendencias en sistema, no aisladas.



02

Inteligencia Artificial · estado del arte y disciplinas profesionales

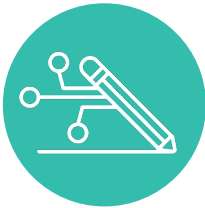
Fundamentos · LLMs · agentes · MLOps ·
seguridad · regulación · taller aplicado



Inteligencia Artificial · estado del arte y disciplinas profesionales

Definición operativa de Inteligencia Artificial

No discutimos qué es "inteligencia"; discutimos qué problemas resuelven los sistemas.



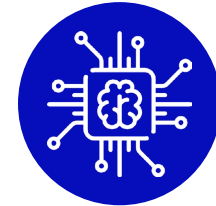
Definición de trabajo

- Conjunto de técnicas que permiten a una máquina realizar tareas que, ejecutadas por un humano, atribuiríamos a inteligencia: percibir, razonar, aprender, decidir, comunicarse.
- Operativa, no filosófica. La pregunta "¿es realmente inteligente?" no produce decisiones de negocio.



Cuatro capacidades que definen la IA moderna

- Aprender de ejemplos · Machine Learning clásico (1990s–presente, vigente).
- Generalizar · Deep Learning, AlexNet 2012 abrió la era moderna.
- Generar contenido original · IA generativa, Transformers 2017 → ChatGPT 2022.
- Actuar con autonomía · IA agéntica, frontera 2024–presente.



Tres palancas del salto reciente

- Cómputo: NVIDIA H100/B200, TPU v5/v6, miles de millones de FLOPs accesibles.
- Datos: corpus masivos (Common Crawl, datasets curados, datos sintéticos).
- Arquitectura: Transformer escalable (Vaswani et al. 2017).

Inteligencia Artificial · estado del arte y disciplinas profesionales

Línea del tiempo de la IA · hitos seleccionados

70 años de campo, 5 años de cambio drástico estructural.

2020–presente

Era generativa y agéntica

- Nov 2022 · ChatGPT — adopción masiva, 100 M usuarios en 2 meses.
- 2023 · GPT-4, Claude 1, Llama 2, Mistral 7B; explosión de open weights.
- 2024 · Multimodalidad nativa (GPT-4o, Gemini), modelos de razonamiento (o1, R1).
- Reciente · Era agéntica — Devin, Computer Use, agentes con tool calling masivo.

2000–2020

Era estadística y profunda

- 2006 · Hinton et al. revitalizan deep learning con deep belief networks.
- 2012 · AlexNet (Krizhevsky, Sutskever, Hinton) gana ImageNet con CNN — "momento Big Bang".
- 2014 · Goodfellow propone GANs (Generative Adversarial Networks).
- 2017 · Vaswani et al. publican Transformer en NeurIPS.
- 2018–2020 · BERT, GPT-2, GPT-3 — escalado de Transformers a tamaños inéditos.

1950–2000

Periodo simbólico

- 1950 · Turing publica "Computing Machinery and Intelligence" (Mind).
- 1956 · Conferencia de Dartmouth — McCarthy acuña "Artificial Intelligence".
- 1980s · Sistemas expertos (XCON, MYCIN); primer boom comercial.
- 1990s · "AI Winter" — recortes de financiación; trabajo silencioso en SVMs.

Inteligencia Artificial · estado del arte y disciplinas profesionales

Línea del tiempo de la IA · hitos seleccionados

70 años de campo, 5 años de cambio drástico estructural.

Aprendizaje supervisado

- Datos etiquetados (input → output esperado); función objetivo: minimizar pérdida sobre etiquetas.
- Ejemplos: clasificación de fraude, scoring de crédito, detección de defectos en visión.
- Limitante: etiquetar es caro; calidad de etiquetas determina calidad del modelo.

Aprendizaje por refuerzo

- Agente · entorno · acciones · recompensas; función objetivo: maximizar recompensa acumulada.
- Hitos: TD-Gammon (1992), AlphaGo (2016), AlphaZero (2017), AlphaStar (2019).
- Aplicación contemporánea crítica: RLHF — Reinforcement Learning from Human Feedback (Christiano et al. 2017).

Aprendizaje no supervisado

- Sin etiquetas; función objetivo: encontrar estructura (clusters, densidad, factores latentes).
- Técnicas: k-means, GMM, PCA, autoencoders, contrastive learning.
- Aplicación menos visible pero crítica: pre-entrenamiento de LLMs es no supervisado (next-token prediction).

Cómo se combinan en LLMs modernos

- Paso 1 · Pre-entrenamiento no supervisado con corpus masivo (next-token).
- Paso 2 · Fine-tuning supervisado con instrucciones humanas curadas (SFT).
- Paso 3 · RLHF o DPO (Direct Preference Optimization) para alineamiento con preferencias.

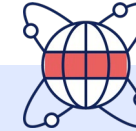
Inteligencia Artificial · estado del arte y disciplinas profesionales

ML clásico vs Deep Learning · cuándo cada uno



ML clásico (gradient boosting, regresión, árboles)

- Datos estructurados, tabulares (filas y columnas).
- Algoritmos interpretables (SHAP, feature importance).
- Funciona con miles de ejemplos.
- Coste de inferencia bajo, latencia ms.
- Estado del arte en banca, scoring, fraude tabular.
- Frameworks: XGBoost, LightGBM, CatBoost, scikit-learn.
- Sigue siendo la mayoría de IA en producción 2024.



Deep Learning (CNNs, Transformers)

- Datos no estructurados (texto, imagen, audio, vídeo).
- Caja negra; explicabilidad limitada (atribución, gradient maps).
- Necesita millones de ejemplos y GPUs.
- Coste de inferencia alto, latencia 100ms–segundos.
- Estado del arte en visión, NLP, generación, biología (AlphaFold).
- Frameworks: PyTorch (dominante), JAX, TensorFlow (legacy).
- Crece pero NO sustituye al ML clásico universalmente.

Inteligencia Artificial · estado del arte y disciplinas profesionales

Arquitectura Transformer · la pieza clave

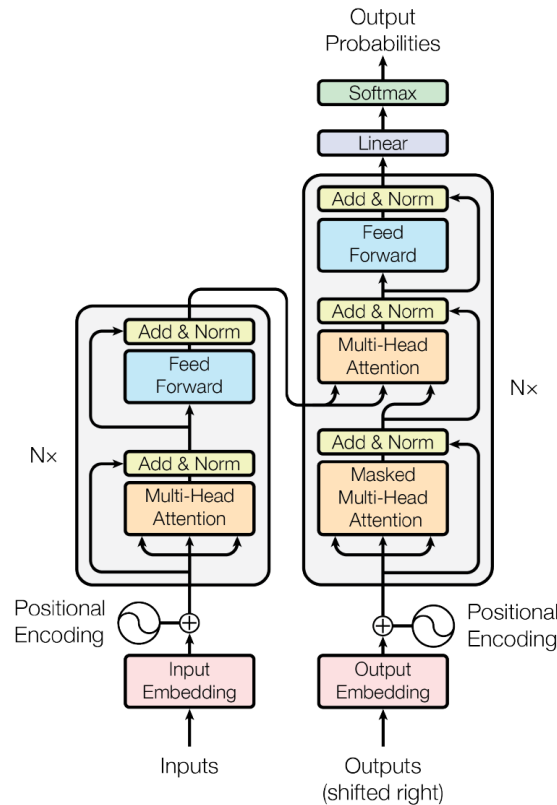
Atención multi-cabeza, paralelización, generalidad.

La idea central

- Procesar todos los tokens en paralelo, no secuencialmente como RNNs/LSTMs.
- Cada token "presta atención" a todos los demás vía mecanismo de scaled dot-product attention.
- Capas de self-attention + feed-forward apiladas en bloques residuales.

Tres consecuencias estructurales

- Paralelización · entrenamiento eficiente en GPUs/TPUs (a diferencia de RNNs secuenciales).
- Scaling laws · más parámetros + más datos = mejor desempeño de forma predecible (Kaplan et al. 2020).
- Generalidad · misma arquitectura aplicada a texto, imagen (ViT), audio (Whisper), código, biología (AlphaFold), proteínas.



Variantes principales

- Encoder-only · BERT, RoBERTa — clasificación, embeddings.
- Decoder-only · GPT, Llama, Claude — generación autoregresiva (dominante en 2024).
- Encoder-decoder · T5, original Transformer — traducción, summarization.

Optimizaciones recientes

- Flash Attention (Dao et al. 2022) — memoria y velocidad.
- Mixture of Experts (Mixtral, DeepSeek) — más parámetros, mismo cómputo activo.
- Grouped-query attention, sliding window — eficiencia en contextos largos.

Inteligencia Artificial · estado del arte y disciplinas profesionales

Cómo se entrena un LLM moderno · tres etapas

01

Etapas 1 · Pre-training

- Objetivo: predecir siguiente token sobre corpus masivo (next-token prediction).
- Datos: trillones de tokens (Common Crawl filtrado, libros, código, papers académicos).
- Cómputo: cluster de miles de GPUs durante meses; coste ~\$50-200M para frontera.
- Resultado: modelo que sabe MUCHO pero no sigue instrucciones.

02

Etapas 2 · Fine-tuning supervisado (SFT)

- Objetivo: enseñar al modelo a seguir instrucciones humanas.
- Datos: miles a millones de ejemplos curados por humanos (instruction tuning).
- Cómputo: días en cluster pequeño; mucho más barato que pre-training.
- Resultado: modelo que sigue instrucciones razonablemente.

03

Etapas 3 · Alineamiento por preferencias (RLHF / DPO)

- Objetivo: alinear el modelo con preferencias humanas matizadas (utilidad, seguridad, no toxicidad).
- Datos: cientos de miles a millones de comparaciones "respuesta A vs B".
- RLHF · entrenar reward model; usar PPO. DPO · alternativa que evita el reward model.
- Constitutional AI (Anthropic) · usar IA para evaluar respuestas según principios escritos.



Quién puede hacer cada etapa

- Etapa 1 frontier: ~10 organizaciones globales (OpenAI, Anthropic, Google, Meta, xAI, Mistral, Alibaba, DeepSeek, ByteDance, NVIDIA).
- Etapas 2 y 3: cualquier organización con datos curados y unos miles de € en cómputo.

Capacidades y límites de los LLMs frontier



Capacidades sólidas (estado del arte)

- Comprensión y generación de lenguaje natural en docenas de idiomas.
- Generación y revisión de código (productividad medida +20–40 % con uso correcto).
- Extracción de información de documentos no estructurados (contratos, facturas, informes).
- Resumen, traducción, paráfrasis con preservación de matices.
- Razonamiento multipaso explícito (modelos o1, R1, Claude 3.5 Sonnet con thinking).
- Multimodalidad: texto + imagen + audio + vídeo + código en un mismo modelo.



Límites y modos de fallo

- Alucinaciones · generación confiada de información falsa.
- Razonamiento de muchos pasos sin estructura: precisión decae exponencialmente.
- Matemáticas exactas: el modelo predice tokens, no calcula. Solución: tool use.
- Fecha de corte de entrenamiento: información posterior requiere búsqueda externa o RAG.
- Sesgo y representación: refleja sesgos del corpus de entrenamiento.
- Contexto limitado (128k–2M tokens según modelo): no recuerda conversaciones largas sin RAG.

Inteligencia Artificial · estado del arte y disciplinas profesionales

Panorama de modelos cerrados (frontier)

Arquitectura comercial dominante; ofrecidos vía API.

OpenAI

- GPT-4o · multimodal frontier (texto, imagen, audio, vídeo).
- o-series (o1, o3) · modelos de razonamiento explícito; orden de magnitud más caros por respuesta.
- Asociación estratégica con Microsoft (Azure OpenAI Service).
- API: platform.openai.com.

Anthropic

- Claude 3.5 Sonnet, Claude 3 Opus, Haiku · alineamiento via Constitutional AI.
- Particularmente fuerte en código, razonamiento, instrucciones largas.
- Computer Use API · primer modelo comercial que controla un ordenador (GUI).
- API: anthropic.com.

Google DeepMind

- Gemini 1.5 Pro, 2.0 Flash · multimodal nativo (texto, imagen, audio, vídeo).
- Ventana de contexto hasta 2 M tokens (la mayor del mercado).
- Integración profunda con Google Workspace, Search, Vertex AI.
- API: ai.google.dev.

Otros relevantes

- xAI · Grok-2, Grok-3 (con razonamiento) — desarrollado por equipo de Elon Musk.
- Mistral · Le Chat (modelos cerrados aparte de los abiertos), opción europea.
- Cohere · Command R+ (foco enterprise), Aya (multilingüe).



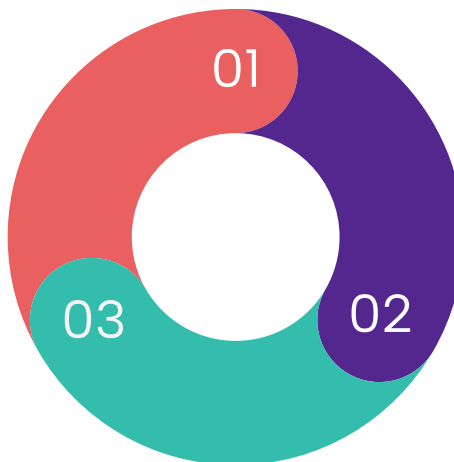
Inteligencia Artificial · estado del arte y disciplinas profesionales

Panorama de modelos abiertos (open weights)

Pesos descargables, despliegue propio, soberanía y coste a escala.

Familias principales

- Llama (Meta) · 8B, 70B, 405B; licencia permisiva con restricciones para empresas grandes.
- Mistral / Mixtral (Francia) · 7B, 8x7B (MoE), 8x22B; MoE eficiente.
- Qwen (Alibaba) · 0.5B–72B, multilingüe potente, fuerte en chino e inglés.
- DeepSeek (China) · DeepSeek-V3 (685B MoE, 37B activos), DeepSeek-R1 (razonamiento abierto)
- Phi (Microsoft) · modelos pequeños eficientes (3–14B) entrenados con datos sintéticos curados.
- Gemma (Google) · 2B–27B, alternativa Google a la familia Llama.



Cuándo elegir abierto

- Soberanía de datos · información sensible no debe salir a APIs externas.
- Volumen alto · coste por token vía API >> coste de operar modelo propio.
- Personalización profunda · fine-tuning específico de dominio.
- Cumplimiento AI Act alto riesgo · trazabilidad y auditoría completa.

Infraestructura común

- Hugging Face Hub · catálogo, transformers, accelerate, PEFT (LoRA, QLoRA).
- Servidores: vLLM, TGI (HuggingFace), TensorRT-LLM (NVIDIA), Ollama (local).
- Orquestación: LangChain, LlamaIndex, LangGraph (agentes).

Inteligencia Artificial · estado del arte y disciplinas profesionales

Multimodalidad nativa · texto, imagen, audio, vídeo, código

Los modelos foundation modernos procesan varias modalidades en un único pase.

Capacidades multimodales actuales

- Texto · todos los modelos foundation modernos.
- Imagen → texto · GPT-4o, Claude 3.5, Gemini, Llama 3.2 Vision.
- Imagen → imagen · Stable Diffusion 3, FLUX.1, DALL-E 3, Imagen 3.
- Audio · Whisper (transcripción), ElevenLabs (síntesis), Suno/Udio (música).
- Vídeo · Sora (OpenAI), Veo 2 (Google), Runway Gen-3, Kling, Hailuo MiniMax.
- 3D · DreamFusion, Stable3D, NVIDIA Edify 3D.

Casos prácticos en empresa

- Lectura automática de facturas, contratos, recibos con preservación de estructura.
- Análisis de llamadas (voz → texto → sentimiento, temas, acciones).
- Resumen de reuniones grabadas con identificación de hablantes.
- Generación de assets visuales para marketing (con C2PA para autoría).
- BI conversacional · pregunta en lenguaje natural sobre datos estructurados.

Limitaciones a tener presentes

- Generación de vídeo aún cara y limitada en duración (~10-60 s).
- Music gen con derechos de autor pendientes de jurisprudencia (Suno/Udio demandados).
- Generación de imágenes y deepfakes plantean retos de autenticidad (C2PA emergente).

Inteligencia Artificial · estado del arte y disciplinas profesionales

Modos de fallo · cómo se rompen los LLMs en producción

Diagnóstico de fallos típicos y mitigaciones probadas.

Alucinaciones (hallucination)

- El modelo genera información falsa con apariencia confiada (citas inventadas, APIs inexistentes, datos numéricos plausibles pero incorrectos).
- Mitigación: RAG con citas verificables, prompt instructions "si no sabes, di no sé", LLM-as-judge para detección automática.
- Frameworks: Ragas, TruLens, OpenAI Evals, Anthropic eval.

Drift de comportamiento

- El proveedor actualiza el modelo silenciosamente y tu output cambia.
- Mitigación: pin a versiones específicas (gpt-4o-2024-11-20, claude-3-5-sonnet-20241022); regression testing automático.

Costes runaway

- Bug en orquestación → agente itera 100 veces; factura mensual ×100.
- Mitigación: presupuestos por request y por usuario, alertas en tiempo real, max iterations en agentes.

Inyección de prompts (security)

- Atacante inserta instrucciones en datos que el LLM lee (correo, web, documento).
- Mitigación: separación contexto sistema/usuario, sanitización, guardrails I/O, principio de mínimo privilegio en tool use.
- Referencia: OWASP Top 10 for LLM Applications.

Sesgo y representación

- El modelo refleja sesgos del corpus; en decisiones sensibles puede discriminar grupos.
- Mitigación: tests de fairness desagregados por grupo, evaluación humana, supervisión.

Inteligencia Artificial · estado del arte y disciplinas profesionales

RAG · Retrieval-Augmented Generation

Patrón empresarial dominante para integrar conocimiento propio en LLMs.

El problema y la idea

- Problema: el LLM no conoce los documentos, datos y políticas de tu organización.
- Idea: recuperar contexto relevante en tiempo de consulta y pasarlo al modelo.
- No requiere reentrenar; actualización instantánea; trazabilidad por citas.

Arquitectura básica (4 pasos)

1. Indexar · documentos → chunks → embeddings → base vectorial.
2. Recuperar · query → embedding → top-K chunks por similitud.
3. Generar · LLM con prompt = pregunta + chunks + instrucciones.
4. Citar · respuesta enlaza a documentos fuente.

Patrones avanzados

- Hybrid search · combina denso (vectorial) + esparso (BM25) para precisión.
- Reranking · cross-encoder (Cohere Rerank, ColBERT) sobre top-K para reordenar.
- Contextual retrieval (Anthropic 2024) · enriquecer chunks con contexto del documento.
- GraphRAG (Microsoft 2024) · grafo de conocimiento + retrieval → mejor en preguntas globales.
- Agentic RAG · agente decide qué buscar y cuándo.

Stack común

- Bases vectoriales: Pinecone, Weaviate, Qdrant, Milvus, pgvector (PostgreSQL).
- Frameworks: LangChain, LlamaIndex, Haystack.
- Embeddings: OpenAI text-embedding-3, Cohere Embed v3, BGE (open), Voyage.

Inteligencia Artificial · estado del arte y disciplinas profesionales

Fine-tuning · cuándo tiene sentido y cuándo no

Está sobrevalorado. La regla práctica es "empieza por prompt, luego RAG, fine-tune solo si todo lo demás falla".

Cuándo fine-tunear (sí)

- Estilo o tono específico que prompt no captura consistentemente.
- Formato de salida estricto y repetitivo (JSON con esquema fijo).
- Dominio con vocabulario propio (legal español, ICD-10 médico, CAD ingeniería).
- Coste a escala alta · modelo pequeño fine-tuneado puede ser 10-100× más barato que API.

Técnicas modernas (PEFT)

- LoRA (Hu et al. 2021) · Low-Rank Adaptation; reentrena solo matrices pequeñas (~1 % de parámetros)
- QLoRA (Dettmers et al. 2023) · LoRA + cuantización 4-bit; permite fine-tune de 70B+ en 1-2 GPUs
- DPO (Direct Preference Optimization) · alineamiento sin reward model intermedio

Cuándo NO fine-tunear

- Para enseñar conocimiento puro · usa RAG (mucho más barato y actualizable).
- Si los datos cambian con frecuencia · reentrenar continuo es inviable.
- Si tienes pocos ejemplos · <500 ejemplos de calidad → sobreajuste, beneficio nulo.

Stack práctico

- Frameworks: HuggingFace TRL, axolotl, LLaMA-Factory, Unsloth (eficiencia extrema).
- Serving del modelo fine-tuneado: vLLM con LoRA adapters cargados dinámicamente.

Inteligencia Artificial · estado del arte y disciplinas profesionales

Matriz de decisión · prompt · RAG · fine-tuning · agente

Empezar por lo más simple. Cada nivel implica más coste e infraestructura.

Nivel 1 · Prompt engineering (siempre primero)

01

- Cuándo: tareas que requieren conocimiento general + razonamiento estándar.
- Coste: marginal (solo tokens). Tiempo: horas-días.
- Resuelve aprox. 60 % de casos empresariales.

02 Nivel 2 · RAG (si Nivel 1 no basta por conocimiento privado)

- Cuándo: el modelo necesita información que no tiene (documentos internos, datos cambiantes).
- Coste: base vectorial + indexación + reranking. Tiempo: días-semanas.
- Resuelve aprox. 30 % adicional.

Nivel 3 · Fine-tuning (solo si lo anterior no resuelve)

03

- Cuándo: estilo, formato, dominio o coste a escala alta justifica el esfuerzo.
- Coste: datos curados + cómputo + mantenimiento. Tiempo: semanas-meses.
- Resuelve el 10 % restante.

04 Nivel 4 · Agentes (ortogonal a los anteriores)

- Cuándo: la tarea requiere acciones (no solo respuestas) y descomposición autónoma.
- Coste: orquestación + observabilidad + costes de iteración. Tiempo: semanas-meses.
- Combinable con cualquier nivel anterior.

Inteligencia Artificial · estado del arte y disciplinas profesionales

Del chatbot al agente · cinco etapas evolutivas

Cada etapa amplía capacidad pero exige más infraestructura, observabilidad y guardrails.



Anatomía de un agente moderno

LLM + memoria + herramientas + bucle de razonamiento + guardrails.

Componentes

- Cerebro · LLM (frontier o fine-tuneado) que decide qué hacer.
- Memoria · corto plazo (contexto), largo plazo (RAG), estructurada (estado del agente).
- Herramientas · APIs propias, código, búsqueda web, file system, otros agentes.
- Bucle de razonamiento · ReAct loop (Reason → Act → Observe → repeat).
- Guardrails · presupuesto de tokens, límite de iteraciones, listas blancas/negras de acciones.

Diseño del bucle

1. Recibe objetivo y estado inicial.
2. Razona: ¿qué herramienta usar? ¿con qué argumentos?
3. Ejecuta tool call.
4. Observa resultado.
5. Itera hasta cumplir objetivo o alcanzar límite.

Anti-patrones a evitar

- Bucles infinitos · agente itera sin progreso (límite de iteraciones obligatorio).
- Costes desbocados · sin límite de tokens, una iteración mal puede multiplicar factura.
- Propagación de errores · un fallo en planificación rompe toda la cadena
- Confianza ciega entre agentes · prompt injection puede colarse de un agente a otro.

Inteligencia Artificial · estado del arte y disciplinas profesionales

Agentes IA en producción · casos verificables

Estado real (no marketing). La mayoría aún en piloto; algunos productivizados.

IDE y desarrollo software

- GitHub Copilot Workspace · agente que navega repo, propone cambios, abre PR.
- Cursor · IDE construido alrededor de agentes; rondas de funding 2024-2025 muy altas.
- Devin (Cognition) · agente especializado en software engineering.
- Claude Code (Anthropic) · agente CLI de Anthropic.

Investigación y análisis

- Perplexity Pro Search · investigación con citas verificables.
- OpenAI Deep Research · investigación profunda autónoma con informe final.
- Anthropic Research mode · alternativa equivalente.

Atención al cliente

- Intercom Fin · resuelve 30-50 % de tickets sin escalado humano (datos publicados).
- Salesforce Agentforce · agentes para Service Cloud y Sales Cloud.
- Zendesk AI Agents · oferta análoga.

Backoffice y procesos

- Agentes de procurement, expedientes, KYC en banca con human-in-the-loop.
- Adoption en retail, seguros, sector público con casos piloto crecientes.



Inteligencia Artificial · estado del arte y disciplinas profesionales

MLOps vs LLMOps · disciplinas distintas

LLMs introducen retos operativos nuevos que MLOps tradicional no cubre.

Diferencias clave

- Outputs no deterministas · evaluación requiere LLM-as-judge o humanos, no asserts.
- Prompts son código vivo · versionables, revisables, testeables (PromptOps).
- Modelos cambian sin que hagas nada · hyperscaler actualiza, tu output cambia.
- Costes no lineales · agentes pueden multiplicar coste 100× con un bug.
- Seguridad nueva · prompt injection, jailbreaks, fuga de datos vía contexto.

Tres pilares operativos

- Observabilidad · logs por request (tokens, latencia, calidad percibida, coste).
- Evaluación · suites automáticas + LLM-as-judge + evaluación humana periódica.
- Guardrails · filtros entrada/salida, policies, content moderation, PII detection.

Stack representativo

- Observabilidad: LangSmith, Helicone, Langfuse, Datadog LLM Observability, Arize Phoenix.
- Evaluación: Ragas (RAG), TruLens, OpenAI Evals, custom test suites.
- Guardrails: NeMo Guardrails (NVIDIA), Guardrails AI, Lakera Guard, AWS Bedrock Guardrails.

Inteligencia Artificial · estado del arte y disciplinas profesionales

Riesgos · sesgo, alucinación, privacidad, IP

Casos públicos verificables y mitigaciones conocidas.

Sesgo · casos documentados

- Amazon CV screening (2018) · sistema retirado por penalizar mujeres.
- COMPAS (sentencias, USA) · ProPublica (2016) documentó sesgo racial en predicciones de reincidencia.
- Mitigación: tests de fairness desagregados (AIF360, Fairlearn), datos de entrenamiento balanceados, auditorías independientes.

Privacidad · casos publicados

- Samsung (2023) · prohibición temporal de ChatGPT tras pago de código confidencial.
- Italia (Garante, 2023) · suspensión temporal de ChatGPT por RGPD.
- Mitigación: Azure OpenAI con cláusula no-reentrenamiento, anonimización previa, DLP.

Alucinaciones · casos públicos

- Mata v. Avianca (2023) · bufete sancionado por presentar sentencias inventadas por ChatGPT.
- Stack Overflow ban inicial (2022) · respuestas plausibles pero incorrectas.
- Mitigación: RAG con citas, LLM-as-judge, supervisión humana en decisiones críticas.

Propiedad intelectual · litigios abiertos

- NYT vs OpenAI (diciembre 2023) · primera demanda de prensa importante.
- Getty Images vs Stability AI (2023) · imágenes con copyright en entrenamiento.
- Suno y Udio demandados por RIAA (2024) · música generada con corpus posiblemente infringente.
- Mitigación: contratos con cláusula de indemnización IP (Microsoft, Google, OpenAI ofrecen).

Inteligencia Artificial · estado del arte y disciplinas profesionales

AI Act · cuatro niveles de riesgo · obligaciones

Reglamento (UE) 2024/1689 · aplicable escalonadamente hasta agosto 2027.



Inteligencia Artificial · estado del arte y disciplinas profesionales

AI Act · cronograma y plan de cumplimiento operativo

Cinco pasos para cualquier organización que use IA.

Cronograma escalonado

- Febrero 2025 · prácticas prohibidas vigentes; alfabetización en IA obligatoria para personal.
- Agosto 2025 · obligaciones para modelos GPAI (proveedores).
- Agosto 2026 · sistemas de alto riesgo NUEVOS deben cumplir plenamente.
- Agosto 2027 · sistemas de alto riesgo PREEXISTENTES en mercado deben cumplir plenamente.

Plan en 5 pasos

1. Inventario · ¿qué sistemas IA hay, dónde, quién los usa, para qué?
2. Clasificación · ubicar cada sistema en la pirámide de riesgo.
3. Análisis de huecos · qué falta para cumplir según nivel.
4. Plan con plazos · quién hace qué, cuándo, con qué presupuesto.
5. Formación mínima · personal que interactúa con IA debe estar formado (es obligación legal).

Preguntas para directivos

- ¿Hay un AI Officer o responsable claro en la organización?
- ¿Está documentado qué sistemas IA hay y bajo qué nivel?
- ¿Hay procesos de evaluación de conformidad para nuevos sistemas?
- ¿Las decisiones de alto impacto pasan por supervisión humana significativa?

Inteligencia Artificial · estado del arte y disciplinas profesionales

IA aplicada por sector · estado del despliegue real

Casos verificables, no demos. Distinción entre piloto y producción a escala.

Salud

- Diagnóstico por imagen · IDx-DR (FDA, retinopatía), Aidoc (PE en TC), Annalise.ai.
- Documentación clínica · Ambient AI (Nuance DAX, Abridge); ahorro reportado ~30 % tiempo.
- Drug discovery · AlphaFold 3 (Abramson et al. Nature 2024), Insilico Medicine en ensayos clínicos.

Banca y servicios financieros

- Detección de fraude · ML clásico + deep learning combinados; pérdidas reducidas 30–50 %.
- Asistentes virtuales · Erica de Bank of America, Aida de SEB; millones de usuarios activos.
- KYC/AML · automatización con LLMs + tool calls a registros oficiales.

Industria 4.0

- Mantenimiento predictivo · acoustic + visión + sensórica IoT con modelos foundation.
- Control de calidad por visión · Cognex, Zebra Aurora con vision foundation models.
- Cobots con learning · Universal Robots, ABB con instrucciones en lenguaje natural.

Desarrollo software (sector más transformado)

- GitHub Copilot · 1.8M+ developers pagantes; +20–40 % productividad medida.
- Cursor · IDE construido sobre agentes; valoración multi-billón en 2025.
- Devin, Claude Code, Codex · agentes especializados en SWE.

Inteligencia Artificial · estado del arte y disciplinas profesionales

Convergencia · IA + otras tecnologías

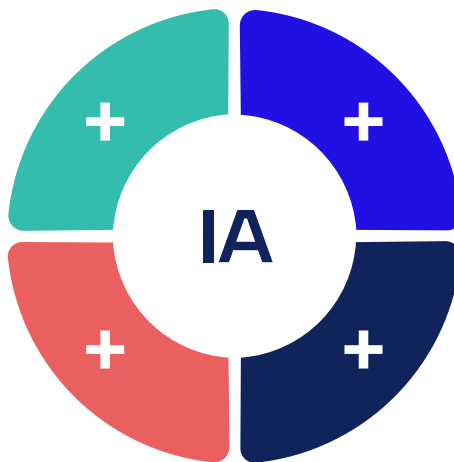
Los puntos de convergencia son donde se concentra el valor económico nuevo.

IA + Cloud + Datos

- La pila de la última década. Lakehouses + LLMs + RAG = patrón empresarial dominante.
- Snowflake Cortex, Databricks Foundation Models, AWS Bedrock, Azure AI Foundry.

IA + Biotecnología

- AlphaFold 3 (Abramson et al. 2024), AlphaProteo · diseño de proteínas a medida.
- Insilico Medicine · primer fármaco diseñado con IA en ensayo clínico fase 2.
- Recursion, Deep Genomics, BenevolentAI · pipelines IA-first.



IA + Robótica

- Foundation models aplicados a control: RT-1, RT-2, RT-X (Google DeepMind + 21 instituciones).
- Tesla Optimus, Figure 02, 1X NEO, Apptronik Apollo · robots humanoides comerciales emergentes.
- Open X-Embodiment dataset · ~1M episodios robot publicados.

IA + Hardware especializado

- NVIDIA Blackwell, Cerebras WSE-3, Google TPU v6, Meta MTIA, AWS Trainium.
- Edge AI · Apple Neural Engine, Google Edge TPU, Qualcomm AI; modelos pequeños eficientes.

Inteligencia Artificial · estado del arte y disciplinas profesionales

Recursos para profundizar · curated

Para mantenerse en la frontera. Selección minimalista, frecuencia mensual.

Newsletters de calidad

- Latent Space · latent.space — engineering perspective, signal alto
- Import AI · Jack Clark (Anthropic policy lead) — política y técnica
- AI Snake Oil · Narayanan + Kapoor (Princeton) — análisis crítico
- Ahead of AI · Sebastian Raschka — técnica con profundidad

Investigación primaria

- arXiv cs.AI / cs.LG · arxiv.org
- Papers With Code · paperswithcode.com — estado del arte por benchmark
- alphaXiv · alphaxiv.org — discusión sobre papers de arXiv
- AK on X (Hugging Face papers daily)

Reportes y benchmarks

- Stanford AI Index (anual) · aiindex.stanford.edu
- Hugging Face Open LLM Leaderboard · huggingface.co/spaces/HuggingFaceH4/open_llm_leaderboard
- Chatbot Arena (LMSYS) · chatarena.com
- Artificial Analysis · artificialanalysis.ai

Libros recomendados

- "AI Engineering" · Chip Huyen (2024) — el manual práctico más reciente
- "Designing Machine Learning Systems" · Chip Huyen (2022)
- "Hands-on Large Language Models" · Alammam, Grootendorst (2024)



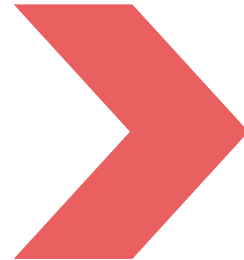
Inteligencia Artificial · estado del arte y disciplinas profesionales

Cierre del Bloque 2 · transición a Blockchain

Para mantenerse en la frontera. Selección minimalista, frecuencia mensual.

Qué nos llevamos

- Marco técnico: ML, Deep Learning, Transformers, scaling laws.
- Patrones empresariales: prompt → RAG → fine-tune → agentes.
- Disciplinas operativas: LLMOps, observabilidad, evaluación, guardrails.
- Marco regulatorio: AI Act, niveles de riesgo, plan de cumplimiento.
- Casos verificables por sector y referencias para profundizar.



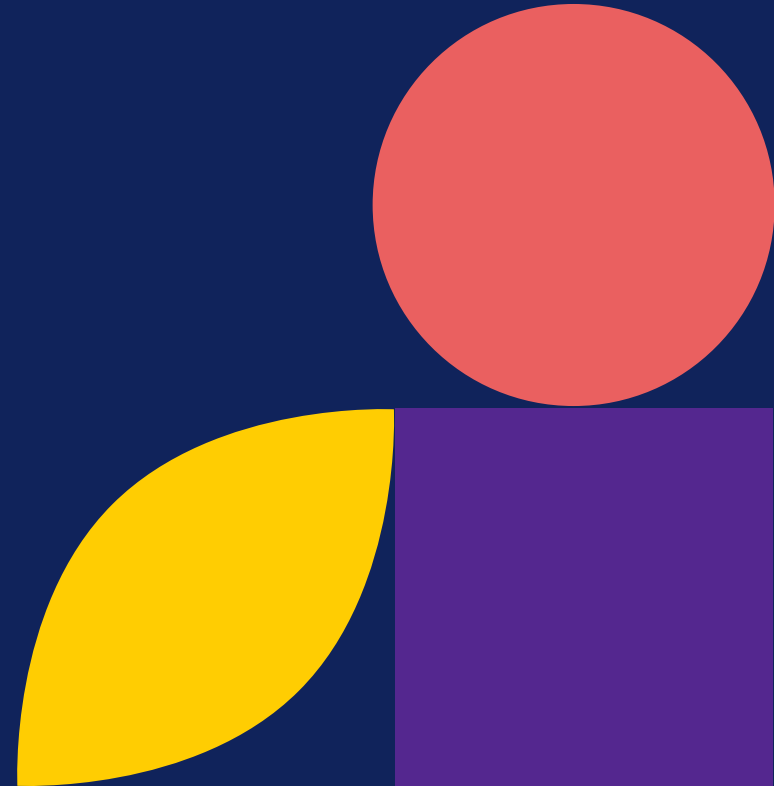
Próximo bloque

- Blockchain y Web3 · 1 hora · una tendencia que cruzó el chasm en aplicaciones empresariales.
- Pausa recomendada: 10-15 minutos antes de continuar.

03

Blockchain · más allá de la especulación

Criptografía aplicada · activos tokenizados
· infraestructura programable



Blockchain · más allá de la especulación

Qué es realmente una blockchain

Una definición precisa, sin marketing.



Definición

- Base de datos distribuida, replicada en muchos nodos, donde el orden y validez de las escrituras lo determina un protocolo de consenso, no un administrador.
- Cada bloque referencia criptográficamente al anterior — modificar un bloque pasado exige rehacer toda la cadena de hashes.



Propiedades emergentes (no inherentes)

1. Inmutabilidad práctica — cuesta más reescribir que el valor protegido.
2. Resistencia a la censura — depende del grado de descentralización real.
3. Auditabilidad pública — todo el histórico es verificable por cualquiera.
4. Liquidación atómica 24/7 — sin intermediarios de compensación.



Lo que NO es

- No es una base de datos rápida (es entre 100 y 100.000× más lenta que Postgres).
- No es "gratis" (cada escritura paga gas o consume cómputo del consenso).
- No es anónima por defecto — la mayoría de cadenas son seudónimas y trazables.

Blockchain · más allá de la especulación

De Nakamoto a la era de la tokenización

Diecisiete años de iteración.

Era 1 · Bitcoin

- Whitepaper de Nakamoto resuelve el problema de doble gasto sin autoridad central.
- Proof of Work como mecanismo de consenso · primer ejemplo de oro digital.
- Lecciones: una cadena con un solo activo nativo y reglas fijas.

2008–2014

2015–2020

2021–2024

Era actual

Era 2 · Programabilidad y Ethereum

- Ethereum introduce smart contracts Turing-completos — la blockchain como ordenador mundial.
- ICOs (2017) — burbuja de financiación · DeFi summer (2020) — primitivas financieras nativas.

Era 3 · Escalado y consolidación (2021–2024)

- Merge de Ethereum (sept 2022) — paso a Proof of Stake, –99,9% energía.
- Auge de capas 2 (rollups) · MiCA en UE · ETFs spot de BTC y ETH (2024).

Activos reales tokenizados

- Treasuries tokenizados (BlackRock BUIDL, Ondo OUSG) superan 5.000M USD.
- Stablecoins liquidan billones · infraestructura financiera real, no especulativa.



Blockchain · más allá de la especulación

Anatomía técnica · qué hay debajo del capó

Cuatro piezas · hashing y Merkle ilustrados a la derecha.

01

Hashing criptográfico

- SHA-256 (Bitcoin), Keccak-256 (Ethereum).
- Cada bloque incluye el hash del anterior.
- Cualquier cambio en el input cambia el output completo.

02

Árbol de Merkle

- Agrega N transacciones en un solo hash raíz.
- Light clients verifican inclusión con $\log(N)$ hashes.
- Permite SPV (Simplified Payment Verification) en móvil.

03

Firmas digitales

- ECDSA, EdDSA, BLS (agregación clave para PoS).
- Clave privada = posesión del activo.

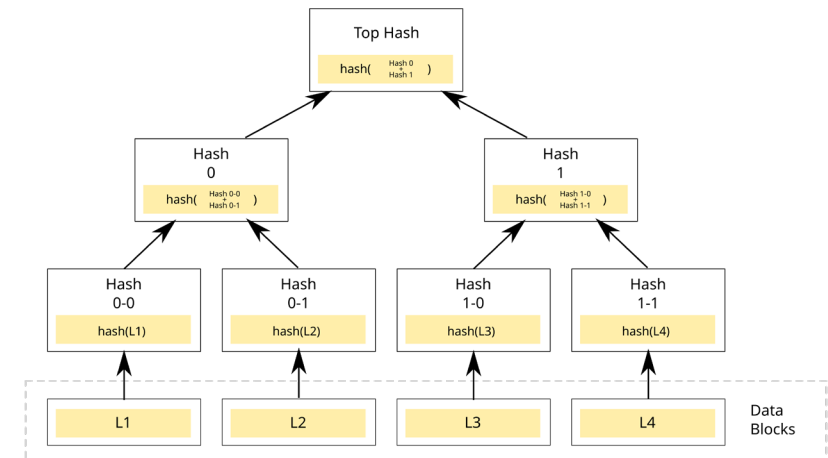
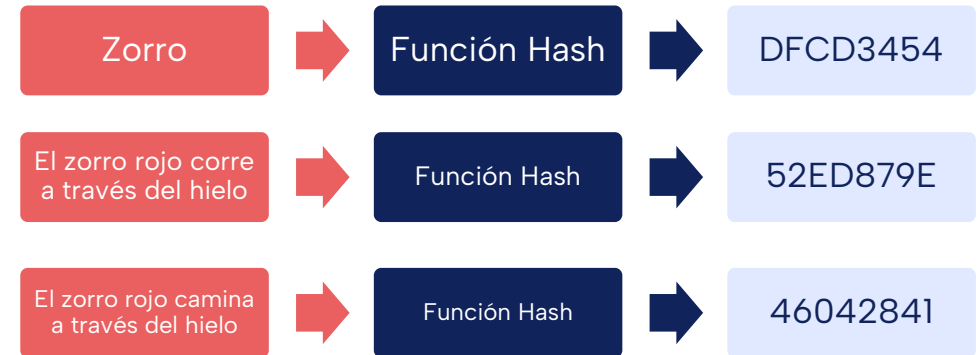
04

Consenso

- Tolerancia bizantina · 1/3 (PoS) o 1/2 (PoW).
- Lamport (1982) demostró el límite teórico.

Entrada

Valor Hash



Blockchain · más allá de la especulación

Proof of Work vs Proof of Stake · tradeoffs reales

Una definición precisa, sin marketing.

Proof of Work (Bitcoin)

- Consenso por gasto físico de energía.
- Seguridad: atacante necesita >51% hashrate (~150 GW).
- Coste energético: ~150 TWh/año.
- Resistencia máxima a censura.
- Sin slashing — pierde sólo electricidad.
- Ideal para activos de reserva.



Proof of Stake (Ethereum, Solana, Cosmos)

- Consenso por colateral económico bloqueado.
- Seguridad: atacante necesita >33% stake (~120B USD).
- Coste energético: 99,9% menos que PoW.
- Validators identificables (más censurable).
- Slashing destruye capital del atacante.
- Permite finalidad rápida y sharding.



Blockchain · más allá de la especulación

Públicas, permissionadas, híbridas · cuándo cada una

Decisión arquitectónica antes de cualquier piloto.

Pública sin permiso (Bitcoin, Ethereum, Solana)

- Cualquiera puede leer, escribir (pagando), validar y construir aplicaciones.
- Máxima resistencia a censura · máxima auditabilidad · máxima descentralización.
- Coste por escritura · privacidad limitada · gobernanza compleja.
- Caso de uso: dinero programable, activos tokenizados con liquidez global, identidad portable.

Permissionada (Hyperledger Fabric, R3 Corda, Quorum)

- Acceso restringido por consorcio · validators conocidos y vetted.
- Throughput alto · privacidad por canales · sin token nativo necesario.
- Caso de uso: trade finance entre bancos, supply chain entre socios, registros B2B.

Híbrida (rollups con secuenciador centralizado, sidechains)

- Ejecución privada o controlada, finalidad pública en una L1 abierta.
- Combina performance permissionada con auditabilidad pública.
- Tendencia dominante en empresa: Polygon CDK, Arbitrum Orbit, Optimism Superchain.

Blockchain · más allá de la especulación

Smart contracts · código como infraestructura

Lo que los hace útiles · y lo que los rompe.

Qué son

- Programas que se ejecutan determinísticamente en cada nodo de la red.
- Inmutables tras deployment (salvo proxies upgradeables) · transparentes (bytecode público).
- Componibles — un contrato puede llamar a otro · DeFi nace de esta primitiva.

Tooling moderno

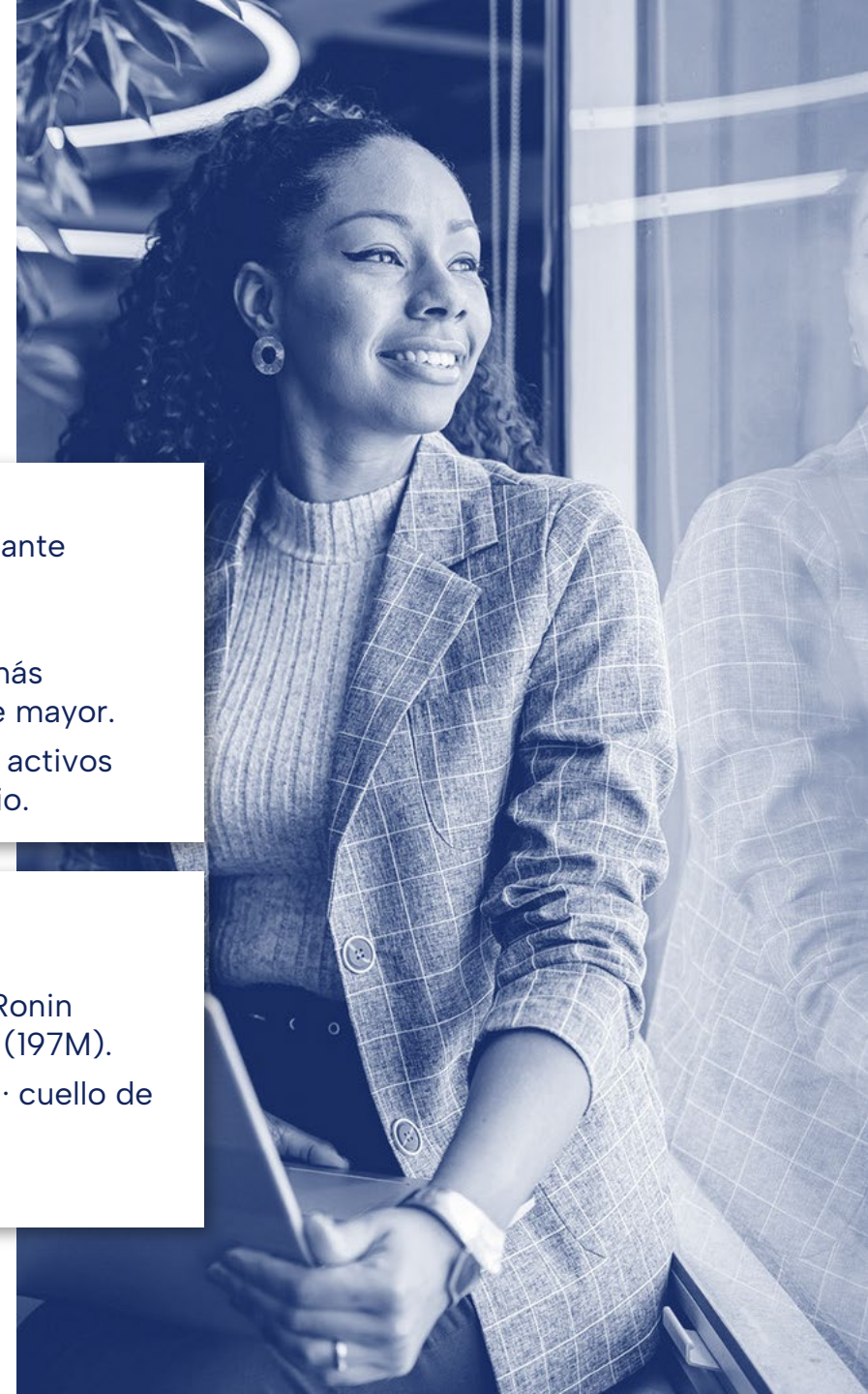
- Foundry (testing rápido en Solidity) · Hardhat · Anchor (Solana).
- Tenderly (debugging de tx) · OpenZeppelin Defender (operaciones).
- Halmos, Certora (verificación formal) · Slither (static analysis).

Lenguajes y entornos

- Solidity / EVM · ecosistema dominante (Ethereum + 20+ cadenas EVM-compatibles).
- Rust / SVM (Solana, Sui, Aptos) · más performante, curva de aprendizaje mayor.
- Move (Aptos, Sui) · diseñado para activos digitales seguros desde el principio.

Lo que falla

- Bugs lógicos cuestan millones — Ronin (625M), Wormhole (320M), Euler (197M).
- Auditoría sigue siendo bottleneck · cuello de botella humano.



Blockchain · más allá de la especulación

Capa 2 · cómo se escala una blockchain pública

Optimistic rollups vs ZK rollups · el debate central de 2025.

Por qué L2

- Ethereum L1 procesa ~15 tx/s · imposible para uso global.
- L2 ejecuta off-chain pero hereda seguridad de L1 publicando datos y proofs.
- Tras EIP-4844 (blobs, marzo 2024), coste por tx en L2 cae 10-100×.

Optimistic rollups (Arbitrum, Optimism, Base)

- Asumen optimísticamente que las tx son válidas · permiten 7 días para impugnar (fraud proof).
- Latencia de retiro alta · throughput muy alto · EVM-equivalentes (sin fricción para devs).
- TVL combinado: >40B USD a inicios de 2025.

ZK rollups (zkSync, Starknet, Polygon zkEVM, Linea)

- Cada lote de tx se acompaña de una prueba criptográfica de validez.
- Finalidad inmediata · privacidad nativa posible · más complejos técnicamente.
- Coste de prover en bajada continua — antes era prohibitivo, ahora competitivo.

Hacia dónde va

- App-chains: cada app grande tendrá su propio rollup · Coinbase Base, Sony Soneium, Worldcoin World Chain.
- Interoperabilidad entre L2s aún no resuelta · bridges siguen siendo el punto débil.

Blockchain · más allá de la especulación

Stablecoins, CBDCs y MiCA · el dinero programable

El segmento más importante de blockchain hoy.

Stablecoins · escala real

- Capitalización total: ~190B USD a inicios de 2025 (USDT 140B, USDC 40B).
- Volumen anual de liquidación: >10 billones USD — supera a Visa+Mastercard combinadas.
- Tipos: respaldadas por fiat (USDT, USDC), por cripto (DAI), algorítmicas (mayoría colapsadas tras Terra).

Casos de uso reales (no especulativos)

- Remesas: Filipinas, Argentina, Nigeria — coste 1-3% vs 6-9% Western Union.
- Liquidación B2B internacional · settlement instantáneo 24/7.
- Reserva de valor en economías con inflación alta (USDT en Argentina, Venezuela, Turquía).

CBDCs · monedas digitales de banco central

- Eurozona: digital euro en fase preparatoria (BCE, decisión 2025-2026).
- China e-CNY: piloto desde 2020, no usa blockchain pública (centralizado).
- 98% de bancos centrales investigan CBDC · 11 lanzados, 39 en piloto.

MiCA · marco europeo

- Markets in Crypto-Assets · vigente desde 30 dic 2024.
- Reservas 1:1 obligatorias · auditorías mensuales · whitepaper aprobado.
- Limita stablecoins no-EUR significativas a 1M tx/día — fricción real para USDT/USDC.

Blockchain · más allá de la especulación

Tokenización de activos reales · la nueva ola institucional

De la curiosidad a la asignación estratégica.

Qué se está tokenizando · lo serio

- Treasuries USA: BlackRock BUIDL (>500M), Ondo OUSG, Franklin Templeton FOBXX.
- Private credit: Centrifuge, Maple — yield 8–12% para inversores crypto.
- Real estate: Lofty (residential), RealT, RedSwan (commercial).
- Commodities tokenizados: oro (PAXG, XAUT), uranium (Madison).

Promesa real

- Liquidación instantánea T+0 vs T+2 tradicional · capital working más eficiente.
- Acceso 24/7 · fraccionalización · transparencia de holdings.
- Componibilidad: usar treasury tokenizado como colateral en DeFi.

Por qué ahora

- Tipos altos hacen atractivo el yield de bonos del Tesoro on-chain (vs 0% en stablecoin).
- Custodios institucionales (BNY Mellon, Anchorage) abren la puerta a balance corporativo.
- Frameworks regulatorios maduran: MiCA en EU, ETF aprobados en USA.

Limitaciones honestas

- Mayoría tokeniza el wrapper, no el activo subyacente — dependes del emisor.
- Liquidez secundaria todavía limitada · regulación KYC en cada jurisdicción.
- El "sueño" de tokenizar inmobiliario residencial sigue siendo más promesa que realidad.



Blockchain · más allá de la especulación

DeFi · finanzas sin intermediarios · estado actual

Qué funciona · qué no · qué aprender.

Primitivas que funcionan en producción

- DEX (Uniswap, Curve) — >2 billones USD en volumen acumulado · AMMs reemplazan order books.
- Lending (Aave, Compound) — colateralizado, sin contraparte humana · liquidaciones automáticas.
- Stablecoins descentralizadas (DAI, sDAI) — refugio de exposure a USD.
- Liquid staking (Lido, Rocket Pool) — yield de PoS sin lock-up.

Lo que sigue siendo fricción

- UX terrible para usuarios no-cripto · wallets, gas, signing → muralla.
- Riesgo de smart contract sigue presente · DeFi insurance es subdesarrollada.
- Compliance: cumplir KYC + Travel Rule en pools sin permisos es estructuralmente difícil.

TVL y madurez

- TVL global DeFi: ~100B USD (caído desde el peak de 180B en 2021, pero estabilizado).
- Aave V3 funciona en 12+ cadenas · multichain ya es default.
- Auditorías rigurosas · bug bounties de millones · seguros (Nexus Mutual).

Tendencia: DeFi institucional

- Aave Arc, Compound Treasury — pools KYC con whitelists.
- Maple Finance · undercollateralized lending a market makers regulados.
- Convergencia con TradFi: bonos tokenizados como colateral, no solo cripto-nativo.

Blockchain · más allá de la especulación

Casos en producción · más allá de las finanzas

Cuándo blockchain aporta valor real.

1

Identidad y credenciales (verifiable credentials)

- European Digital Identity Wallet (EUDI) · estándares basados en blockchain pública.
- Diplomas universitarios on-chain · MIT, EPFL, Universidad de Múnich.
- ENS (Ethereum Name Service) · 3M+ nombres registrados como identidad portable.

2

Supply chain con verificación criptográfica

- Origen de medicamentos · MediLedger (Pfizer, Genentech) cumple DSCSA.
- Trazabilidad alimentaria · IBM Food Trust (Walmart, Carrefour, Nestlé).
- Diamantes · Everledger registra >2M piedras.

3

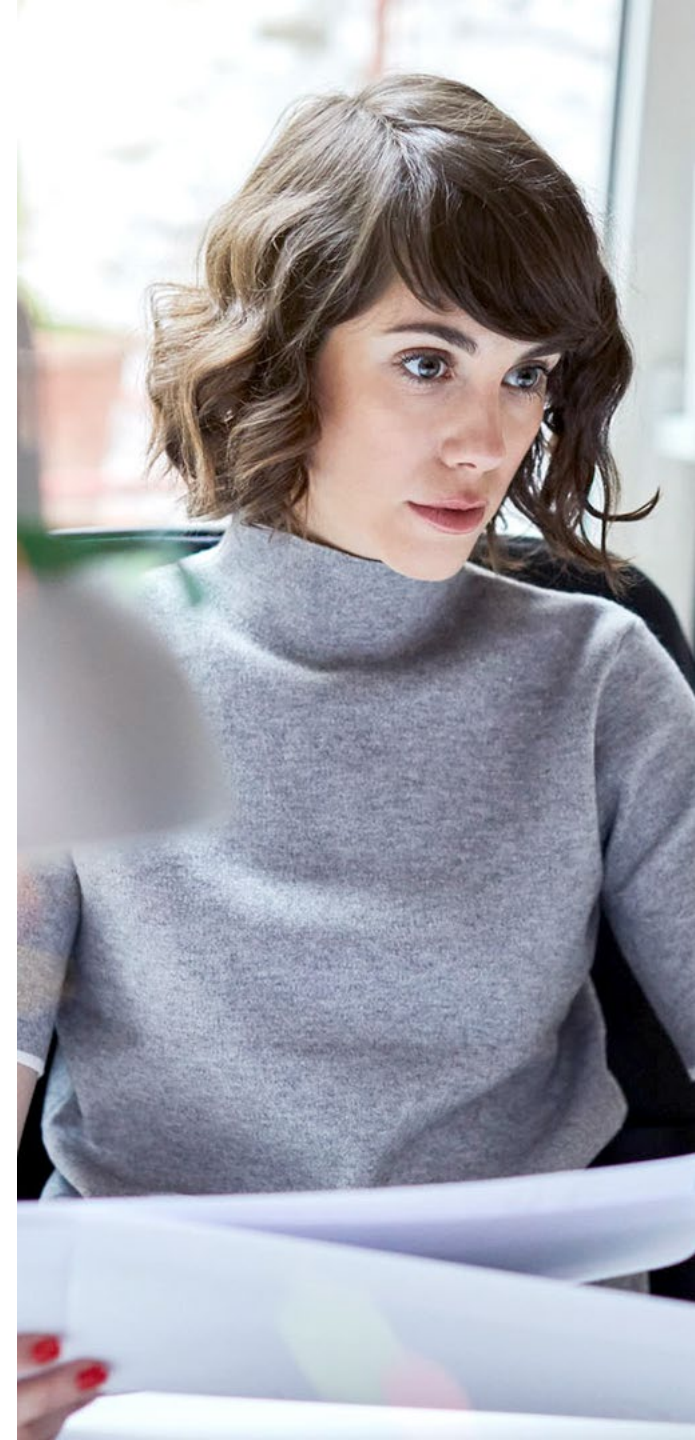
Royalties y propiedad intelectual

- Música: Audius, Sound.xyz · payout directo a artistas · sin sello.
- Coleccionables digitales con royalties enforceables on-chain.

4

Públicos y gobierno

- Estonia · backbone de su digital state usa KSI Blockchain (Guardtime) desde 2012.
- Voto on-chain en pequeña escala · West Virginia (militares en el extranjero, 2018).
- Registros de propiedad: Suecia, Georgia, Honduras (pilotos, no producción nacional).



Blockchain · más allá de la especulación

Crítica honesta · cuándo NO usar blockchain

La pregunta que debe hacerse en cada proyecto.



Antipatrones comunes

- "Blockchain interno" sin múltiples partes desconfiadas → es un log con hashes, usad Postgres con WORM.
- Tokenizar todo solo porque suena moderno → coste y complejidad sin valor añadido.
- Pretender que la inmutabilidad resuelve el GIGO (garbage in, garbage out).



Riesgos sistémicos

- Concentración de validators · Lido + Coinbase = >40% stake Ethereum.
- Custodia de claves: pérdida = pérdida del activo, sin recuperación.
- MEV (Maximal Extractable Value): ordenamiento de tx puede ser manipulado.



Limitaciones reales

- Throughput: incluso L2s rápidas son 1000× más caras que SQL para writes simples.
- Privacidad: cadena pública revela patrones que regulación o competencia pueden explotar.
- Latencia: finalidad típica es segundos-minutos, no microsegundos.
- Reversibilidad: no hay "deshacer" — bug = pérdida permanente.



Test de dos preguntas (Wüst-Gervais)

- ¿Hay múltiples writers que no confían entre sí? · ¿No hay un tercero confiable disponible?
- Si ambas son SÍ, blockchain puede aportar · Si alguna es NO, hay opción más simple.

Blockchain · más allá de la especulación

Síntesis del bloque

01

Blockchain ya no es promesa: es infraestructura de pagos (stablecoins) y tokenización de activos (treasuries) en uso institucional.

La elección entre red pública, permitisionada o híbrida es la decisión arquitectónica más importante — y la que más se equivoca.

02

03

Smart contracts son código en producción con dinero real — exigen rigor de software crítico (auditoría, verificación formal, observabilidad).

MiCA y la regulación EU dan certidumbre operativa — el viento regulatorio sopla a favor de los actores serios.

04

05

La pregunta correcta no es "¿cómo usamos blockchain?" sino "¿en qué procesos hay múltiples partes que necesitan compartir un registro verificable?"

04

Realidad Extendida · RV, RA y dispositivos espaciales

Spatial computing · gemelos digitales ·
interfaces post-pantalla



Realidad Extendida · RV, RA y dispositivos espaciales

RV, RA, MR, XR · cuatro letras y un mismo continuum

Milgram (1994) sigue siendo la mejor brújula taxonómica.

Continuum (Milgram)

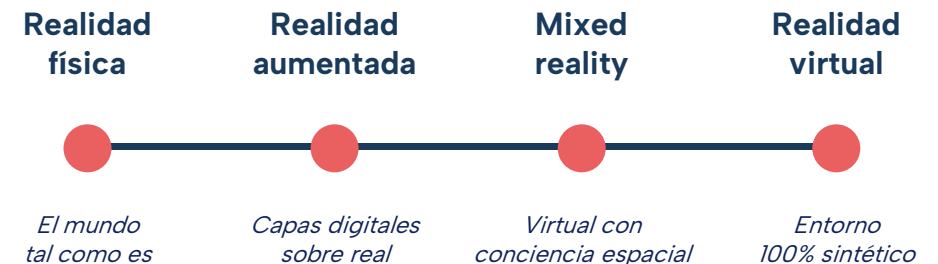
- Real → RA → MR → RV · eje, no categorías discretas.
- Cambia la proporción de píxeles sintéticos vs reales.

Definiciones operativas

- RV · 100% sintético · headset opaco.
- RA · capas sobre real · móvil / gafas transparentes.
- MR · virtual con conciencia de geometría real.
- XR · paraguas · spatial computing es su concreción 2024.

Lo que NO son

- Pantalla 3D ≠ RV (falta tracking de cabeza).
- Filtro Snapchat ≠ RA productiva (falta persistencia).
- Metaverso ≠ RV (es servicio, no tecnología).



Continuum realidad ↔ virtualidad (Milgram & Kishino, 1994)

Realidad Extendida · RV, RA y dispositivos espaciales

Hardware 2024–2025 · el momento de la maduración

Por primera vez los dispositivos llegan al umbral usable.

Headsets MR insignia

- Apple Vision Pro (3.500 USD) · 4K por ojo · eye tracking precisísimo · M2+R1 chip.
- Meta Quest 3 (500 USD) · passthrough color · standalone · masivo (>20M unidades 2024).
- Quest 3S (300 USD) · cuña de mercado mainstream · sept 2024.



Gafas RA óptica transparente

- Meta Orion (prototipo 2024) · gafas reales con waveguide y LLM nativo · no comercial todavía.
- Magic Leap 2 (3.300 USD) · enterprise · medical, defensa.
- Snap Spectacles V5 (kit dev) · campo de visión 46° · gestos manos.

Componentes que han madurado

- Pantallas micro-OLED (Sony) >3000 ppi · pixel density indistinguible del ojo humano.
- Eye tracking sub-grado · permite foveated rendering (renderizar en alta solo donde miras).
- Hand tracking sin guantes · gestos como input principal.
- Passthrough color de baja latencia (<10ms) · mixed reality realista.

Realidad Extendida · RV, RA y dispositivos espaciales

Casos productivos · dónde aporta hoy

Más allá del gaming, qué funciona en empresa.

Formación · ROI demostrado

- Cirugía simulada · Osso VR, FundamentalVR · reduce errores 40–60 % en quirófano (Harvard Med 2020).
- Operadores industriales · Boeing, Siemens, BMW reducen tiempo formación 40 %.
- Soft skills · entrevistas, presentaciones, negociación · simulación con IA (PIXL, Bodyswaps).
- Ahorro ROI: PWC (2020) — formación XR es 4× más rápida que aula y 275 % más confiada.

Diseño y revisión 3D

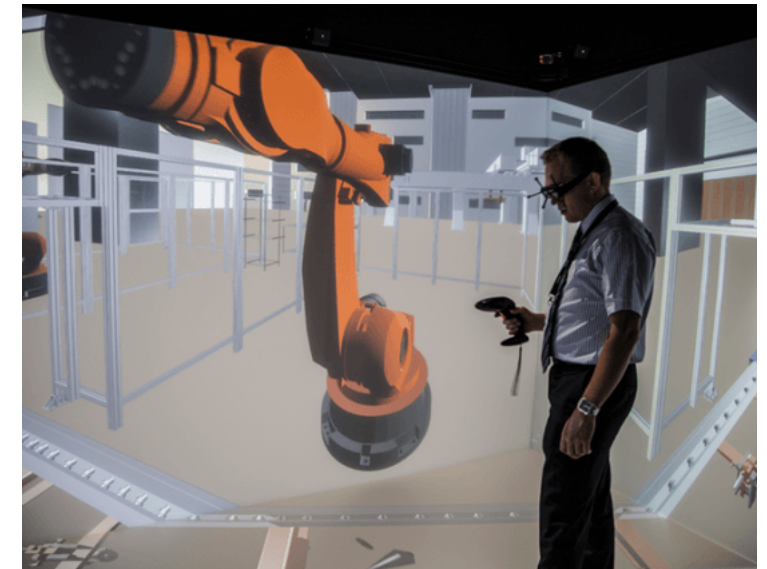
- BMW, Volvo, Boeing revisan prototipos en RV antes de fabricar — ahorros de millones.
- Arquitectura · clientes recorren el edificio antes de obra · reduce cambios costosos.
- Twinmotion, Enscape, Iris VR para arquitectos · workflows establecidos.

Asistencia remota y mantenimiento

- Operario con HoloLens recibe overlays de experto remoto · ABB, GE, Boeing.
- Reducción 30–50% tiempo resolución · aumento first-time-fix rate.

Salud y rehabilitación

- Tratamiento de fobias y PTSD (Bravemind, USC) · evidencia clínica robusta.
- Manejo de dolor crónico · EaseVRx (FDA aprobado 2021) · primera prescripción XR.



This Photo by Unknown Author is licensed under [CC BY-SA](#)

Realidad Extendida · RV, RA y dispositivos espaciales

Spatial computing · la apuesta de Apple y la respuesta del mercado

Más que un producto · una nueva categoría de cómputo.

Qué cambia técnicamente

- El SO entiende el espacio 3D del usuario (room mesh, semantic understanding, persistencia).
- Las apps son volúmenes en el espacio, no rectángulos en pantalla.
- Input multimodal: ojos miran (selección), manos pinch (acción), voz (comando).
- Persistencia espacial: un objeto digital queda donde lo dejas físicamente.

Vision Pro · qué hace bien y qué no

- Bien: fidelidad visual, eye/hand tracking, ergonomía corta sesión, ecosistema iOS.
- Mal: peso (650g), batería externa, precio, contenido ecosistema todavía pobre.
- ROI corporativo: emerging para colaboración remota, design review, training.

Respuesta competitiva

- Samsung+Google+Qualcomm · headset Android XR (2025) · respuesta directa al Vision Pro.
- Meta Quest 3 + visionOS-style features · presión de precio masivo.
- Microsoft pivote: Mesh para Teams · MR colaborativo en Microsoft 365.

Tendencia clave

- Del headset al "forma de gafas" (Orion, Spectacles) · 5-10 años para versión consumer.
- LLM nativos en gafas · asistente contextual por voz/visión · Meta/Google/Apple invirtiendo.



Realidad Extendida · RV, RA y dispositivos espaciales

Gemelos digitales · el lado serio de XR industrial

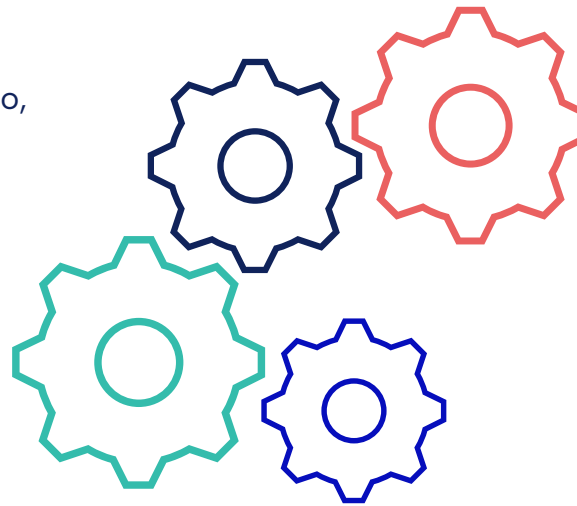
Donde XR converge con IoT, simulación y BIM.

Qué es un digital twin

- Réplica virtual de un activo, proceso o sistema físico, sincronizada en tiempo real con sensores.
- No es CAD ni BIM — incorpora datos vivos de IoT y permite simulación predictiva.
- Niveles: descriptivo (visualización) · diagnóstico · predictivo · prescriptivo (autonomous).

Casos en producción

- Siemens · gemelo de fábricas completas · MindSphere + NX · usado en BMW, Anheuser-Busch.
- GE Vernova · turbinas eólicas · cada turbina tiene su gemelo, ahorro 8-12% mantenimiento.
- Singapur · ciudad entera modelada · Virtual Singapore · planificación urbana.



Plataformas del mercado

- NVIDIA Omniverse · USD universal scene description · física precisa · colaboración multiusuario.
- Microsoft Azure Digital Twins · ontologías DTDL · integración con Azure IoT.
- Unity Industrial Collection · Unreal Twinmotion · enfoque en visualización.

Por qué importa ahora

- Convergencia con IA generativa · gemelos que se diseñan/optimizan automáticamente.
- OpenUSD como estándar abierto · soportado por Apple, Pixar, NVIDIA, Adobe.
- Coste de cómputo cae · simulación de fábrica entera viable en cluster pequeño.

Realidad Extendida · RV, RA y dispositivos espaciales

Métricas técnicas · qué pedirle a un dispositivo XR

Para diferenciar producto serio de demo.

01

Latencia motion-to-photon

- Umbral de mareo (sickness): >20ms induce malestar en mayoría de usuarios.
- Vision Pro y Quest 3 están en 11-13ms · clase enterprise.
- Headsets baratos sin tracking adecuado superan 30-50ms · inutilizable para sesiones largas.

02

Campo de visión y resolución

- FOV humano: ~210° horizontal · headsets 2024: 100-110° (Vision Pro 100°, Quest 3 110°).
- Resolución: 3K-4K por ojo · pixel density >35 ppp (pixels per degree) para no ver "screendoor"
- Apple Vision Pro: 23M píxeles totales · estado del arte.

03

Frame rate y refresco

- Mínimo aceptable: 90 Hz · óptimo: 120 Hz · profesional: 144+ Hz.
- Refresco consistente importa más que pico (jitter genera mareo).
- Foveated rendering permite 90 Hz en hardware modesto · clave económicamente.

04

Tracking y comfort

- 6DoF inside-out (cámaras en headset) ya estándar · sin estaciones externas.
- Eye tracking <2° error · hand tracking <5mm precision.
- Peso óptimo <500g · IPD ajustable · refresh comfort breaks recomendado cada 30 min.

Realidad Extendida · RV, RA y dispositivos espaciales

WebXR y estándares abiertos · contenido portable

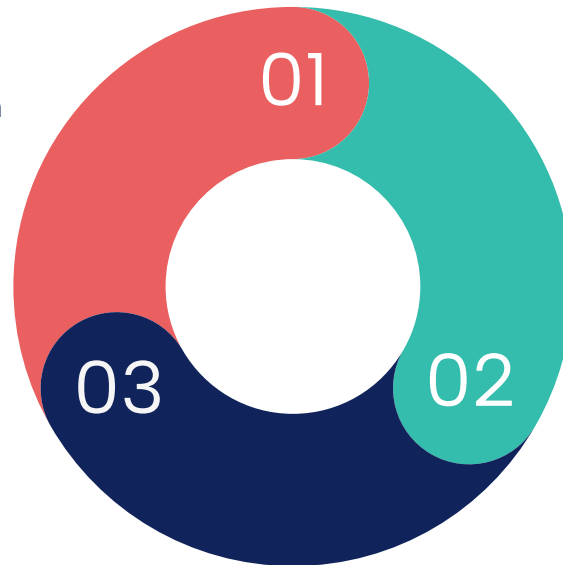
No quedarse atrapado en un ecosistema.

WebXR Device API

- Estándar W3C para experiencias XR en navegador · sin instalar apps.
- Soportado en Chrome, Edge, Quest browser, Vision Pro Safari (parcial).
- Frameworks: A-Frame, Three.js + WebXR, Babylon.js, PlayCanvas.

OpenUSD · contenido 3D portable

- Universal Scene Description · originado en Pixar 2016 · open-source.
- Estándar de facto en industrial (NVIDIA, Apple, Microsoft, Adobe).
- Permite gemelos digitales y assets reusables entre Omniverse, Vision Pro, Unity.



OpenXR · Khronos Group

- API estándar para acceder a runtimes XR · sin acoplar a SDK propietario.
- Soportado por Meta, Microsoft, Valve, ByteDance · Apple aún no (visionOS propio).
- Reduce coste de portar entre headsets de 60% a <20 %.

Lecciones para empresa

- Apostar por web/standards reduce dependencia de hardware específico.
- Para ROI a 5+ años, el estándar abierto > rendimiento máximo de plataforma cerrada.

Realidad Extendida · RV, RA y dispositivos espaciales

Riesgos · seguridad, salud, privacidad, ética

Lo que un comité de ética debe revisar antes de despliegue masivo.

Salud física y bienestar

- Cybersickness · 20–50% de usuarios afectados en sesiones >30 min · decreciente con buen hardware.
- Fatiga ocular y vergence-accommodation conflict · investigación abierta sobre efectos largo plazo.
- Lesiones por colisión con entorno físico · Quest 3 mitiga con guardian system.

Privacidad masiva (XR genera más datos que cualquier dispositivo)

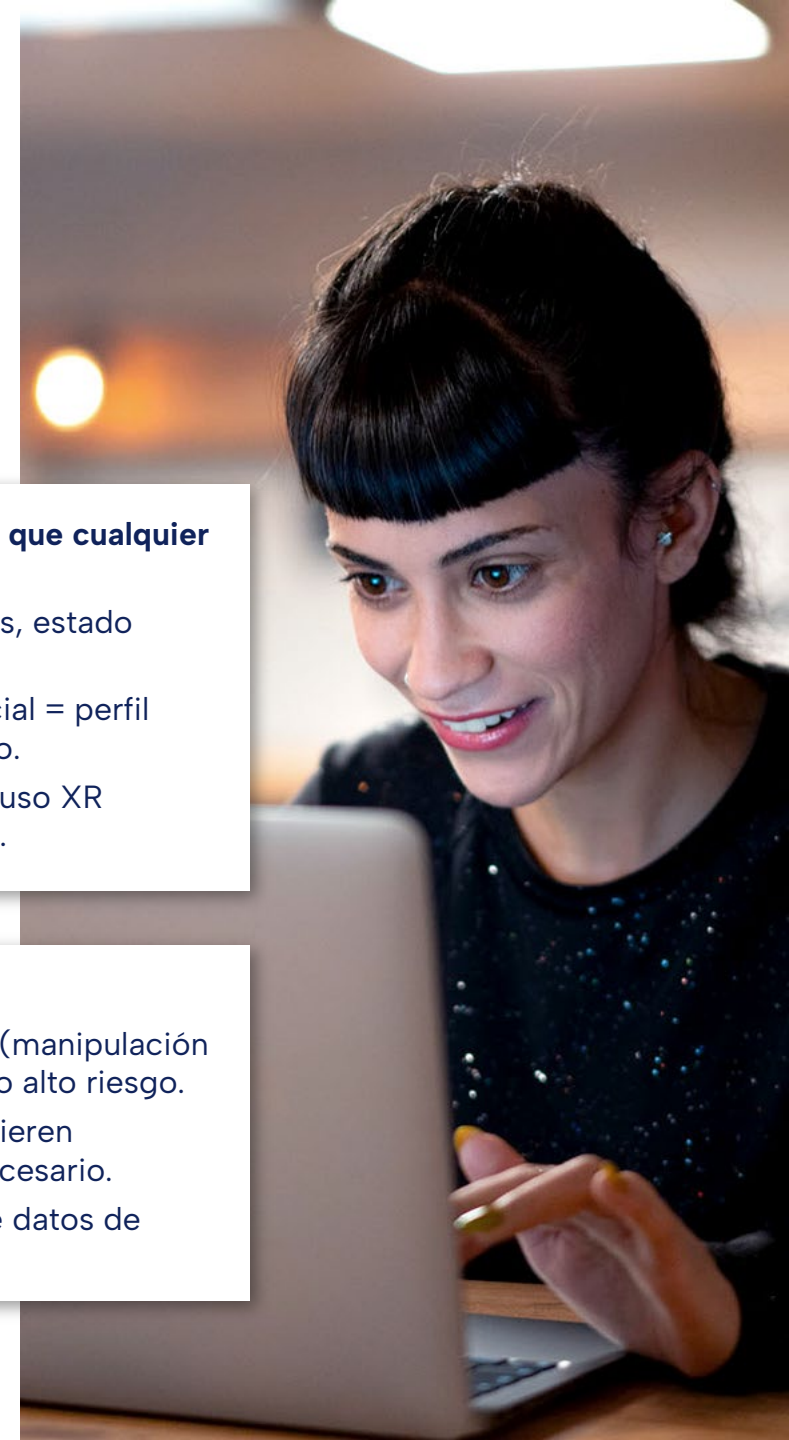
- Eye tracking revela atención, intereses, estado emocional, posibles patologías.
- Hand tracking + voz + posición espacial = perfil biométrico extremadamente detallado.
- Estudio Stanford 2023: 5 minutos de uso XR identifican usuario con 95% precisión.

Seguridad y manipulación

- Inmersión total amplifica persuasión · publicidad en RV mucho más efectiva (Stanford VHIL).
- Deepfakes en avatares · suplantación de identidad social en metaverso.
- Riesgo niños · contenidos no apropiados, adicción, exposición a desconocidos.

Regulación emergente

- EU AI Act categoriza ciertos usos XR (manipulación emocional, biometric inference) como alto riesgo.
- GDPR aplica · datos biométricos requieren consentimiento explícito y mínimo necesario.
- FTC USA · investigación a Meta sobre datos de menores en Quest 2023–2024.



Realidad Extendida · RV, RA y dispositivos espaciales

Stack profesional XR · capas y elecciones

Mapa para arquitectos que evalúan adopción.

Capa de hardware

- Headsets: Vision Pro, Quest 3/3S, Pico 4 Enterprise, HoloLens 2 (legacy), Magic Leap 2.
- Controladores: hand tracking nativo > controladores físicos · evolución clara.

Capa de runtime / SDK

- OpenXR (estándar) · Meta SDK · visionOS frameworks (RealityKit, ARKit).
- Cross-platform: Unity XR Interaction Toolkit · Unreal OpenXR plugin.

Capa de motor 3D y herramientas

- Unity XR · dominante en formación, salud, simulación.
- Unreal Engine 5 · alta fidelidad gráfica · Twinmotion para arquitectura.
- Web: Three.js + WebXR · A-Frame · Babylon.js.
- Industrial: NVIDIA Omniverse + USD · gemelos digitales.

Capa de servicios

- Backend MR multiusuario: Photon, Normcore, Coherence.
- Streaming XR (CloudXR, NVIDIA): renderizado en cloud, headset thin client.
- Asset libraries: Sketchfab, TurboSquid, Quixel Megascans.

Realidad Extendida · RV, RA y dispositivos espaciales

Cómo adoptar XR sin gastar 1M€ en una experiencia que falla

Hoja de ruta práctica para 2025–2027.

Fase 1 · Exploración (3–6 meses, <30K€)

- Inventario de procesos candidatos: formación compleja, diseño 3D, asistencia remota, simulación de riesgo.
- Compra 5–10 Quest 3 + suscripción a Demo apps (Osso VR, Bodyswaps, Mesh).
- Sesiones con stakeholders · identificar 1–2 casos con ROI claro.

Fase 2 · Piloto (6–12 meses, 50–200K€)

- Construir 1 caso end-to-end con métricas pre/post · KPIs duros (tiempo, errores, satisfacción).
- Empezar con Unity + OpenXR + Quest 3 · evitar over-engineering inicial.
- Privacy Impact Assessment · políticas de uso · formación a usuarios pilotos.

Fase 3 · Escalado (12–24 meses, 200K–1M€)

- Si fase 2 mostró ROI, expandir · estandarizar arquitectura · MDM para fleet de devices.
- Integración con sistemas corporativos (LMS, MES, CRM, gemelo digital).
- Plan de obsolescencia · headsets se actualizan cada 2–3 años · dispositivos no son CapEx eterno.

Anti-patrones a evitar

- Comprar 500 headsets antes de validar 1 caso de uso.
- Encargar contenido custom a agencia sin métricas de éxito.
- Asumir que personal usará XR sin entrenamiento e incentivos.



Realidad Extendida · RV, RA y dispositivos espaciales

Síntesis del bloque

01

XR vive su transición de "juguete tecnológico" a "infraestructura productiva" gracias a hardware finalmente usable y a casos de uso con ROI medido.

02

La distinción entre RV (aislamiento), RA (aumento), MR (fusión) y spatial computing es clave — cada uno tiene casos donde gana y donde no aporta.

03

Apostar por estándares abiertos (OpenXR, OpenUSD, WebXR) protege la inversión a 5-10 años contra fragmentación de plataformas.

04

Los gemelos digitales son donde el dinero real está en industrial · convergencia con IoT, IA, simulación.

05

Privacidad y salud son riesgos serios · ningún despliegue masivo sin Privacy Impact Assessment y políticas claras.



¡Gracias!