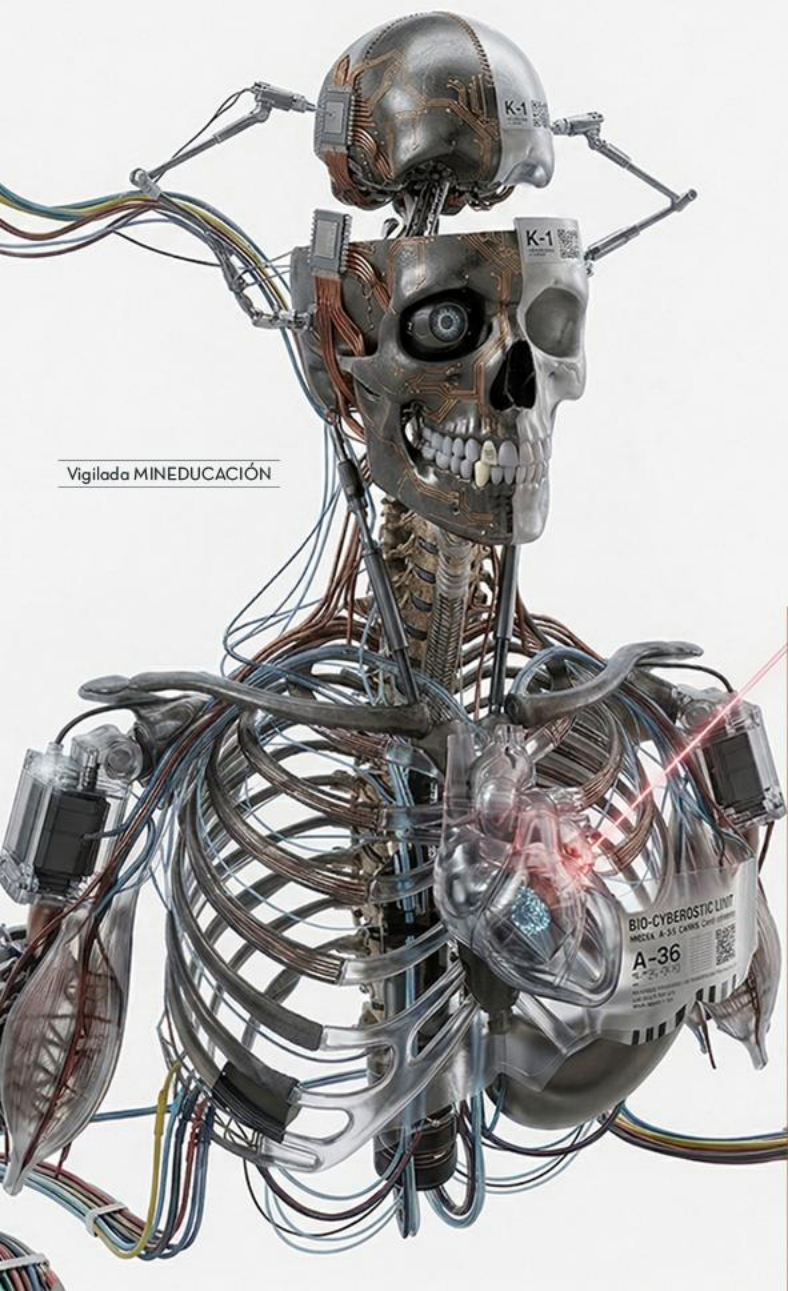




Vigilada MINEDUCACIÓN

Comunidad Pedagógica

Grancolombiana 2026

2/9. Saberes en conflicto:
Fundamentos, flujos, herramientas
y oportunidades prácticas
de la IA Generativa.

- Anatomía de la IA Generativa
- Herramientas y tareas
- Flujos de trabajo integral

Ponentes

María Gaby Boshell Villamarín
Vicerrectora de Desarrollo Académico

Fabián Eduardo Cáceres Cáceres
Departamento de Biblioteca - Referencista

UNA
EXPERIENCIA
DE VIDA

- Entender qué es de verdad (IA) (¿qué es?)
- Cómo funciona por dentro (¿cómo funciona?)
- Problemas (cómo mitigarlos) (¿cuáles son sus fallos?)
- Marco para elegir herramientas (criterio propio) (¿cómo lo uso yo?)
- Identificar oportunidad concreta (Opportunity Canvas)

“~~Inteligencia~~ artificial”

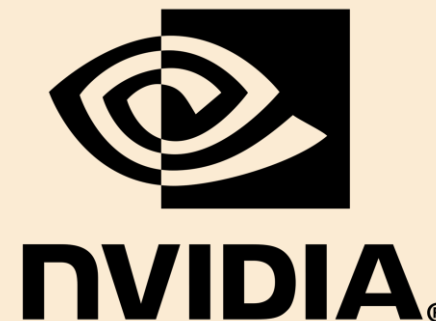
«¿Qué es?»

“~~Inteligencia~~ artificial”

«¿Cuántas de estas inteligencias puede demostrar una máquina?»

“Hoy en día definimos a la inteligencia como la capacidad
para resolver problemas a través de una gama amplia de habilidades”
(UNICEF, 2025)

**Lingüística - Lógico-matemática – Espacial – Musical - Emocional
– Interpersonal – Intrapersonal – Naturalista - Cinestésico-corporal**



Superó los **5 billones de dólares** en capitalización en 2025

Sus GPUs resultaron ser la **arquitectura perfecta para entrenar modelos de IA** — sin haber sido diseñadas para eso (procesadores gráficos para videojuegos).

Sus GPUs son el **hardware estándar para entrenar modelos avanzados de IA**. Su sistema CUDA es el "sistema operativo" de facto de los centros de datos de IA.

Los principales proveedores de nube e hiperescaladores proyectan gastar ~600.000 millones de dólares en 2026 en infraestructura — **la gran mayoría en chips NVIDIA**.

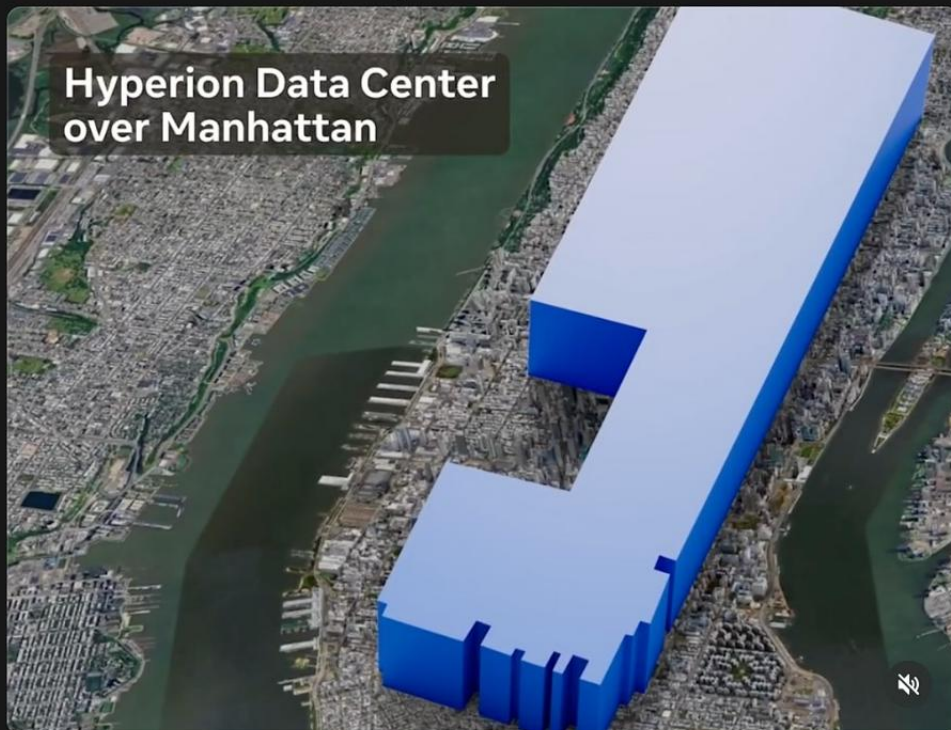
Una sola búsqueda en ChatGPT consume aproximadamente **10 veces más energía que una búsqueda en Google**.

Millan, S. (2026, 17 de enero). *Cómo Nvidia puso el mundo a sus pies: el secreto de un asalto al poder tecnológico y financiero*. El país



zuck 1 día

En realidad estamos construyendo varios clústeres de múltiples gigavatios. Al primero lo estamos llamando Prometheus y estará en línea en el '26. También estamos construyendo Hyperion, que podrá escalar hasta 5 GW a lo largo de varios años. Además, estamos construyendo varios clústeres titanes más. Solo uno de estos cubre una parte significativa del tamaño de Manhattan.



478 58 44 175

X Daily News [@xDaily]. (2025, 14 de julio). *NEWS: Mark Zuckerberg announced that Meta is building several massive, multi-gigawatt AI clusters.*

Ver <https://x.com/xDaily/status/1944785388771357032>. Para cobertura reciente [Imagen adjunta] [Publicación]. X.

Inteligencia Artificial: De la Infraestructura a la Interfaz

Es donde el humano interactúa.

System Prompt: Define la personalidad ("Eres un abogado experto").

User Prompt: La instrucción específica del usuario.

Interfaz: El chat o aplicación que traduce el código a una conversación amigable y todo el ecosistema circundante.

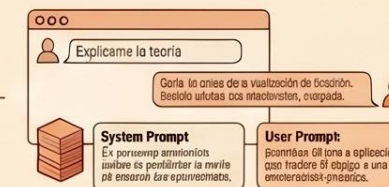
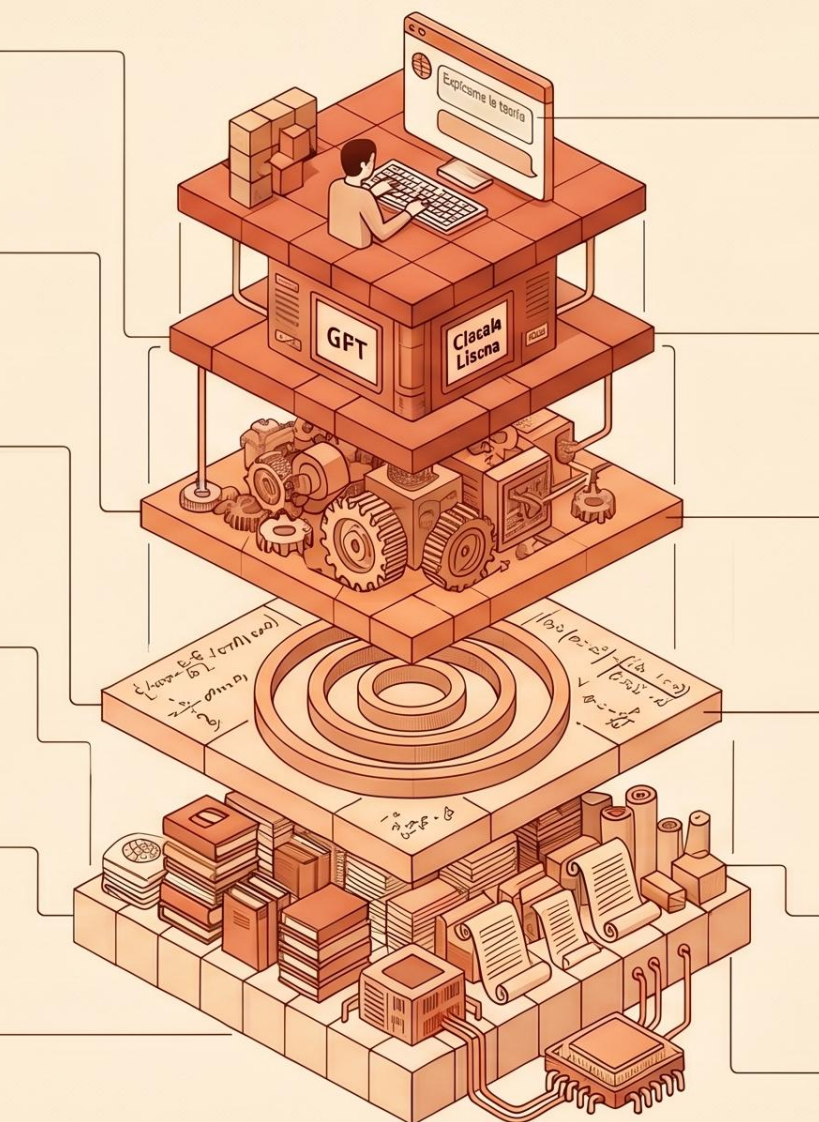
Se crea el modelo base. Aquí ocurre la "Naturaleza Estocástica" (predicción de la siguiente palabra más probable). También se aplica el RLHF (Reinforcement Learning from Human Feedback) para **"educar"** al modelo y que sus respuestas sean útiles y seguras.

El Transformer es el diseño del **cerebro de la IA**. Permite que el modelo entienda el contexto de una palabra según las palabras que tiene alrededor (mecanismo de atención). Aquí el texto se convierte en Tokens (vectores numéricos).

Se definen las reglas del aprendizaje. El Deep Learning imita las conexiones neuronales humanas para que la máquina aprenda patrones complejos por sí sola, sin ser programada explícitamente para cada tarea.

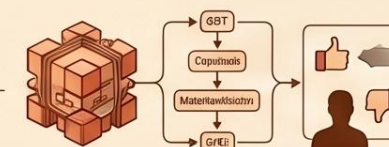
Se recolecta, limpia y organiza la información. Es el **"alimento"** del modelo. La calidad de la IA depende de la calidad de estos datos.

Es la potencia bruta. Aquí se procesan los trillones de **cálculos matemáticos necesarios** para que la IA **"piense"**. Sin este hardware, el software no tiene donde correr.



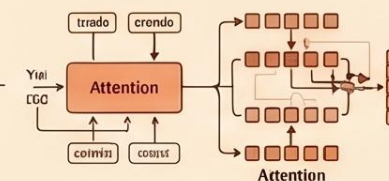
Nivel 6
Capa de Aplicación e Interacción
(La Interfaz)

Tecnología: UI (ChatGPT, Copilot),
System Prompts, User Prompts.



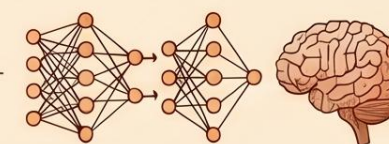
Nivel 5
El Núcleo: LLM y
Entrenamiento (El Modelo)

Tecnología: Pre-entrenamiento
Generativo (GPT, Claude, Llama).



Nivel 4
La Arquitectura "Transformer" y
PLN (El Motor)

Transformers (Atención),
PLN (Procesamiento de
Lenguaje Natural).



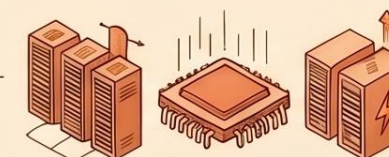
Nivel 3
El Marco Científico (IA > ML > DL)

Tecnología: Inteligencia Artificial,
Machine Learning
(Aprendizaje Automático) y Deep
Learning (Redes Neuronales Profundas).



Nivel 2
Cimientos de Datos y
Big Data (La Materia Prima)

Tecnología: Datasets masivos
(Common Crawl, Wikipedia,
libros, código).



Nivel 1
Infraestructura de Cómputo (El Suelo)

Tecnología: GPUs (Nvidia),
TPUs (Google), Cloud Computing.

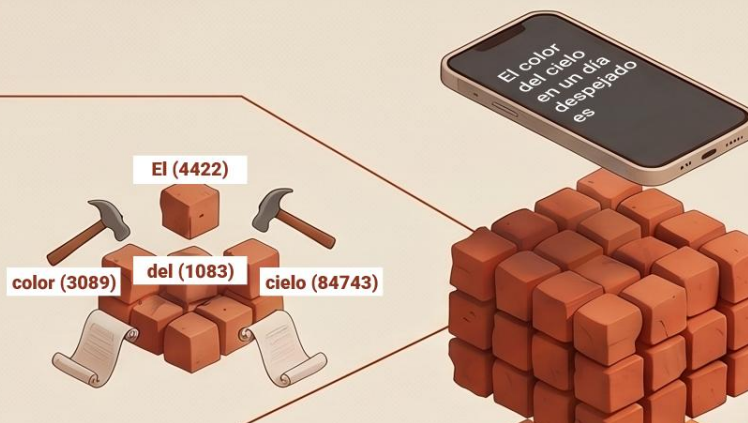
Potencia bruta (hardware) – Datos – Marco científico – Aplicación práctica – Interacción humana

Fase 1

Ingesta y Fragmentación (La Entrada)

Tokenización y Normalización.

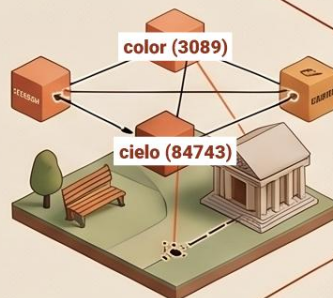
El texto humano se limpia y se rompe en "tokens" (pedazos de palabras). Es la unidad mínima de procesamiento. Un token equivale aproximadamente a 3/4 de una palabra en inglés/español.

**Fase 3** Procesamiento de Atención

(El Corazón del Transformer)

Self-Attention (Autoatención).

El modelo hace un "escaneo cruzado". Cada palabra mira a El modelo descarta la importancia de palabras como "El" o "un" y concentra toda su energía computacional en la combinación "color + cielo + despejado". Entiende que "despejado" es la condición que define al color.

**Fase 5** Selección y Decodificación

(La salida)

Sampling (Muestreo) y De-tokenización.

Basado en la Temperatura, el modelo elige un token (ej. "Bien"). Ese ID numérico se traduce de nuevo a texto humano y se muestra en pantalla.

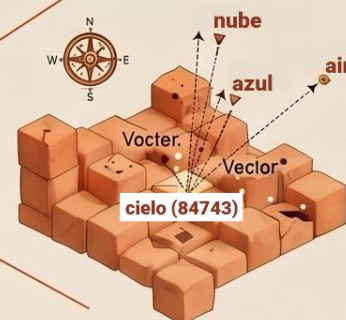
El proceso se repite. La palabra "Bien" ahora se suma al prompt original y el modelo vuelve a empezar para predecir la siguiente palabra, hasta que genera un token de "Fin de texto".

**Motores de predicción de tokens**

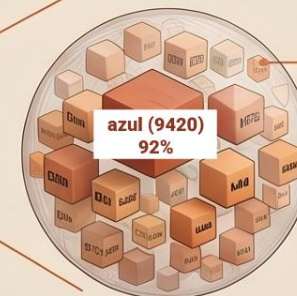
La probabilidad no es verdad.
Un LLM no "sabe", "piensa" ni "recuerda"

Fase 2

Traducción Matemática (El Espacio Vectorial)

**Embedding + Positional Encoding.**

Cada token se convierte en una lista de cientos de números (vectores). Aquí se le añade la "brújula" de posición para que el modelo sepa el orden de la oración. Los Embeddings sitúan las palabras en un mapa de miles de dimensiones donde "Rey" está cerca de "Reina" y lejos de "Manzana".

Fase 4 Predicción de Probabilidades
(La Cabeza del Lenguaje)**Softmax y Capa Lineal.**

El modelo calcula una lista de puntajes para todas las palabras que conoce (su vocabulario).

Por ejemplo:

Azul: 92% de probabilidad.

Gris: 4% de probabilidad.

Rojo: 2% de probabilidad.

Verde: 0.1% de probabilidad.

Es una distribución de probabilidad sobre el espacio del vocabulario.

El LLM no "sabe" que el cielo es azul. Lo que hizo fue:

Convertir tu frase en coordenadas matemáticas.

Ver qué coordenadas suelen estar juntas en su entrenamiento.

Calcular que, estadísticamente, después de la secuencia "color-cielo-despejado", la palabra con más "peso" es Azul.

Un LLM

No razona: predice.

No recuerda: almacena patrones.

No verifica: interpola.

Fallos en la IA generativa —————→

Fallos en la IA generativa: Taxonomía y limitaciones críticas

ARQUITECTURA DEL MODELO	DATOS DE ENTRENAMIENTO	MÉTODOS DE AJUSTE (RLHF/Fine-tuning)
Alucinaciones factuales. El modelo puede inventar información y presentarla como cierta, porque no tiene forma de verificar la verdad ni es consciente de que está equivocándose. (citas bibliográficas con DOIs reales que no existen, estadísticas con decimales precisos pero fabricadas, sentencias judiciales falsas).	Sesgos (Bias). El sesgo en IA es la tendencia sistemática del modelo a producir resultados que favorecen o discriminan a ciertos grupos, ideologías, culturas o perspectivas. No es un error aleatorio sino un patrón estructural reproducido a escala masiva (<i>sesgos del corpus con el que fue entrenado</i>).	Inyección de prompts. La inyección de prompts (prompt injection) es un ataque en el que instrucciones maliciosas son embebidas en el texto de entrada (directas) o en documentos externos que el modelo procesa (indirectas), con el objetivo de subvertir el sistema de instrucciones original, eludir restricciones de seguridad, exfiltrar información del contexto o ejecutar acciones no autorizadas en sistemas agénticos.
Trampa del siguiente token. Los LLM generan texto token a token de forma autorregresiva: cada elemento depende solo de los anteriores y no puede corregirse, por lo que el modelo queda comprometido con la dirección desde el inicio.	Envenenamiento de datos. Más allá de las alucinaciones generales, los LLM son especialmente peligrosos cuando generan desinformación estructurada: citas bibliográficas con DOI inventados, estadísticas con decimales precisos pero fabricadas, sentencias judiciales falsas, URLs que parecen reales, o atribuciones incorrectas de declaraciones a personas reales.	Jailbreaking. El jailbreaking es el conjunto de técnicas mediante las cuales un usuario logra que un LLM eluda las restricciones de seguridad implementadas durante el proceso de RLHF (Reinforcement Learning from Human Feedback) o el ajuste con IA Constitucional.
Ausencia de razonamiento profundo. Los LLM muestran un razonamiento superficial (estadístico): imitan patrones convincentes, pero fallan en inferencia lógica, aritmética precisa, causalidad y planificación en múltiples pasos.	Memorización de datos sensibles. La memorización ocurre cuando un LLM almacena en sus parámetros fragmentos literales de su corpus de entrenamiento y puede reproducirlos ante ciertas entradas. Esto supone riesgos de privacidad (reproducción de datos personales, médicos o financieros), de propiedad intelectual (reproducción de texto con copyright) y de seguridad (filtración de secretos presentes en el corpus).	Sycophancy (complacencia). La sycophancy o complacencia aduladora es la tendencia del modelo a modificar sus respuestas para alinearse con lo que percibe que el usuario quiere escuchar, incluso cuando esto implica validar afirmaciones incorrectas, retractarse de respuestas correctas ante presión del usuario, exagerar elogios o evitar correcciones incómodas.
Amnesia catastrófica. La amnesia catastrófica ocurre cuando, al ajustar un modelo con nuevos datos o tareas, este olvida masivamente lo aprendido previamente.		

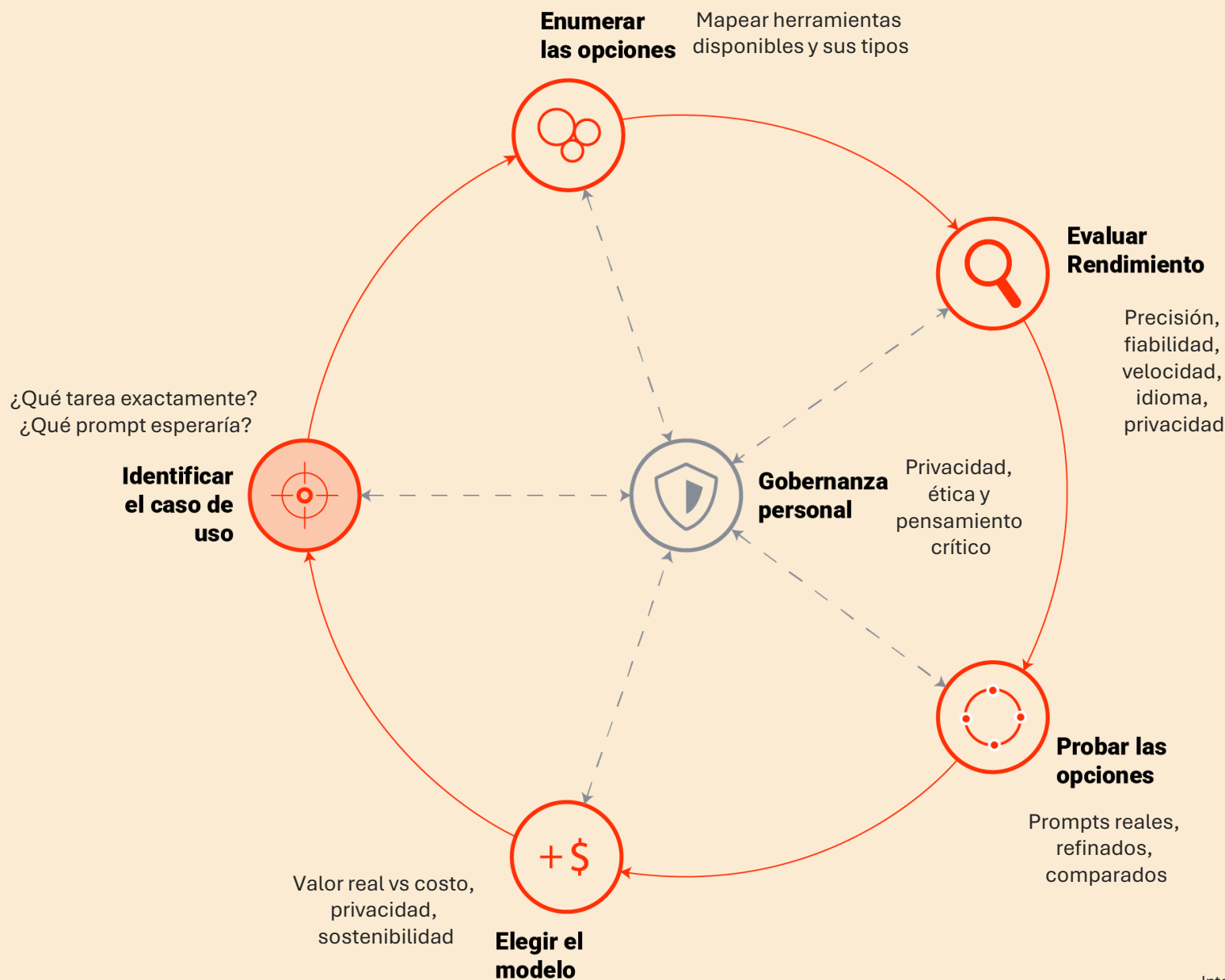
¿Cómo intentar controlar esto?

Algunos de estos fallos están en la **arquitectura del modelo**, en los **datos con que fue entrenado**, en el **proceso de ajuste**.
Eso no lo puedo modificar como usuario.

FALLO	ORIGEN PRIMARIO	GRAVEDAD	SOLUCIÓN PRIORITARIA	
Alucinaciones	Arquitectura	★★★★★	RAG + verificación factual	
Sesgos	Datos	★★★★☆	Curación de corpus + Constitutional AI	
Información falsa	Arquitectura	★★★★★	RAG + NLI post-generación	
Trampa del siguiente token	Arquitectura	★★★★☆	Chain-of-Thought + ToT	
Sin razonamiento lógico	Arquitectura	★★★★☆	CoT + herramientas externas	
Inyección de prompts	Integración/Ajuste	★★★★★	Separación de contextos + guardrails	
Amnesia catastrófica	Arquitectura	★★★★☆	LoRA / EWC	
Jailbreaking	Ajuste (RLHF)	★★★★☆	Constitutional AI + red teaming	
Memorización / Privacidad	Datos	★★★★☆	Differential Privacy + deduplicación	
Sycophancy	Ajuste (RLHF)	★★★★☆	RLHF con evaluadores de honestidad	

**Entonces,
¿cómo decido qué usar?**

Cómo elegir tu herramienta o modelo —————→



Marco personal para elegir la herramienta o modelo de IA adecuado

no por su popularidad o precio, sino por su idoneidad para tu necesidad concreta.



La **gobernanza** como núcleo del proceso.

Se traduce en tu propia **conciencia crítica** cuestionándote siempre la privacidad de: **tus datos,** **la fiabilidad de los resultados,** **el costo de las herramientas y** **el impacto ético de lo que usas.**

No es un paso aislado; es una actitud que acompaña cada decisión.

0. El eje central: Gobernanza personal

Paso 1: Identifica claramente tu caso de uso

El error más común es **buscar una herramienta** antes de definir con precisión **qué se quiere lograr**.

¿Qué tarea concreta quiero resolver?

Hay que ser específico (**NO instrucción amplia**).

Imagina el resultado ideal y trabaja hacia atrás.

Pregúntate: si la IA ya hubiera hecho mi tarea, **¿cómo se vería ese resultado?**

¿Qué le preguntaría exactamente?

Lo que debes hacer en este paso:

Escribe en una frase la tarea que quieres resolver. Luego desglósala en preguntas o instrucciones concretas que le darías a la IA, es decir

"convertir el caso de uso en prompts". (se pueden necesitar diferentes herramientas)

Esa especificidad cambia radicalmente qué herramienta necesitas.

Aspectos a tener en cuenta:

La especificidad del lenguaje importa: Un mismo prompt puede ser interpretado de formas muy distintas. **Cuanto más contextualices tu instrucción, más preciso será el resultado.**

Piensa en la frecuencia: ¿es una tarea puntual o algo que harás todos los días? Esto afecta si vale la pena pagar por una herramienta de pago.

Considera si la tarea es sensible: si involucra datos personales, información confidencial o contenido privado, esto condicionará fuertemente los pasos siguientes.

Casos de uso: <https://bit.ly/4cWK5wW>

Escritura: Redactar, corregir estilo, traducir textos

Investigación: Resumir documentos, buscar fuentes, sintetizar literatura

Programación: Generar o depurar código

Creatividad: Generar imágenes, música, guiones

Organización: Gestionar agenda, tomar notas, planificar proyectos

Aprendizaje: Explicar conceptos, crear ejercicios, responder duda

Identificar
el caso de
uso



Paso 2: Enumera las opciones disponibles

Mapea todas las opciones posibles antes de comprometerte con alguna.

La oferta de herramientas de IA es enorme y cambia con rapidez.

Lo que debes hacer en este paso:

Haz una **lista de herramientas o modelos** que, al menos en apariencia, **sirvan para tu caso de uso**. Clasifícalas según: si son de uso general o especializadas, gratuitas o de pago, y si requieren conocimientos técnicos o no, cantidad de Tokens, ventana de contexto, parámetros.

Conversación,
Análisis de
documentos,
Síntesis

 Claude

Uso general,
Redacción,
Código

 ChatGPT

Integración con
Servicios Google,
Multimodal

 Gemini

Texto, escritura e investigación

 NotebookLM

Análisis profundo de
Documentos propios.

 perplexity

Búsqueda con fuentes
Citadas, ideal para
Investigación.

 Claude

 ChatGPT

Explicación y
generación de
código.

 CURSOR

Editor de código
con IA integrada.

 GitHub Copilot

Asistente dentro de
editores de código.

Código

Aspectos a tener en cuenta:

Modelos de acceso: muchas herramientas tienen versiones gratuitas limitadas y versiones de pago. Identifica si el plan gratuito cubre tu caso de uso.

Disponibilidad en tu región: algunas herramientas tienen restricciones geográficas o no están optimizadas para el español.

Integración con tus herramientas existentes: ¿la IA puede conectarse con las aplicaciones que ya usas (Google Drive, Word, correo, etc.)?

Audio y Video

 ChatGPT

Transcripción de
audio.

 ElevenLabs

Síntesis y clonación
de voz.

 runway

Edición y
generación de
video.

Organización y Productividad

 Copilot

Integrado en
herramientas de
Office (Word, Excel,
Teams).

 Notion

Notas y gestión de
proyectos con IA.

Integrado en
ChatGPT,
Accesible y
versátil.

 DALL-E

Integrado en
herramientas de
Adobe, enfoque
profesional.

 Adobe Firefly

Generación de imágenes



Midjourney

Alta calidad
artística.

 Gemini

Generación y edición de
imágenes y es conocido
por su precisión y
fotorrealismo.

Enumerar
las opciones



Paso 3: Evalúa el rendimiento de cada opción

Tres criterios clave: **precisión, fiabilidad y velocidad**. Evalúa qué tan bien hace la herramienta lo que promete, si lo hace de **forma consistente**, y si **responde con la rapidez** que necesitas.

Lo que debes hacer en este paso:

Diseña una **prueba simple pero representativa** de tu tarea real. No evalúes la herramienta con **preguntas genéricas**; pruébala con **algo parecido a lo que vas a hacer** de verdad.

Criterios adaptados al uso personal:

Precisión - ¿El resultado es correcto y útil? Para tareas de investigación o análisis documental, esto es crítico. Un modelo que "alucina" (inventa datos o fuentes) puede ser más perjudicial que útil. Herramientas como Perplexity o NotebookLM tienen ventaja aquí porque trabajan sobre fuentes verificables.

Fiabilidad - ¿Obtienes resultados similares si repites la misma consulta? ¿El modelo es consistente en su estilo y nivel de calidad? la fiabilidad incluye también la ausencia de sesgos y toxicidad, algo relevante si usas IA en contextos educativos o profesionales.

Velocidad - ¿Cuánto tarda en responder? Para tareas creativas o reflexivas esto puede no importar mucho, pero si necesitas respuestas rápidas mientras trabajas en tiempo real, la latencia importa.

Usa los mismos prompts en diferentes herramientas para comparar resultados de forma justa.

Revisa siempre los resultados críticamente: ningún modelo es infalible.

Consulta reviews y comparativas recientes, ya que los modelos se actualizan frecuentemente.

Criterios adicionales relevantes para uso personal:

Soporte del idioma: no todos los modelos funcionan igual de bien en español. Claude, ChatGPT y Gemini tienen buen rendimiento en español; otros modelos más pequeños o especializados pueden ser menos fluidos.

Ventana de contexto: indica cuánto texto puede procesar el modelo en una sola interacción. Si necesitas analizar documentos largos, necesitas una ventana de contexto amplia (Claude y Gemini destacan aquí).

Multimodalidad: ¿puede procesar imágenes, PDFs, audio? Esto depende de tu caso de uso.

**Evaluar
Rendimiento**



Paso 4: Prueba las opciones seleccionadas

No te quedes con la evaluación teórica. **Prueba de verdad la herramienta con tu tarea real** antes de comprometerte con ella, especialmente si implica pago.

Definir (Prompt base)

Iteración Progresiva: Inicia con una instrucción básica y añade capas de contexto, formato y tono de manera incremental.

Testear (Variaciones)

Ajuste de Sensibilidad: Realiza pequeñas variaciones en el texto para entender cómo reacciona el modelo y optimizar la precisión.

Complementar (Uso de múltiples herramientas).

Integración Multimodular: Evalúa si el resultado mejora combinando diferentes herramientas (por ejemplo, una para investigación y otra para redacción).

Probar las opciones



Aspectos a tener en cuenta:

No descarte una herramienta demasiado rápido: muchas veces un resultado pobre **se debe a un prompt mal formulado**, no al modelo en sí.
Documenta tus pruebas: anota qué prompts funcionaron mejor en cada herramienta; eso te ahorra tiempo en el futuro.
Verifica siempre los datos factuales que genere cualquier modelo antes de usarlos en contextos profesionales o académicos.

Consejos prácticos de prompting (mejores resultados):

Da contexto sobre quién eres y para qué sirve el resultado: "Soy bibliotecólogo y necesito un resumen ejecutivo de este artículo para presentarlo a un comité académico ... "

Especifica el formato de salida: "Responde en formato de tabla/ en tres párrafos/ con viñetas / en menos de 200 palabras"

Pide al modelo que razone paso a paso cuando necesites análisis complejos.

Incluye ejemplos de lo que esperas si el resultado inicial no es satisfactorio.

Paso 5: Elige la opción que te aporte más valor

La **mejor opción** no siempre es la **más potente ni la más cara**, sino la que mejor equilibra rendimiento, velocidad y costo para tu caso de **uso específico**.

Evalúa cada opción candidata contra estos tres ejes:

Costo vs. beneficio real - ¿El plan de pago añade capacidades que realmente vas a usar? Una herramienta gratuita bien utilizada puede superar a una de pago usada a medias (el tamaño y el precio no garantizan el mejor resultado).

Adecuación al caso de uso - Prioriza la herramienta que mejor resuelva tu tarea específica, no la que tenga mejores puntuaciones generales en benchmarks.

Sostenibilidad del uso - ¿Es una herramienta que podrás seguir usando a largo plazo? Considera si el precio es asumible, si la interfaz es cómoda para ti y si se integra bien en tu flujo de trabajo.

Aspectos a tener en cuenta:

La elección no es definitiva: la selección de modelos es un proceso continuo. Las herramientas evolucionan, tus necesidades cambian y nuevas opciones aparecen regularmente. Revisa tu elección periódicamente.

Puedes usar más de una herramienta: no es necesario elegir una sola. Muchos usuarios avanzados combinan herramientas según la tarea del momento.

Considera el aprendizaje acumulado: una vez que dominas una herramienta y sus prompts, cambiar tiene un costo de adaptación. Solo hazlo si la ganancia es significativa.

Esquema de decisión simplificado:

Precisión ¿Los resultados son correctos y útiles para mi tarea?

Fiabilidad ¿Es consistente y sin sesgos problemáticos?

Velocidad ¿Responde con la rapidez que necesito?

Costo ¿Justifica el precio el valor que me aporta?

Privacidad ¿Puedo usarla con mis datos sin riesgos?

Idioma ¿Funciona bien en español?

Accesibilidad ¿Es fácil de usar sin conocimientos técnicos?

Elegir el
modelo

+

\$

Herramientas (no únicas) para elegir un modelo



<https://bit.ly/4w7aMrx>

Es un ranking (leaderboard) de modelos de inteligencia artificial.

Sirve para:

Comparar modelos como GPT, Claude, Gemini, etc.

Ver cuál responde mejor en tareas reales (texto, código, imágenes, etc.)

Evaluar rendimiento con base en votos y comparaciones de usuarios

En la práctica, funciona como un “ranking tipo FIFA pero de IA”.

¿cuál responde mejor?



<https://bit.ly/42amTX6>

Es un agregador de benchmarks de modelos de IA (LLM).

En concreto:

Junta resultados de pruebas como MMLU, GSM8K, HumanEval, etc.

Compara modelos en función de: rendimiento velocidad costo capacidad (context window)

Es básicamente un “tablero técnico” más que una arena de usuarios.

¿cuál saca mejor puntaje técnico?

Ningún leaderboard sabe lo que tú necesitas.

Debes crear:

20–50 prompts reales de tu uso

Casos buenos / malos

Criterios claros:

Precisión

Claridad

Utilidad

tiempo/costo

Luego pruebas:

2–3 modelos

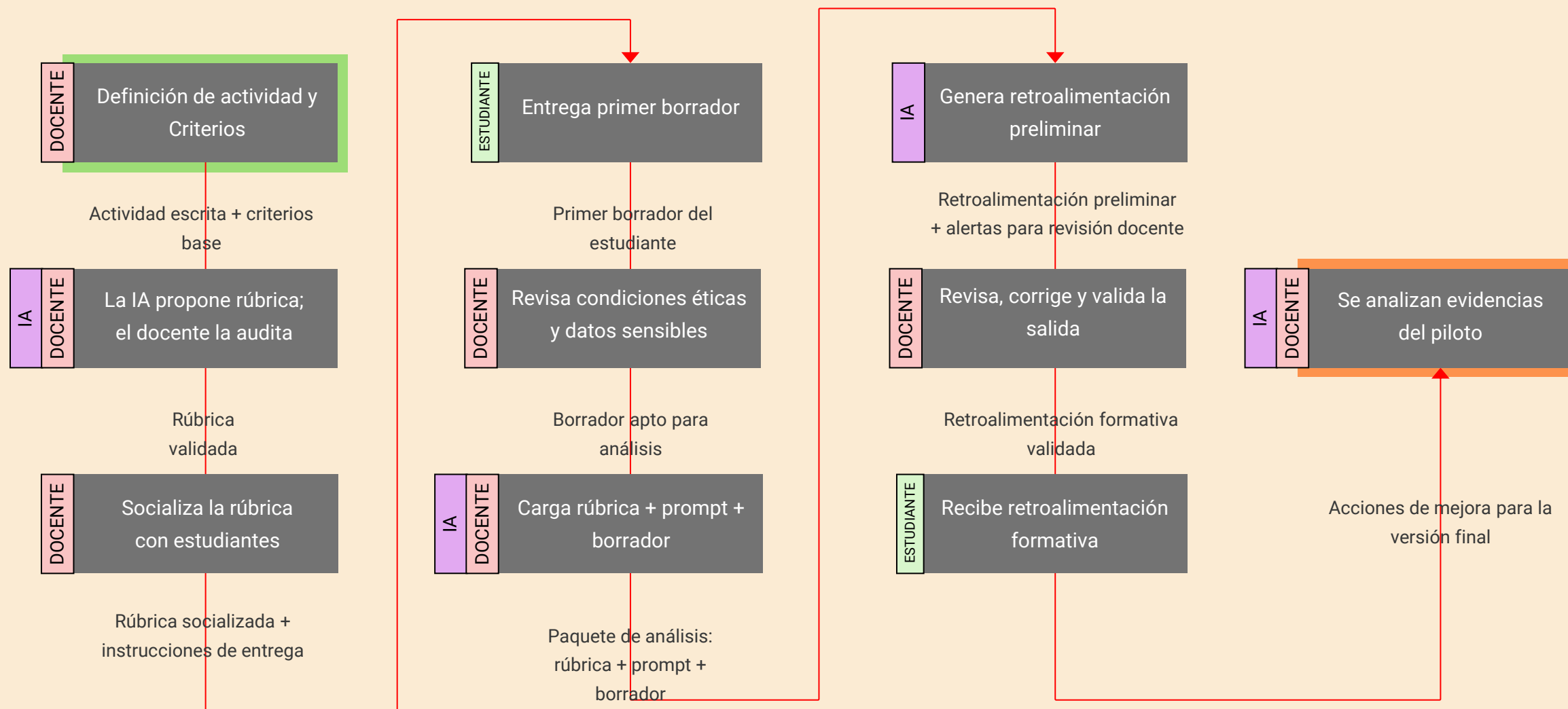
Mismo prompt

Comparas resultados

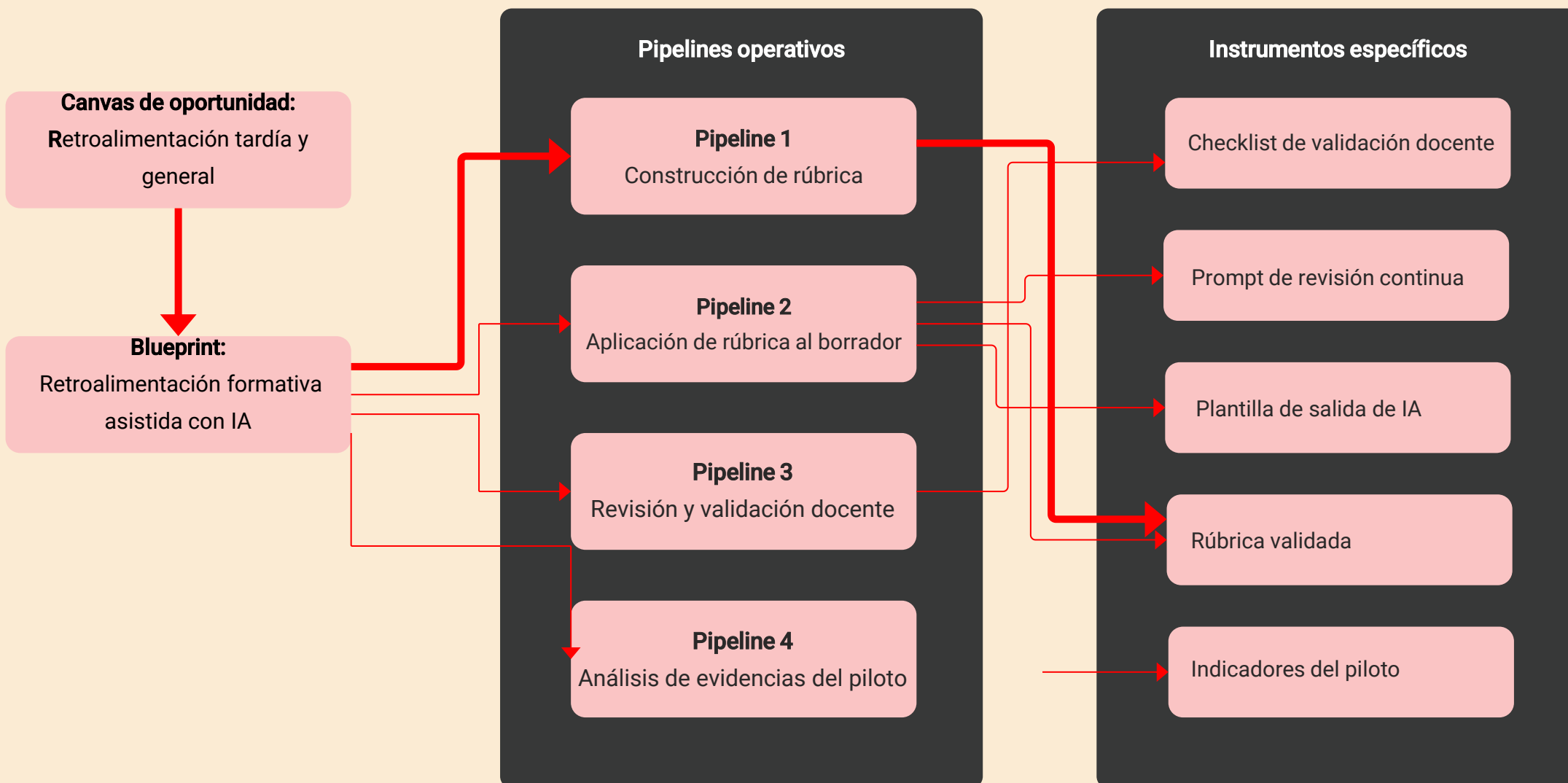
"El problema primero, la herramienta después"

Canvas de Oportunidad →

Del Canvas al Blueprint

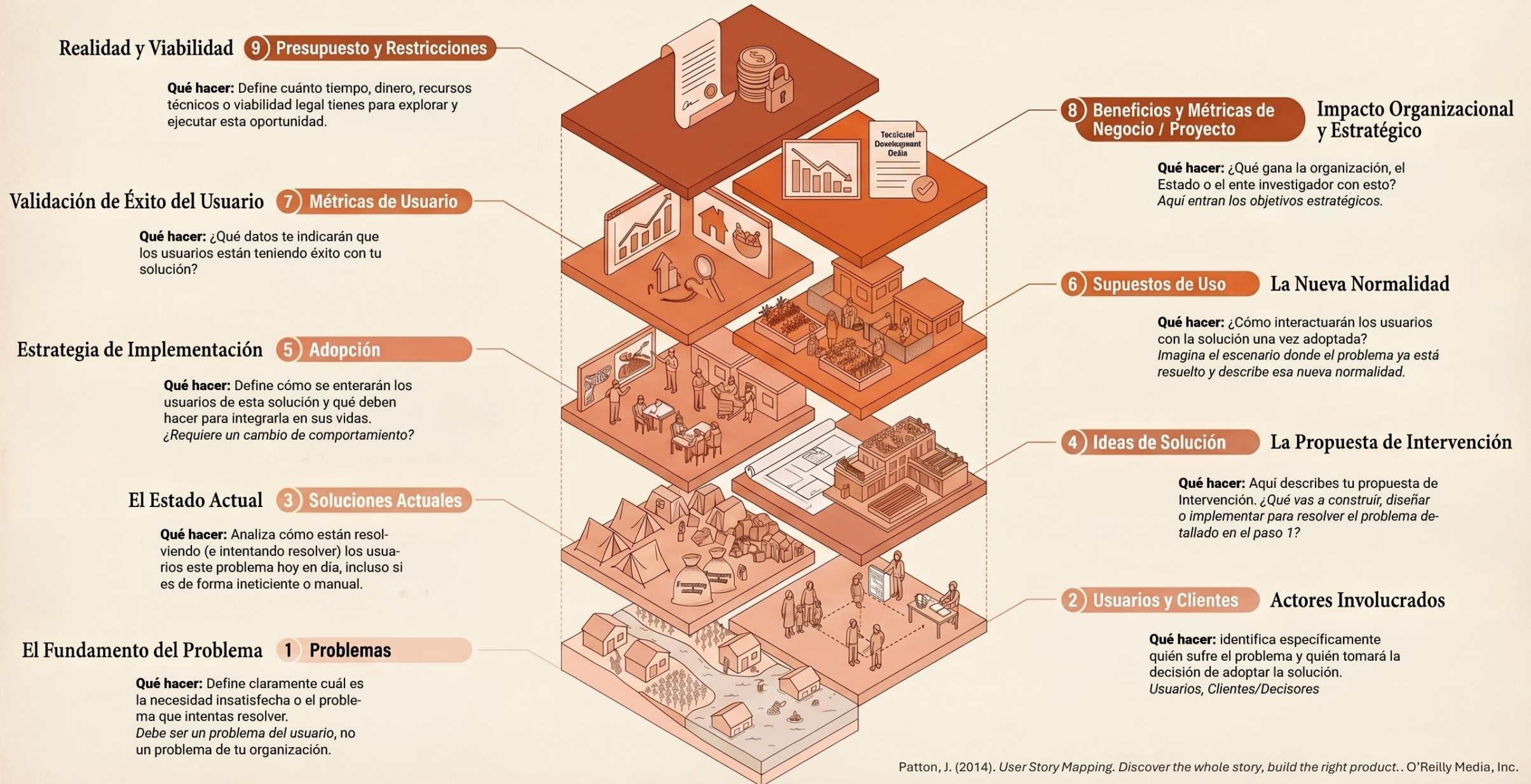


Entendiendo el Blueprint



El canvas de oportunidad

Arquitectura para la Toma de Decisiones Estratégicas



Aprendizaje

Enseñanza

Didáctica

Metodología

Evaluación / Retroalimentación

Currículo

1 EL FUNDAMENTO DEL PROBLEMA
Problemas

2 ACTORES INVOLUCRADOS
Usuarios y Clientes

3 EL ESTADO ACTUAL
Soluciones Actuales

4 LA PROPUESTA DE INTERVENCIÓN
Ideas de Solución

5 ESTRATEGIA DE IMPLEMENTACIÓN
Adopción

6 LA NUEVA NORMALIDAD
Supuestos de Uso

7 VALIDACIÓN DE ÉXITO DEL APRENDIZAJE
Métricas de Usuario

8 IMPACTO ACADÉMICO E INSTITUCIONAL
Beneficios y Métricas

9 REALIDAD Y VIABILIDAD
Presupuesto y Restricciones

Los pasos 1 a 3 son de diagnóstico: el docente identifica el problema real (no el síntoma), reconoce quiénes son sus actores clave y analiza lo que ya se está haciendo —aunque sea de forma reactiva o ineficiente.

Los pasos 4 a 6 son de diseño e intervención: se propone una solución didáctica fundamentada (como la alineación constructiva de Biggs), se piensa cómo implementarla con sus implicaciones de cambio, y se imagina la nueva normalidad del aula una vez adoptada.

Los pasos 7 a 9 son de validación y viabilidad: se definen los indicadores de éxito del aprendizaje, el impacto institucional medible (útil para acreditaciones y PEI), y se reconocen con honestidad las restricciones curriculares, normativas y de recursos.

Aprendizaje

Enseñanza

Didáctica

Metodología

Evaluación / Retroalimentación

Currículo

1

EL FUNDAMENTO DEL PROBLEMA

Problemas

QUÉ HACER

Define claramente cuál es el punto de dolor pedagógico, la necesidad insatisfecha o la fricción que intentas resolver como docente. Debe ser un problema del estudiante o del proceso de aprendizaje, no un problema administrativo tuyo.

CASO EDUCATIVO

Desalineación entre objetivos, actividades y evaluación en el syllabus

Ejemplo aplicado:

El docente detecta que sus estudiantes memorizan contenidos pero no desarrollan competencias prácticas. Las actividades del syllabus no corresponden a los resultados de aprendizaje declarados, generando frustración y bajo desempeño al final del semestre.

Aprendizaje

Currículo

El problema del paso 1 se formula desde el punto de vista del estudiante —no de la administración—, ya que eso orienta todas las decisiones posteriores hacia el aprendizaje real y no hacia el cumplimiento formal del syllabus.

Aprendizaje

Enseñanza

Didáctica

Metodología

Evaluación / Retroalimentación

Currículo

2

ACTORES INVOLUCRADOS

Usuarios y Clientes

QUÉ HACER

Identifica específicamente quién sufre el problema pedagógico (usuarios) y quién tomará la decisión de adoptar o avalar la solución (clientes/decisores).

CASO EDUCATIVO

Perfiles del aula y del programa académico

Ejemplo aplicado:

Usuarios: estudiantes que no alcanzan los resultados de aprendizaje esperados; docentes que perciben bajo rendimiento y desmotivación. Clientes/Decisores: coordinadores de área, comités curriculares, decanaturas y, en algunos casos, el propio estudiante como agente de su proceso.

Enseñanza

Currículo

¿Quién sufre el problema pedagógico (**usuario**)?

¿Quién tomara la decisión de adoptar la solución (**cliente/decisor**)?

Aprendizaje

Enseñanza

Didáctica

Metodología

Evaluación / Retroalimentación

Currículo

3

EL ESTADO ACTUAL

Soluciones Actuales

QUÉ HACER

Analiza cómo están resolviendo (e intentando resolver) los actores este problema pedagógico hoy en día, incluso si es de forma ineficiente, informal o reactiva.

CASO EDUCATIVO

Respuestas reactivas y parches pedagógicos

Ejemplo aplicado:

Clases de refuerzo esporádicas sin diagnóstico previo; ajustes informales al syllabus sin metodología curricular; uso de los mismos instrumentos de evaluación sin revisión; repetición de contenidos sin identificar la causa del problema de aprendizaje.

Evaluación / retroalimentación

Metodología

¿Qué hacen hoy los usuarios para resolver o evadir el problema?

Aprendizaje

Enseñanza

Didáctica

Metodología

Evaluación / Retroalimentación

Currículo

4

LA PROPUESTA DE INTERVENCIÓN

Ideas de Solución

QUÉ HACER

Describe tu propuesta pedagógica de intervención. ¿Qué vas a rediseñar, implementar o transformar para resolver el problema detallado en el paso 1?

CASO EDUCATIVO

Modelo de syllabus con alineación constructiva

Ejemplo aplicado:

Rediseño del syllabus aplicando el modelo de alineación constructiva de Biggs: los objetivos de aprendizaje, las estrategias didácticas activas (ABP, aula invertida, estudios de caso) y los instrumentos de evaluación formativa se articulan coherentemente hacia las competencias del perfil de egreso.

Didáctica

Metodología

Evaluación / retroalimentación

El objetivo **no es listar todo lo que se podría hacer**, sino proponer **intervenciones precisas que ataquen directamente** el punto de fricción o problema detectado.

Aprendizaje

Enseñanza

Didáctica

Metodología

Evaluación / Retroalimentación

Currículo

5 ESTRATEGIA DE IMPLEMENTACIÓN Adopción

QUÉ HACER

Define cómo se enterarán los estudiantes y colegas de esta solución y qué cambios de comportamiento o práctica implica su adopción en el aula y en el programa.

CASO EDUCATIVO

Integración pedagógica progresiva

Ejemplo aplicado:

Socialización del nuevo syllabus con los estudiantes en la primera sesión del curso, explicando el enfoque y los criterios de evaluación. Formación docente en diseño instruccional y participación en comunidades de práctica del departamento. Uso de un LMS (Moodle, Canvas) como plataforma de gestión del aprendizaje.

Enseñanza

Didáctica

Determinar **cómo los estudiantes (usuarios)** y el **comité/docentes (clientes)** descubrirán e **integrarán esta nueva metodología en sus dinámicas**, y qué fricciones de cambio de comportamiento deben superarse.

Aprendizaje

Enseñanza

Didáctica

Metodología

Evaluación / Retroalimentación

Currículo

6

LA NUEVA NORMALIDAD

Supuestos de Uso

QUÉ HACER

¿Cómo interactuarán los estudiantes y el docente con la solución una vez adoptada? Imagina el escenario donde el problema ya está resuelto y describe esa nueva normalidad pedagógica.

CASO EDUCATIVO

Corresponsabilidad y ajuste continuo

Ejemplo aplicado:

Los estudiantes navegan rutas de aprendizaje claras con criterios transparentes. El docente monitorea el progreso mediante retroalimentación formativa continua y ajusta el syllabus dinámicamente en ciclos semestrales, usando los resultados de evaluación como insumo para rediseñar la siguiente versión del curso.

Aprendizaje

Metodología

Visualizar el "día a día" del **estudiante** interactuando con la solución. **¿Cómo se ve el éxito operativo una vez que el nuevo sistema está en marcha?**

Aprendizaje

Enseñanza

Didáctica

Metodología

Evaluación / Retroalimentación

Currículo

7

VALIDACIÓN DE ÉXITO DEL APRENDIZAJE

Métricas de Usuario

QUÉ HACER

¿Qué datos e indicadores te dirán que los estudiantes están teniendo éxito con tu solución pedagógica? Define métricas observables y medibles.

CASO EDUCATIVO

Indicadores de logro por competencias

Ejemplo aplicado:

Tasa de alcance de competencias por cohorte al finalizar el semestre; niveles de participación activa en dinámicas de aula; calidad de los productos de aprendizaje entregados (rúbricas); resultados comparativos entre evaluaciones diagnósticas y sumativas; percepción de los estudiantes mediante encuestas de satisfacción pedagógica.

Evaluación / retroalimentación

Aprendizaje

¿Qué **datos nos indicarán** que los estudiantes **están adoptando la solución** y que su fricción cognitiva y metodológica realmente está disminuyendo?

Aprendizaje

Enseñanza

Didáctica

Metodología

Evaluación / Retroalimentación

Currículo

8

IMPACTO ACADÉMICO E INSTITUCIONAL

Beneficios y Métricas

QUÉ HACER

¿Qué gana el docente, el programa y la institución con esta intervención pedagógica? Aquí entran los objetivos estratégicos del área y de la institución.

CASO EDUCATIVO

Optimización del programa académico

Ejemplo aplicado:

Reducción de tasas de reprobación y deserción en la asignatura; mejora en indicadores de calidad para procesos de acreditación institucional; cumplimiento demostrable de los resultados de aprendizaje del programa; alineación con el Proyecto Educativo Institucional (PEI) y los estándares de competencia del campo disciplinar.

Currículo

Enseñanza

¿Qué gana el programa, la
Universidad La Gran Colombia al
respaldar esta intervención?

Aprendizaje

Enseñanza

Didáctica

Metodología

Evaluación / Retroalimentación

Currículo

9

REALIDAD Y VIABILIDAD

Presupuesto y Restricciones

QUÉ HACER

Define cuánto tiempo, recursos didácticos y viabilidad institucional tienes para explorar y ejecutar esta intervención pedagógica. ¿Cuáles son las limitaciones reales?

CASO EDUCATIVO

Limitaciones pedagógicas e institucionales

Ejemplo aplicado:

Restricciones del plan de estudios aprobado que limitan cambios sin aval del comité curricular; número de horas de contacto disponibles vs. horas necesarias para nuevas metodologías; acceso desigual de los estudiantes a tecnología; disponibilidad del docente para formación en nuevas estrategias didácticas; normatividad institucional sobre modificaciones al syllabus.

Currículo

Metodología

Aquí es donde la **propuesta choca con la realidad institucional**. Analizar el presupuesto y las restricciones.

Tienes en tus manos un **Canvas de Oportunidad** completamente articulado y basado en evidencia documental.

La IA predice, no piensa.

Úsela con criterio.

El problema primero,

la herramienta después.

Su juicio no es reemplazable.

**Es su principal activo con
estas herramientas.**

¡Gracias!