

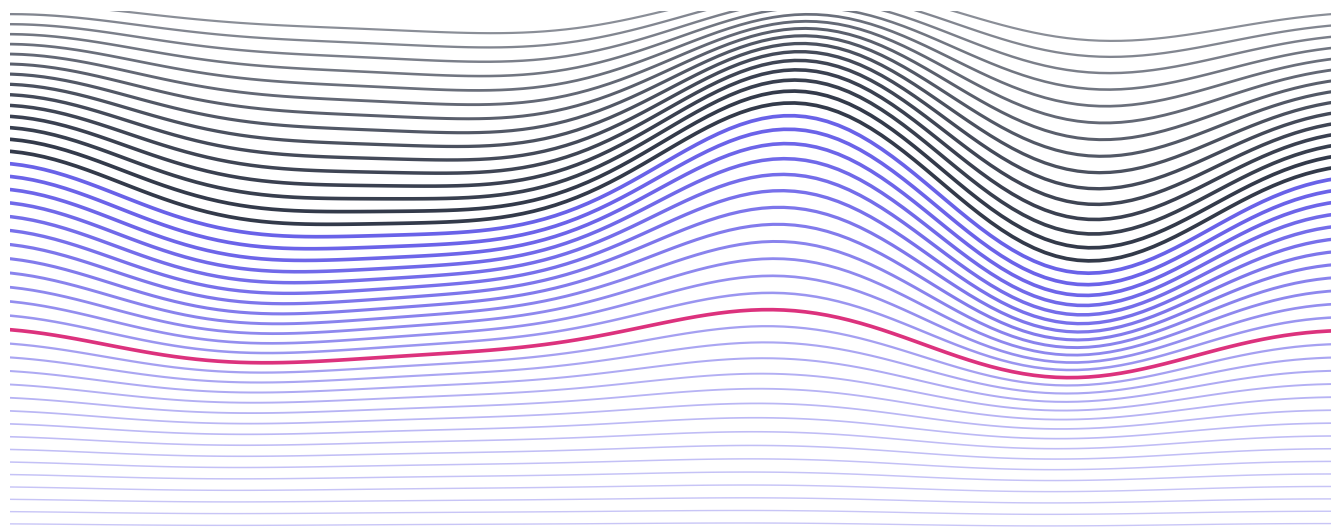


ORIENTACIONES PARA EL USO PEDAGÓGICO DE LA IA

Anexo C

Glosario de términos

Material de consulta no normativo. Acompaña al documento principal; no crea normas ni presenta decisiones oficiales de la Universidad de Guadalajara. Las remisiones con § corresponden a sus secciones.



Contenido

Anexo C · Glosario de términos	3
Referencias del anexo C	36

ANEXO

C

Glosario de términos

Este glosario reúne los términos de estas Orientaciones y del sitio IA y Aprendizaje. Permite recuperar una definición y reconocer relaciones entre conceptos durante la consulta. Las explicaciones, los casos y las actividades desarrollados en el documento y en el sitio amplían ese vocabulario.

Agencia humana. Capacidad de la persona de orientar el proceso, decidir y responder por el resultado cuando trabaja con IA. No desaparece automáticamente porque una herramienta realice una operación: la cuestión pedagógica es qué se delega, con qué propósito y si la persona conserva comprensión y posibilidades de revisión. El riesgo aparece cuando se sustituye precisamente la operación que la actividad buscaba que aprendiera o ejercitara.

En este glosario, la Dirección epistémica describe una forma observable de esa agencia mediante preguntas sobre propósito, comprobación, decisión y responsabilidad. No es sinónimo de Agenciamiento persona-IA, que nombra la configuración completa, ni equivale a prescindir de apoyos técnicos.

Fuentes públicas y alcance: [Guidance for generative AI in education and research](#) (Orientación internacional: agencia humana, interacción pedagógica, proporcionalidad, participación y rendición de cuentas; no es una norma universal ni evidencia causal.)

Ubicación orientativa: §2.2

Agenciamiento persona-IA. El **agenciamiento persona-IA** nombra la configuración completa en la que ocurre un trabajo con IA. No se reduce al par persona-herramienta: incluye la actividad que se realiza, las fuentes disponibles, los criterios con los que se evalúa y las condiciones institucionales que enmarcan la tarea. Cambiar cualquiera de esas piezas cambia lo que la interacción puede producir. El término no significa «trabajo en equipo» ni supone que la persona y el sistema sean equivalentes: cada elemento influye, pero no todos comprenden, deciden o responden de la misma manera, y la responsabilidad no se reparte por partes iguales.

El término procede de la filosofía de Gilles Deleuze y Félix Guattari, donde *agencement* designa un ensamblaje de elementos heterogéneos. La fuente primaria bibliográfica citada en la guía sustenta ese marco conceptual; la aplicación educativa que sigue es una interpretación situada, no una síntesis completa de su obra. Sirve para no atribuir los resultados solo a la herramienta o solo a la persona. Para describir el trabajo cotidiano resulta más transparente hablar de Cocreación persona-IA, que nombra el ciclo que la persona dirige dentro de esa configuración. Ensamblaje cognitivo y Virtual (lo virtual) dialogan con este marco desde otras nociones; no son simplemente términos de la misma tradición.

Fuentes públicas y alcance: [Mil mesetas: capitalismo y esquizofrenia](#) (Deleuze y Guattari (2020), edición bibliográfica citada en la guía para el concepto de *agencement*; la aplicación educativa del glosario es una interpretación situada.)

Ubicación orientativa: §3.1

Agentes de IA. Un **agente de IA** es un componente o sistema que, dentro de límites definidos, puede planificar y encadenar acciones, consultar información, usar herramientas o conservar memoria para cumplir una tarea en varios pasos. Se distingue de una respuesta aislada porque mantiene una tarea y puede actuar por etapas. El término no implica autonomía total: algunos agentes siguen un flujo fijado por sus responsables; otros eligen entre pasos permitidos según lo que encuentran.

Por ejemplo, un agente de apoyo bibliográfico podría buscar tres fuentes, extraer sus datos y preparar una tabla. La persona responsable tendría que aprobar las fuentes antes de incorporarlas a una bibliografía compartida. Dar acceso al catálogo no equivale a autorizar cambios en él.

La distinción importa cuando se asignan permisos. Antes de autorizar un agente, conviene revisar qué herramientas puede abrir, qué datos recibirá, qué acciones necesitan aprobación, cuánto tiempo puede operar y qué registro dejará. Para una tarea breve y predecible suele bastar una instrucción o un flujo fijo. La IA generativa produce respuestas o contenido; un

agente puede añadir acciones y uso de herramientas. Un Tutor inteligente describe una función educativa y puede operar con un flujo fijo o con capacidades de agente. Un Sistema agéntico puede organizar uno o varios agentes junto con herramientas y mecanismos de coordinación.

Alfabetización en IA. La **alfabetización en IA** (*AI literacy*) reúne capacidades cognitivas, técnicas y éticas para comprender qué puede y qué no puede hacer un sistema de IA, cómo trabaja con información y qué efectos puede tener su uso. Se relaciona con la alfabetización digital, pero no es necesariamente una etapa superior: depende de la tarea, de la experiencia previa y de las herramientas disponibles. No se trata solo de aprender a usar herramientas, sino de poder evaluar cuándo es apropiado delegar una tarea a la IA y cuándo resulta contraproducente para el aprendizaje; implica entender nociones como el entrenamiento de los modelos, la privacidad de los datos que se entregan y la imperfección de las respuestas.

Como proceso de formación, las Orientaciones la organizan en tres Literacidades operativa, crítica y de cocreación: una operativa, para usar una herramienta con un propósito delimitado; una crítica, para examinar sus respuestas, datos y efectos; y una de cocreación, para integrar sus aportes sin perder la dirección del trabajo. Como trayecto acumulan capacidades que se ejercitan desde el inicio y vuelven a combinarse según la tarea. No forman una escalera obligatoria ni equivalen automáticamente a los niveles de la Taxonomía de Bloom: la relación con esa taxonomía expresa énfasis posibles, no operaciones exclusivas de cada literacidad, y ninguna de ellas se demuestra solo por usar IA. El Doble imperativo recuerda que la formación aplicada necesita ir acompañada de capacidades humanas que la herramienta no reemplaza.

Fuentes públicas y alcance: [Guidance for generative AI in education and research](#) (Orientación internacional: agencia humana, interacción pedagógica, proporcionalidad, participación y rendición de cuentas; no es una norma universal ni evidencia causal.)

Ubicación orientativa: §3 y §6.7

Alienación epistémica. Situación en la que alguien firma como propio un resultado sobre el que ya perdió el control interpretativo, es decir, cuando ya no puede explicar por qué su trabajo dice lo que dice. Es una de las formas extremas de pérdida de la Dirección epistémica que la investigación ha nombrado; la otra es la Subyugación algorítmica. Los nombres no cambian el fenómeno, pero permiten conversar sobre él con precisión y advertir que el paso de la cocreación dirigida a estas formas de desplazamiento rara vez se decide de manera consciente.

Ubicación orientativa: §3.1

Aprendizaje activo. El **aprendizaje activo** (*active learning*) reúne situaciones en las que la persona hace algo cognitivamente significativo con el contenido: selecciona información, explica relaciones, compara alternativas, formula una pregunta o justifica una decisión. Participa, resuelve, contrasta, discute y reflexiona; no se limita a recibir información. Tampoco se define solo por participar, moverse o usar una herramienta.

Una actividad visible, como subrayar, hacer clic u ordenar tarjetas, puede quedarse en manipulación superficial. Para saber si hay aprendizaje activo conviene observar qué operación realiza la persona y qué evidencia deja: una explicación, una conexión pertinente, una revisión razonada o una conclusión apoyada en el material. El marco ICAP ayuda a distinguir esos modos de implicación.

La IA puede proponer ejemplos, preguntas, contraargumentos o retroalimentación, pero su presencia no vuelve activa una tarea. El criterio sigue siendo lo que la persona debe comprender, producir, comprobar y revisar; también puede diseñarse una alternativa sin IA que conserve esas operaciones. Esta orientación recupera la tradición de la Educación dialógica.

Fuentes públicas y alcance: [Active-Constructive-Interactive: A Conceptual Framework for Differentiating Learning Activities](#) (Marco conceptual para diferenciar actividad pasiva, activa, constructiva e interactiva; la clasificación de conducta no prueba por sí sola aprendizaje.); [The ICAP framework: Linking cognitive engagement to active learning outcomes](#) (Marco ICAP de implicación cognitiva y resultados de aprendizaje; usar como lente de diseño, no como garantía de resultado.)

Ubicación orientativa: §2.4 y §4.3

Aprendizaje basado en problemas. El **aprendizaje basado en problemas** (ABP) es una metodología en la que un problema situado, sin una respuesta única evidente, orienta la investigación y el aprendizaje. El grupo identifica qué necesita saber, busca y contrasta información, propone una respuesta y explica por qué la considera razonable. El problema no es un pretexto para adivinar una solución: permite usar conceptos y revisar las decisiones que se toman con ellos.

La IA puede ayudar a formular preguntas, mostrar escenarios alternativos o señalar vacíos en una primera propuesta. Su respuesta no sustituye la indagación ni decide cuál evidencia es pertinente para el caso. Si se usa, conviene que cada estudiante o equipo pueda mostrar qué información comprobó, qué sugerencia descartó o transformó y cómo llegó a la decisión que defiende.

Aprendizaje digital. Cuando una clase depende de una plataforma, una simulación, un documento compartido o una herramienta de comunicación, hablamos de **aprendizaje digital** si ese recurso interviene de manera significativa en lo que el grupo hace para aprender. Puede ocurrir en un curso presencial, híbrido o en línea. La tecnología puede servir para acceder a materiales, colaborar, practicar, recibir comentarios o revisar avances.

Subir un PDF a una plataforma puede formar parte de una actividad, pero el archivo por sí solo no explica qué hará el grupo. Por ejemplo, una actividad digital puede pedir que dos estudiantes comparen sus anotaciones sobre un texto, revisen una afirmación con las fuentes disponibles y dejen una versión corregida. Lo importante es la acción que la tecnología permite o facilita, no que todo sea interactivo, automático o personalizado. La expresión tampoco significa que una herramienta mejore el aprendizaje por sí misma: hay que revisar acceso, privacidad, carga de trabajo y relación con lo que se espera aprender.

Aprendizaje híbrido. Una estudiante prepara una explicación en línea, la contrasta durante una sesión presencial y después revisa su texto. El trabajo cambia de espacio, pero mantiene una pregunta y aprovecha lo que se produjo antes.

El **aprendizaje híbrido** (*blended learning*) conecta de manera deliberada momentos presenciales y en línea, síncronos y asíncronos. Cada momento utiliza o transforma lo producido en el anterior; no consiste solo en repartir contenidos entre el aula y una plataforma. Diseñar bien esa combinación exige decidir qué aporta cada entorno y cómo se relacionan las actividades.

Presencial y en línea describen el entorno; síncrono y asíncrono indican si las personas coinciden en el tiempo. Una videollamada es en línea y síncrona; un foro con plazos amplios es en línea y asíncrono. Son decisiones distintas. Tampoco existe un reparto obligatorio en el que la pantalla transmita información y el aula concentre el pensamiento crítico: en ambos entornos se puede explicar, debatir, practicar o revisar. Para diseñar una semana, conviene señalar qué producto se retoma al cambiar de momento y qué hará la persona con él; si el acceso o el horario impiden participar, una alternativa debe conservar esa posibilidad de contrastar y revisar.

Fuentes públicas y alcance: [Blended learning: Uncovering its transformative potential in higher education](#) (Marco de blended learning en educación superior; no prescribe un reparto fijo entre modalidades.)

Ubicación orientativa: §4.6

Aprendizaje transformativo. El **aprendizaje transformativo** describe un proceso en el que una persona no solo incorpora información nueva, sino que revisa supuestos y reorganiza su manera de interpretar un problema o experiencia. Puede reconocerse, por ejemplo, cuando explica

qué creía al inicio, qué evidencia o diálogo puso esa idea en cuestión y cómo cambió su forma de actuar o de justificar una decisión. No todo aprendizaje intenso ni toda revisión de una respuesta alcanza esa transformación.

Un estudio de encuesta con estudiantes de tres regiones encontró una asociación entre las formas de relación pedagógica con IA, la vigilancia crítica, la Descarga cognitiva estratégica y experiencias de aprendizaje transformativo. Ese resultado no establece una receta para cada curso ni prueba que delegar una tarea produzca transformación. En una actividad concreta, la pregunta sigue siendo qué necesita investigar, contrastar y decidir la persona; una IA puede ofrecer un contraste o realizar una operación ya dominada, pero no puede revisar los supuestos en su lugar.

Fuentes públicas y alcance: [Pedagogical partnerships with generative AI in higher education: how dual cognitive pathways paradoxically enable transformative learning](#) (Estudio mixto con estudiantes de educación superior en China, Europa y Estados Unidos (N=912); asociación y modelos de mediación, no receta causal general.)

Asistente de IA. Programa que responde a solicitudes para apoyar tareas concretas. Por ejemplo, puede proponer una explicación, resumir un texto o formular preguntas. La persona todavía necesita juzgar la respuesta. Su función de apoyo no implica que sea incapaz de usar herramientas o de ejecutar secuencias de acciones: algunos productos combinan capacidades de asistente y de agente. Por eso conviene preguntar qué puede hacer el sistema concreto, qué permisos tiene y qué aprobaciones exige, sin fiarse del nombre comercial.

Fuentes públicas y alcance: [Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile](#) (Perfil NIST de riesgos y acciones para IA generativa; definición técnica y marco voluntario, no política educativa obligatoria.)

Aula invertida. En un **aula invertida** (*flipped classroom*), una parte de la preparación ocurre antes del encuentro grupal y se retoma allí para aplicar, contrastar o revisar lo aprendido. Por ejemplo, el grupo puede leer un caso y anotar una duda antes de clase; durante la sesión compara interpretaciones y corrige su respuesta con criterios compartidos. No basta con mover un video fuera del horario: las dos partes deben estar conectadas por una tarea que aproveche lo preparado.

La preparación puede ser una lectura, un video, una demostración o un ejercicio breve, según lo que el grupo necesite. El encuentro tampoco tiene que ser siempre presencial ni convertirse por definición en debate o trabajo colaborativo: puede dedicar tiempo a resolver un problema, recibir retroalimentación o practicar una operación difícil. Una IA puede proponer preguntas de estudio o una explicación alternativa para la preparación individual; si se incorpora, la actividad

posterior debe pedir a la persona que compare esa propuesta con el material del curso, explique su decisión o la revise. Generar material automáticamente no garantiza que haya preparación ni comprensión.

Bitácora de decisiones (decision log). Nota breve en la que la persona registra qué pidió, qué recibió, qué conservó, transformó o descartó y por qué. Es una de las cuatro maneras de conocer el Recorrido de un trabajo. No necesita reproducir una conversación completa ni convertirse en vigilancia: puede consistir en decisiones breves, versiones comparadas, fuentes comprobadas y una explicación final. Cuando la actividad lo amerita, adopta la forma de un Registro reflexivo selectivo.

Ubicación orientativa: §5.4

Brecha entre evaluación controlada y práctica real. Distancia entre lo que los sistemas de IA logran en exámenes y pruebas de referencia y lo que aportan en condiciones reales de trabajo. Los informes del periodo la documentan de manera reiterada en varios campos profesionales, y respalda la postura de examinar cada uso en su contexto en lugar de inferir utilidad educativa a partir del desempeño de un sistema en una prueba.

Fuentes públicas y alcance: [Measuring the Impact of Early-2025 AI on Experienced Open-Source Developer Productivity](#) (RCT con 16 desarrolladores y 246 tareas en repositorios maduros; halló 19% más tiempo con IA; no representa a toda la programación ni a empleos de entrada.); [Generative AI and jobs: A 2025 update](#) (Índice global de exposición ocupacional basado en tareas, expertos y predicciones; estima potencial de transformación, no automatización efectiva de puestos novatos.)

Ubicación orientativa: §2.3 y anexo A

Cocreación persona-IA. Hay **cocreación persona-IA** cuando la respuesta de un sistema no se copia ni se entrega tal cual, sino que entra en un ciclo intelectual dirigido por la persona: ella la interroga, la contrasta con otras fuentes, decide qué conservar o descartar y transforma su propio trabajo con lo que aprendió en ese intercambio. La herramienta puede proponer preguntas, formulaciones, objeciones o alternativas, mientras la persona orienta el proceso, comprueba y decide.

El término no convierte a la IA en coautora en el sentido humano ni le transfiere responsabilidad alguna. Describe una relación de trabajo, no un reparto de méritos: la herramienta introduce posibilidades y la persona conserva la Dirección epistémica, es decir, el propósito, la comprobación y la decisión final. Cuando ese ciclo se rompe y la persona solo recibe y valida, ya no hay cocreación sino delegación, con los riesgos que describe la Descarga cognitiva. Las for-

mas extremas de esa pérdida tienen nombre: Subyugación algorítmica y Alienación epistémica. La cocreación ocurre dentro de un Agenciamiento persona-IA, pero nombra el trabajo, no la configuración.

Ubicación orientativa: §3.1

Colonización cognitiva. La **colonización cognitiva** es el riesgo de adoptar como universales las categorías, los ejemplos y los valores incorporados en sistemas de IA construidos en otros contextos, desplazando saberes y formas locales de nombrar, investigar y comprender. Cuando los datos, las fuentes o los criterios de un sistema no incluyen de manera suficiente ciertas variantes del español, lenguas originarias o saberes locales, sus respuestas pueden representarlos de forma pobre, equivocada o simplemente omitirlos.

El riesgo no está en usar esas herramientas, sino en dejar de notar lo que imponen. Si una comunidad acepta sus categorías como si fueran neutrales, puede ir desplazando sus propias formas de conocer. La pregunta protectora es concreta: antes de incorporar una clasificación o un criterio sugerido por un modelo, ¿tiene sentido para esta comunidad y quién quedó fuera de sus datos? Ese cuidado conecta el Sesgo algorítmico con la defensa de la Justicia epistémica.

Ubicación orientativa: §6.2 y §8.5

Criterios para elegir tecnología. Cuatro criterios con nombre propio para comparar opciones tecnológicas. La *portabilidad* y la *interoperabilidad* piden que datos y contenidos puedan exportarse y migrarse sin pérdida. La *transparencia* prefiere sistemas cuyo funcionamiento, límites y condiciones puedan conocerse. El *costo total en el tiempo* considera lo que cuesta permanecer y lo que costaría salir, no solo el precio inicial. La *reversibilidad* exige poder cambiar de proveedor o internalizar el servicio sin perder el trabajo acumulado. Ninguno garantiza por sí solo una migración sencilla; juntos permiten reconocer dependencias antes de que se acumulen y forman parte de la Soberanía tecnológica. La Proporcionalidad añade el costo material a esa comparación.

Fuentes públicas y alcance: [Guidance for generative AI in education and research](#) (Orientación internacional: agencia humana, interacción pedagógica, proporcionalidad, participación y rendición de cuentas; no es una norma universal ni evidencia causal.)

Ubicación orientativa: §8.2

Declaración de uso de IA. Nota breve que indica la herramienta utilizada, el propósito y la etapa del uso, y confirma que la persona revisó el resultado y responde por él. Puede bastar con señalar qué herramienta se utilizó, para qué se consultó, qué se comprobó y qué decisiones que-

daron en manos de la persona. No debería funcionar como una confesión de culpa, sino como una parte normal de la descripción del proceso; pedirla sin enseñar cómo hacerla deja una regla difícil de cumplir.

Los comités editoriales de referencia y las políticas de las grandes editoriales científicas usan un formato de este tipo, con una distinción operativa útil: la corrección de estilo asistida no suele requerir declaración, mientras que la generación de contenido sí, siempre con revisión humana y responsabilidad íntegra de quien firma. Algunas escuelas de posgrado piden una declaración estructurada en tesis y protocolos. Una fórmula semejante puede adaptarse a los trabajos escolares, de modo que el estudiantado practique desde el aula la transparencia que le pedirá su vida académica y profesional. Declarar es distinto de hacer visible el Recorrido: la declaración dice que intervino la IA; el recorrido muestra cómo.

Fuentes públicas y alcance: [Guidance for generative AI in education and research](#) (Orientación internacional: agencia humana, interacción pedagógica, proporcionalidad, participación y rendición de cuentas; no es una norma universal ni evidencia causal.)

Ubicación orientativa: §5.7 y §6.1

Defensa oral. Conversación evaluativa en la que la persona sostiene su trabajo, explica sus decisiones y responde preguntas. Es una de las cuatro maneras de conocer el Recorrido. No demuestra por sí sola todo el aprendizaje, pero ayuda a distinguir una comprensión propia de un producto que la persona no puede sostener. En el posgrado, el examen de grado recupera este papel cuando la fluidez de un texto ya no demuestra por sí sola el esfuerzo intelectual de quien lo firma; la literatura propone convertirlo de confirmación ceremonial en indagación que pida razonamiento propio.

Ubicación orientativa: §5.4 y §5.7

Dependencia de proveedor (vendor lock-in). La **dependencia de proveedor** (*vendor lock-in*) es la situación en la que cambiar de servicio se vuelve impracticable: los datos y el trabajo acumulado no pueden recuperarse en un formato utilizable, o el costo de migrar, en dinero, tiempo o formación, supera lo que la institución puede asumir. La dependencia rara vez se decide; se acumula con cada integración que no previó una salida.

Para reducirla conviene comparar opciones antes de adoptar una herramienta con los Criterios para elegir tecnología: portabilidad e interoperabilidad, transparencia, costo total en el tiempo y reversibilidad. Preguntar cómo se cambiaría de proveedor antes de entrar es la forma más sencilla de no tomar, sin decirlo, la decisión de no poder salir. Estos cuidados forman parte de la Soberanía tecnológica.

Ubicación orientativa: §8.2

Deriva metacognitiva. La **deriva metacognitiva** es la pérdida gradual de la capacidad de notar cuánto se ha dejado de pensar al delegar en una IA. La persona conserva una imagen optimista de su competencia mientras cede el control del proceso: ya no identifica qué parte resolvió, qué recibió y qué necesitaría revisar por cuenta propia. No nombra cualquier uso repetido de una herramienta ni permite diagnosticar una pérdida de capacidad a partir del producto final; lo característico es que la delegación se vuelva opaca.

La deriva se aloja en la distancia entre lo que una persona cree en general sobre su propio pensar, una imagen relativamente estable, y lo que experimenta al delegar cada paso, una señal cambiante. Por eso conviene intervenir en el momento de delegar, preguntándose qué operación se traslada y si es una que ya se domina, y no únicamente al entregar. La deriva no es el resultado inevitable de la Descarga cognitiva: la descarga describe el traslado de una operación mental; la deriva, la pérdida de seguimiento sobre ese traslado.

Por ejemplo, una estudiante puede pedir a una IA que resuma cada fuente de un tema y usar esas síntesis para redactar una conclusión. Si al revisar su texto no puede mostrar qué pasajes contrastó, qué resumen corrigió o qué criterio aplicó, el producto por sí solo no permite saber cómo razonó. Registrar una fuente revisada, una decisión propia o una pregunta que quedó abierta puede devolver visibilidad al proceso; el caso no demuestra una incapacidad permanente ni que toda intervención de IA haya sido perjudicial. Términos vecinos: Metacognición, Pereza metacognitiva.

Fuentes públicas y alcance: [Beware of metacognitive laziness: Effects of generative artificial intelligence on learning motivation, processes, and performance](#) (Estudio empírico sobre efectos de IA generativa en motivación, procesos y desempeño; usar el término sólo para el contexto y diseño reportados.)

Ubicación orientativa: §7.1

Descarga cognitiva. La **descarga cognitiva** (*cognitive offloading*) ocurre cuando una persona traslada parte del esfuerzo de procesar, almacenar o estructurar información a una herramienta externa, como una libreta, una calculadora o un sistema de IA. No es un error por definición: recordar una cita con un gestor bibliográfico o pedir a una IA que ordene una tabla puede dejar atención disponible para una tarea distinta.

La cuestión es qué operación se traslada, para qué y qué hace después la persona con el tiempo o la atención que liberó. Si delega una operación ya dominada para comparar evidencias, revisar un argumento o explicar una decisión, la descarga puede apoyar ese trabajo. Si delega

justamente la operación que la actividad buscaba que practicara, el producto puede mejorar sin que eso permita inferir aprendizaje. La atrofia de una capacidad o la Deriva metacognitiva no son resultados necesarios: son riesgos que dependen de la operación, la frecuencia, los apoyos y la posibilidad de revisar lo que se hizo.

La descarga tampoco equivale por sí misma a la deriva: la deriva aparece cuando la persona deja de advertir qué está delegando y qué ya no puede hacer o justificar sin la herramienta. Tampoco se opone a la Ganancia cognitiva: puede haber traslado de una operación y, a la vez, un resultado favorable que la persona explica, aplica y defiende después de la interacción. La expresión Deuda cognitiva procede de un estudio concreto y se define aparte.

Fuentes públicas y alcance: [Cognitive Offloading](#) (Revisión conceptual de la descarga cognitiva como delegación a recursos externos; el valor depende de la tarea y no implica beneficio o daño automático.)

Ubicación orientativa: §7.1

Descarga cognitiva estratégica. La **descarga cognitiva estratégica** es una modalidad deliberada de Descarga cognitiva: la persona delega una operación que ya domina para dedicar tiempo y atención a un trabajo de mayor exigencia. Por ejemplo, puede encargar a una IA que ordene datos que sabe organizar y concentrarse en interpretar una diferencia o comprobar una conclusión. El propósito y el uso de la atención liberada permiten distinguirla de una delegación que sustituye justamente lo que la actividad buscaba enseñar. Liberar tiempo no garantiza por sí solo un mayor aprendizaje: hay que observar qué hace la persona con ese margen.

Fuentes públicas y alcance: [Pedagogical partnerships with generative AI in higher education: how dual cognitive pathways paradoxically enable transformative learning](#) (Estudio mixto con estudiantes de educación superior en China, Europa y Estados Unidos (N=912); asociación y modelos de mediación, no receta causal general.); [Cognitive Offloading](#) (Revisión conceptual de la descarga cognitiva como delegación a recursos externos; el valor depende de la tarea y no implica beneficio o daño automático.)

Ubicación orientativa: §7.1

Desempeño y aprendizaje. El **desempeño** es lo que una persona logra durante una tarea: resolver más ejercicios, escribir un texto más claro, terminar más rápido. El **aprendizaje** es lo que después puede comprender, explicar y volver a hacer en otra situación. Una herramienta puede mejorar el desempeño inmediato sin producir el mismo efecto en el aprendizaje, y un mejor producto después de que una IA generativa redacta o revisa parte de la tarea no demuestra por sí solo un mayor aprendizaje.

En una actividad concreta, recibir una solución completa puede producir un resultado inmediato distinto de recibir una pista que permite continuar el razonamiento. Ese contraste ayuda a decidir qué función tendrá la IA y en qué momento intervendrá, pero no autoriza a generalizar los resultados de un experimento a cualquier asignatura, herramienta o estudiante. La pregunta útil ante un resultado con IA es qué puede hacer la persona cuando la herramienta ofrece sólo una pista o deja de estar disponible y qué evidencia sería suficiente para sostener esa interpretación, sin convertir ese umbral en una prueba concluyente. A escala de grupo o institución, un desfase semejante puede explorarse como Paradoja de productividad.

Fuentes públicas y alcance: [Generative AI without guardrails can harm learning: Evidence from high school mathematics](#) (Experimento en matemáticas de bachillerato sobre soluciones completas frente a pistas y evaluación sin IA; no generalizar a todos los cursos o herramientas.)

Ubicación orientativa: §2.2

Detector automático de texto generado por IA. Sistema que intenta estimar, a partir de patrones estadísticos, si un texto fue producido por un modelo de lenguaje. No observa quién escribió el texto ni reconstruye cómo se produjo. Puede señalar como generado por IA un texto humano: ese error se llama **falso positivo**. En una evaluación de 2023, varios detectores clasificaron erróneamente muestras de personas que escribían en inglés sin ser hablantes nativas; el mismo estudio mostró que determinadas reescrituras podían eludir la detección. Son resultados situados en las herramientas y muestras examinadas, no una tasa de error válida para cualquier detector, idioma o fecha. [Liang y colaboradores, 2023](#).

Por esas razones, una señal de detector no puede fundar por sí sola una decisión con consecuencias académicas. A lo sumo abre una pregunta que se revisa con evidencia independiente, con las versiones y notas que la persona pueda mostrar y con una conversación sobre el trabajo. La cuestión educativa sigue siendo qué comprendió, practicó y puede explicar la persona, no solo si intervino IA. Términos relacionados: Escritura híbrida, Marca de agua, Herramientas de evasión (humanizadores), Pos-plagio, Trazabilidad.

Fuentes públicas y alcance: [GPT detectors are biased against non-native English writers](#) (Evaluación de detectores y escritura en inglés de personas no nativas; respalda falsos positivos y sesgo en ese contexto, no una tasa universal de acierto.)

Ubicación orientativa: §5.6 y anexo B

Deuda cognitiva. Expresión usada en un experimento concreto para nombrar una posible consecuencia de delegar parte del trabajo de escritura a un asistente generativo. El estudio circula como preprint y tiene un alcance limitado; por eso la expresión funciona aquí como una hipóte-

sis de investigación, no como diagnóstico de quien usa IA ni como daño automático. Puede orientar preguntas sobre qué operación se delega y qué evidencia de comprensión conserva la persona, en relación con la Descarga cognitiva.

Fuentes públicas y alcance: [Your Brain on ChatGPT: Accumulation of Cognitive Debt when Using an AI Assistant for Essay Writing Task](#) (Preprint: 54 participantes en sesiones iniciales y 18 en la cuarta; EEG y tareas de escritura; resultado exploratorio, no diagnóstico ni ley general.)

Ubicación orientativa: §7.1

Dirección epistémica. La **dirección epistémica** es la capacidad de conservar el rumbo intelectual de un trabajo cuando interviene una IA: mantener el propósito, comprobar lo que la herramienta entrega, tomar las decisiones importantes y responder por el resultado. Puede observarse mediante cuatro preguntas cotidianas: ¿la persona sabe qué intenta comprender o producir?, ¿comprueba lo que recibe?, ¿toma las decisiones importantes?, ¿puede responder por el resultado?

El concepto no mide cuánta IA se usó, sino quién condujo el proceso. Una persona puede apoyarse mucho en una herramienta y conservar la dirección, si formuló lo que buscaba, contrastó las respuestas y puede explicar por qué aceptó, modificó o descartó cada aporte. También puede usarla poco y perderla, si entrega lo primero que recibió sin poder sostenerlo. Por eso este criterio distingue una Cocreación persona-IA que apoya el pensamiento de una delegación que puede sustituirlo, y ayuda a revisar si la interacción deja una Ganancia cognitiva o una Descarga cognitiva sin trabajo posterior. Atraviesa las tres Literacidades operativa, crítica y de co-creación. Es la forma observable que este glosario da a la Agencia humana; se relaciona con la Soberanía cognitiva cuando la persona sostiene un juicio propio, pero no presupone que ambos términos nombren la misma capacidad.

Ubicación orientativa: §3.1

Diseño inverso. El **diseño inverso** (*backward design*), formulado por Grant Wiggins y Jay McTi-ghe, es un método para planificar cursos, unidades o actividades comenzando por el aprendizaje que se quiere observar, no por la herramienta o la dinámica disponible. Ordena tres preguntas: qué necesita aprender la persona; qué acción o explicación permitiría reconocer ese aprendizaje; y solo entonces qué actividad y qué apoyos necesita el grupo. Si interviene IA generativa, esta tercera pregunta debe precisar su tarea, por ejemplo proponer una objeción o comentar un borrador sin producir la conclusión.

El método ayuda a revisar la alineación entre propósito, evaluación y experiencia, pero no garantiza por sí solo una tarea profunda. Si se incorpora IA, su función se decide después de definir qué evidencia corresponde a la persona y bajo qué condiciones podrá interpretarse; si la conversación comienza por permitir o prohibir la herramienta, el orden ya se invirtió. La Ficha de diseño cognitivo traduce ese orden en cinco preguntas para una actividad concreta. Para precisar la operación cognitiva de cada objetivo sirve la Taxonomía de Bloom.

Ubicación orientativa: §4.2

Doble imperativo. Ordenación de la formación que el informe del Digital Education Council propone como un doble desarrollo: habilidades aplicadas de IA, como diseñar flujos de trabajo asistidos, formular instrucciones y evaluar y adaptar lo que devuelve el sistema, junto con capacidades humanas duraderas, entre ellas la experiencia disciplinar, el pensamiento sistémico y el planteamiento de problemas. Es una propuesta del informe, no un mandato de la UdeG. El mismo informe describe una posible caída inicial de productividad mientras se aprende a revisar los resultados; aquí, en cambio, la Paradoja de productividad es una lente operativa de Orientaciones para preguntar si producir más se acompaña de aprender más.

Fuentes públicas y alcance: [AI Adoption May Reduce Productivity Before Improving It, New DEC Report Finds](#) (Informe del Digital Education Council (13 mayo 2026): propone el desarrollo conjunto de habilidades aplicadas de IA y capacidades humanas duraderas; su posible descenso inicial de productividad no equivale a la paradoja educativa operativa del glosario.)

Ubicación orientativa: §3.5

Educación dialógica. Tradición pedagógica anterior a la IA que entiende el pensamiento como una práctica que se desarrolla entre personas mediante preguntas, escucha, contraste y elaboración compartida, no en la recepción pasiva. Una actividad con IA puede incorporar preguntas que ayuden a revisar una posición, pero no equivale por ese hecho a un buen diálogo: este requiere interlocución humana, condiciones de escucha y una responsabilidad compartida por lo que se examina. El término orienta el Aprendizaje activo cuando el diseño crea esas posibilidades, no cuando solo añade una respuesta automática.

Ubicación orientativa: §4.3

Encuestas sintéticas. Uso de modelos de lenguaje como sustitutos de las personas encuestadas. Las evaluaciones metodológicas disponibles advierten que estos modelos tienden a homogeneizar las respuestas, distorsionar las submuestras y reducir la diversidad de perspectivas, y

señalan el problema inverso: personas que responden encuestas con chatbots contaminan datos que se creían humanos. Interesa a las ciencias sociales como objeto de crítica metodológica, no como método recomendado.

Fuentes públicas y alcance: [Synthetic Replacements for Human Survey Data? The Perils of Large Language Models](#) (Estudio con ChatGPT comparado con ANES 2016–2020 sobre 11 grupos sociopolíticos; menor variación y coeficientes divergentes en subgrupos; no generaliza a todo LLM o población.)

Ubicación orientativa: anexo A

Ensamblaje cognitivo. Expresión de N. Katherine Hayles para estudiar las relaciones entre personas y sistemas técnicos como una configuración conjunta. Junto con el agenciamiento de Deleuze y Guattari y las asociaciones entre participantes humanos y no humanos que describe Latour, forma el marco conceptual que permite observar la relación completa. Estas obras no demuestran que una actividad con IA produzca aprendizaje; ofrecen el marco para mirarla sin atribuir el resultado solo a la persona o solo a la herramienta.

Fuentes públicas y alcance: [Unthought: The Power of the Cognitive Nonconscious](#) (Monografía que desarrolla cognición no consciente y cognitive assemblages humano-técnicos; no convierte la síntesis del glosario en término universal.)

Ubicación orientativa: §3.1

Escritura híbrida. Texto en el que intervienen, en proporciones variables, la escritura de la persona y la de un sistema. Esa combinación no se deja separar de modo fiable por un Detector automático de texto generado por IA a partir del producto final. La pregunta útil no es solo si intervino IA, sino qué hizo, comprendió y declara la persona (Declaración de uso de IA, Recorrido).

Ubicación orientativa: §5.6

Esfuerzo productivo. El **esfuerzo productivo** (*productive struggle*) es la dificultad que forma una capacidad, a diferencia de la dificultad que solo frustra o hace perder tiempo. Redactar un primer borrador torpe, equivocarse en un cálculo y encontrar el error, o intentar una traducción antes de consultar el diccionario son esfuerzos de este tipo: la operación que cuesta trabajo es justamente la que se está aprendiendo.

El concepto orienta una decisión central al trabajar con IA: en qué momento y para qué tarea introducir una ayuda de IA generativa. Protegerlo suele significar dejar que la persona practique primero la operación que necesita aprender, pero no impone esa secuencia en cualquier caso.

Una pista temprana, por ejemplo una pregunta que señale qué criterio revisar, puede sostener la práctica sin resolverla; una solución completa que entrega el razonamiento o el resultado puede reemplazar justamente el ejercicio que se buscaba practicar. Si la IA libera trabajo ya dominado para que la persona pueda explicar, comparar o revisar, funciona como Descarga cognitiva estratégica; la pertinencia de una pista temprana depende de qué evidencia necesita producir la actividad. Fuera del aula, el estrechamiento del Primer escalón laboral da una razón adicional para no delegar durante la formación la práctica que forma el criterio.

Ubicación orientativa: §4.1

Evaluación auténtica. La **evaluación auténtica** reúne principios de diseño que buscan relacionar una tarea con prácticas, problemas o decisiones significativas para una disciplina o un contexto. No es un formato único ni una propiedad binaria: una evaluación puede incorporar distintos grados y formas de autenticidad según su propósito, disciplina y etapa de aprendizaje. Su valor depende del razonamiento que moviliza, no de que el producto parezca profesional.

Suele pedir que la persona aplique conocimiento, tome decisiones, produzca algo pertinente y justifique sus elecciones. La evidencia observable puede incluir el producto, el proceso seguido, la respuesta a preguntas, la revisión posterior o la manera de usar criterios para juzgar calidad. La autenticidad no garantiza por sí sola aprendizaje, integridad académica, inclusión o preparación profesional; su diseño exige explicitar qué apoyos necesita el grupo, qué parte debe realizarse sin asistencia, qué herramientas son pertinentes y cómo se distinguirá una ejecución aparente de una comprensión suficiente. Ante la IA generativa, la pregunta no es cómo volver «irreemplazable» a una persona, sino qué evidencia permitirá interpretar su aplicación, juicio y justificación; según el propósito, la IA puede excluirse, usarse como objeto de crítica o integrarse con reglas y responsabilidades explícitas.

Fuentes públicas y alcance: [Authentic assessment: from panacea to criticality](#) (Análisis crítico de evaluación auténtica en educación superior; apoya tratarla como principios situados, no como formato o garantía binaria.)

Ubicación orientativa: §5.3

Evaluación de dos vías. Arquitectura de evaluación a nivel de programa que combina **momentos asegurados**, donde se verifica en persona el logro individual sin IA, como defensas, coloquios, prácticas y avances presentados presencialmente, con **trabajo abierto** en el que la IA participa de manera declarada y el recorrido queda a la vista. La distribución se decide entre quienes conocen el programa, no desde una regla general, y responde a cuatro preguntas: qué

capacidad promete la titulación que ninguna herramienta debería sustituir, en qué momento tiene sentido verificarla, qué formato la muestra con menos artificio y qué queda en trabajo abierto y cómo se hará visible.

Tratar la integridad como todo o nada puede empujar hacia el examen presencial todo lo que importa. Convertirlo todo en examen presencial tampoco resuelve el problema: limita las formas de evidencia y deja fuera la formación para contextos donde la IA está disponible. Los momentos breves sin IA reciben el nombre de Puntos de seguridad.

Ubicación orientativa: §5.5, §5.7 y anexo B

Ficha de diseño cognitivo. Cinco preguntas que acompañan el diseño de una actividad con IA: qué debe aprender la persona, qué operación debe realizar por sí misma, qué puede delegar sin vaciar el aprendizaje, qué actividad de mayor nivel hace posible esa delegación y cómo mostrará que comprendió, verificó y decidió. Traduce el orden del Diseño inverso en una revisión explícita de la relación entre operaciones, apoyos y evidencias. La operación protegida se decide antes que la delegación, y la delegación se justifica por la actividad de mayor nivel que puede hacer posible.

Ubicación orientativa: §4.2

Ganancia cognitiva. La **ganancia cognitiva** es el desenlace favorable de la interacción con IA: una mejora que la persona puede mostrar en su razonamiento después de trabajar con la herramienta, porque la respuesta la obligó a comprender, relacionar o revisar algo que ahora puede explicar. Puede consistir en detectar un supuesto, corregir una explicación, conectar ideas o formular una pregunta más precisa. Obtener una respuesta más rápida o mejor redactada no basta; la rapidez del resultado no la demuestra.

La ganancia puede coexistir con la Descarga cognitiva: una persona puede trasladar una operación y, a la vez, conservar el trabajo de juzgar. Si la respuesta sirve para contrastar, verificar y revisar una explicación propia, puede haber ganancia; si se acepta sin examinarla, no hay evidencia suficiente para afirmarla. No es automática: depende de la tarea, de los conocimientos de la persona, de la orientación docente y de cómo se use la respuesta.

Por ejemplo, una estudiante parte de un borrador donde afirma que una política pública tuvo un único efecto. Pide a la IA un contraargumento, pero no lo copia: lo contrasta con el texto del curso, verifica una referencia y descubre que su afirmación era demasiado general. Después reescribe el párrafo y anota qué cambió y por qué. El texto final puede ser más claro, pero la mejora importante es que puede explicar el supuesto que corrigió y defender la nueva versión

sin depender de la conversación con la IA. «Ganancia cognitiva» es una síntesis editorial para nombrar un cambio que debe observarse y comprobarse en cada actividad, no un resultado atribuible a la herramienta por sí sola.

Ubicación orientativa: §3.1

Gemelo digital. Réplica computacional de un sistema físico que se actualiza con datos reales. En varios campos de la ingeniería se describe su evolución de réplica pasiva a sistema acoplado con IA generativa para diseño, simulación y mantenimiento predictivo. Aparece en el glosario porque el panorama por campos profesionales lo usa; no forma parte del vocabulario pedagógico del documento.

Ubicación orientativa: anexo A

Gobernanza distribuida. La **gobernanza distribuida** organiza las decisiones sobre IA entre los ámbitos que cuentan con información y responsabilidad para revisarlas, conservando claro quién responde por cada parte. Por ejemplo, una institución puede establecer condiciones generales sobre privacidad o contratación de servicios; un programa puede traducirlas a necesidades disciplinares; y un curso puede decidir si la IA generativa propondrá ejemplos, comentará un borrador o no intervendrá en una actividad concreta. La distribución exacta depende de las atribuciones y procesos de cada institución.

Distribuir no significa diluir responsabilidades ni descartar reglas comunes. Exige aclarar quién responde por cada decisión, qué criterios comparte el conjunto y cómo se revisan los acuerdos cuando cambia la tecnología, la evidencia o una necesidad docente. Puede combinar principios generales con criterios revisables, en vez de pretender que una lista fija de herramientas resuelva todos los casos. Dos instrumentos la acompañan: el Modelo de madurez institucional para decidir por dónde empezar y los Indicadores de seguimiento para conversar sobre tendencias.

Ubicación orientativa: §6.5

Herramientas de evasión (humanizadores). Servicios que reescriben texto para alterar rasgos que algunos detectores asocian con contenido generado. Su resultado y sus efectos dependen de la herramienta, el texto y el criterio de detección; no permiten inferir de manera fiable quién escribió ni qué comprendió una persona. Su existencia refuerza la conclusión de que un detector o una marca no deben convertirse en prueba individual y desplaza la integridad hacia la formación, la transparencia y el recorrido.

Ubicación orientativa: §5.6

Huella ambiental. Costos materiales del entrenamiento y el uso de la IA generativa: energía, agua y emisiones asociadas a los centros de datos, además de los minerales y las cadenas de suministro que sostienen la infraestructura. Frente a la imagen de una «nube» inmaterial, cada consulta moviliza una infraestructura física.

Las cifras de consumo por consulta que circulan deben leerse con cautela, porque dependen del modelo, de la longitud de la interacción, del equipo, de la ubicación y de la metodología, y con frecuencia mezclan entrenamiento, inferencia y consumo general de los centros de datos. Lo que no cambia es la materialidad de la infraestructura, cuyos datos deben fecharse y actualizarse. El criterio de Proporcionalidad lleva esta consideración a la elección de herramienta.

Fuentes públicas y alcance: [Data centres and data transmission networks](#) (Datos y actualización de IEA sobre consumo eléctrico de centros de datos a junio de 2026; no proporciona un consumo fijo por consulta.); [Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile](#) (Perfil NIST de riesgos y acciones para IA generativa; definición técnica y marco voluntario, no política educativa obligatoria.)

Ubicación orientativa: §7.4

IA generativa. La **IA generativa** reúne sistemas que producen texto, imágenes, código, audio u otros contenidos a partir de una instrucción, según patrones aprendidos en grandes volúmenes de datos. Muchos usan modelos fundacionales, como los modelos de lenguaje de gran escala, pero el término describe una capacidad de generación, no una garantía de comprensión, intención o verdad: no comprende ni verifica automáticamente aquello que produce.

En educación, una respuesta generativa puede servir como borrador, ejemplo, explicación alternativa o punto de contraste. No equivale a una búsqueda: puede expresar con seguridad información falsa, incompleta o sin fuente. Quien la usa necesita decidir para qué la pide, comprobar las afirmaciones pertinentes fuera de la herramienta y responder por lo que incorpora al trabajo. Ajustar una instrucción puede orientar el resultado, pero no elimina sesgos ni vuelve fiable una respuesta por sí solo.

Nota sobre la sigla IAG. Algunos textos emplean IAG para «inteligencia artificial generativa». Este glosario y las Orientaciones escriben «IA» y, cuando la precisión lo pide, «IA generativa»; la sigla IAG se conserva únicamente cuando forma parte de un nombre oficial citado.

Fuentes públicas y alcance: [Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile](#) (Perfil NIST de riesgos y acciones para IA generativa; definición técnica y marco voluntario, no política educativa obligatoria.)

ICAP. Marco propuesto por Michelene Chi y Ruth Wylie para examinar qué hace una persona con el conocimiento a partir de conductas visibles. Distingue cuatro modos de implicación: **pasivo**, cuando la persona recibe información, por ejemplo al leer una respuesta de IA sin trabajar con ella; **activo**, cuando manipula lo dado al seleccionar, subrayar u ordenar; **constructivo**, cuando añade algo que no estaba en el material, como una explicación propia, una relación, una pregunta o una corrección; e **interactivo**, cuando construye una explicación mediante turnos en los que los participantes aportan ideas nuevas y elaboran sobre las del otro.

La hipótesis central del marco propone que, en actividades comparables y bien diseñadas, la implicación interactiva suele favorecer más aprendizaje que la constructiva, esta más que la activa y la activa más que la pasiva. No es una garantía: un grupo donde una persona resuelve y las demás copian no trabaja de manera interactiva, y una pantalla con botones no produce por sí sola actividad cognitiva. Conversar con una IA puede apoyar un trabajo constructivo cuando la persona contrasta y reformula. El marco admite un agente computacional que responda pertinentemente como interlocutor; para considerar interactivo el trabajo exige aportaciones constructivas de ambas partes y un intercambio que las desarrolle. La presencia de una pantalla o de una IA no determina el modo de participación. Tampoco existe una correspondencia fija entre ICAP, la Taxonomía de Bloom y el Modelo SAMR.

Fuentes públicas y alcance: [The ICAP framework: Linking cognitive engagement to active learning outcomes](#) (Marco ICAP de implicación cognitiva y resultados de aprendizaje; usar como lente de diseño, no como garantía de resultado.); [Active-Constructive-Interactive: A Conceptual Framework for Differentiating Learning Activities](#) (Marco conceptual para diferenciar actividad pasiva, activa, constructiva e interactiva; la clasificación de conducta no prueba por sí sola aprendizaje.)

Ubicación orientativa: §4.3

Indicadores de seguimiento. Cifras registradas con definiciones estables y comparadas en el tiempo, como quién usa qué herramienta, qué programas cuentan con orientaciones, quién recibió formación y dónde persisten barreras, que permiten conversar sobre tendencias en lugar de impresiones. No deben convertirse en metas aisladas, porque un mayor uso no demuestra un mayor aprendizaje; el seguimiento cumple su función cuando lo observado conduce a continuar, modificar, ampliar, pausar o detener una iniciativa.

Ubicación orientativa: §9.3

Ingeniería de prompts. Cuando una docente o un estudiante necesita que una herramienta de IA realice una tarea delimitada, puede recurrir a la **ingeniería de prompts**, o diseño de instrucciones: escribir, probar y corregir lo que se pide. Para una tarea sencilla puede bastar una frase

directa: «Resume este párrafo en dos oraciones sin añadir datos». Una tarea más compleja puede necesitar el texto inicial, la audiencia, límites, ejemplos o un formato de respuesta. Lo que se escribe es un Prompt o instrucción.

Por ejemplo: «Compara esta introducción con los tres puntos de la rúbrica. Señala una omisión y explica dónde la detectaste. No reescribas el texto». La respuesta todavía debe revisarse contra la introducción y la rúbrica.

Una instrucción clara puede reducir respuestas irrelevantes, pero no garantiza exactitud ni elimina información inventada. Tampoco sustituye el acceso a fuentes, la recuperación de documentos o la comprobación final. Si el resultado falla porque falta información, hay que añadir el material necesario; si falla porque la afirmación es falsa, hay que contrastarla fuera de la herramienta.

Ubicación orientativa: Prompt o instrucción)

Integridad académica. La **integridad académica** es el compromiso con seis valores en la enseñanza, el aprendizaje y la producción de conocimiento: honestidad, confianza, justicia, respeto, responsabilidad y valentía. Incluye atribuir las aportaciones ajenas, aplicar criterios con equidad, reconocer límites y responder por el trabajo que se presenta.

Cuando interviene IA, esos valores toman una forma concreta: honestidad sobre la tarea que realizó la herramienta —por ejemplo, proponer ideas, redactar o revisar—, comprensión de lo que se presenta y capacidad de reconstruir las decisiones principales del trabajo. Las Orientaciones la fundan en la formación, la transparencia y el Recorrido, no en la vigilancia. Se concreta mediante reglas comprensibles, Declaración de uso de IA cuando corresponda, atribución de fuentes y oportunidades razonables para explicar o demostrar el aprendizaje. Ninguna herramienta o diseño garantiza integridad por sí solo: cada institución y actividad debe definir qué tareas permite encargar a una IA y cómo revisará una duda razonable sin sustituir el debido proceso ni convertir una señal automática en veredicto. La Validez del título es lo que está en juego a escala institucional; el Pos-plagio amplía las preguntas que la integridad debe hacerse.

Fuentes públicas y alcance: [The Fundamental Values of Academic Integrity](#) (Marco público de valores de integridad académica; no establece por sí solo reglas institucionales de uso de IA.)

Ubicación orientativa: §5.6 y §6.5

Investigación profunda (deep research). Modalidad de los sistemas generativos que busca, organiza y sintetiza fuentes en varias pasadas, con capacidades de agente. Puede apoyar una búsqueda bibliográfica si cada fuente y cada afirmación se comprueban fuera de la herramienta.

ta. Delegar partes de lectura o síntesis no vacía automáticamente la formación: su pertinencia depende del propósito, de qué operación se quiere aprender y de la posibilidad de comprobar y discutir el resultado.

Fuentes públicas y alcance: [Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile](#) (Perfil NIST de riesgos y acciones para IA generativa; definición técnica y marco voluntario, no política educativa obligatoria.)

Ubicación orientativa: §5.6 y anexo B

Justicia epistémica. En estas Orientaciones, la **justicia epistémica** se usa para examinar si una comunidad puede producir y validar conocimiento desde su propio contexto: con sus lenguas, sus problemas, sus fuentes y sus criterios de lo que cuenta como saber legítimo. El problema aparece cuando las herramientas o las políticas tratan una sola tradición de conocimiento como la medida de todas, o cuando un sistema jerarquiza lenguas, regiones y tradiciones intelectuales al decidir qué devuelve.

En las decisiones sobre IA, este concepto funciona como bisagra entre dos planos que suelen discutirse por separado. En el plano ético, invita a examinar qué perspectivas privilegia un modelo y a contrastar sus respuestas con la experiencia y las fuentes del contexto, más allá del Sesgo algorítmico puntual. En el plano tecnológico, da razones a la Soberanía tecnológica: conservar la capacidad de elegir y revisar herramientas también protege la posibilidad de conocer desde aquí, frente al riesgo de Colonización cognitiva. En una universidad pública mexicana, la pregunta es qué voces mexicanas, latinoamericanas, indígenas o no anglófonas quedan subordinadas en lo que un sistema devuelve, y qué haría falta para que aparezcan.

Ubicación orientativa: §6.2 y §8.5

Literacidades operativa, crítica y de cocreación. Las Orientaciones entienden por **literacidad** una capacidad para actuar con comprensión y criterio en situaciones concretas, no un listado de funciones de una herramienta. Distinguen tres. La **literacidad operativa** permite comprender qué puede hacer un sistema, reconocer sus límites y orientar su uso: saber pedir, acotar y volver a pedir. La **literacidad crítica** permite examinar las respuestas, los supuestos y las condiciones en que el sistema funciona, así como sus consecuencias: comprobar una afirmación, advertir un sesgo o preguntar qué datos se están entregando. La **literacidad de cocreación** permite elaborar y revisar ideas, explicaciones o soluciones con aportaciones de IA conservando el propósito, el juicio y la responsabilidad por el trabajo.

Las tres se desarrollan juntas, con práctica y acompañamiento, y las atraviesa la Dirección epistémica. Pueden organizarse como una progresión acumulativa, pero desde el primer ejercicio se aprende a delimitar una petición, comprobar una respuesta y explicar una decisión. Elegir que la IA proponga una objeción sin redactar la conclusión, o decidir que no intervenga, también puede expresar una formación lograda. Son la organización formativa de la Alfabetización en IA.

Ubicación orientativa: §3.2 a §3.5

Marca de agua. Una **marca de agua** es información codificada dentro del contenido para señalar o facilitar información sobre su procedencia. Se distingue de los metadatos que viajan por un canal aparte. Puede apoyar tareas de trazabilidad en los contextos y con las herramientas que la implementan, pero no reconstruye por sí sola quién escribió, revisó o comprendió un trabajo. Su presencia o ausencia tampoco constituye una prueba disciplinaria individual ni presupone obligaciones jurídicas universales.

Fuentes públicas y alcance: [Reducing Risks Posed by Synthetic Content: An Overview of Technical Approaches to Digital Content Transparency](#) (NIST AI 100-4 (2024), §3.1.1 p. 5/PDF p. 11: información codificada dentro del contenido y distinción frente a metadatos en un canal separado; no respalda por sí solo una obligación normativa.)

Ubicación orientativa: §5.6

Metacognición. La **metacognición** es la capacidad de observar el propio pensamiento. Al trabajar con IA toma la forma de una disciplina cognitiva: reconocer cuándo la herramienta apoya el razonamiento y cuándo lo está reemplazando, y saber cuándo no usarla. La investigación sobre autorregulación advierte que planear, monitorear y corregir el propio aprendizaje puede volverse una tarea compartida entre la persona y la herramienta, lo que exige enseñar a conservar esa regulación en lugar de cederla.

Una opción de diseño es intervenir en el instante en que se delega, con pausas que preguntan qué está haciendo la herramienta y qué necesita comprobarse, en vez de añadir solo una advertencia genérica al final. Su pérdida gradual recibe el nombre de Deriva metacognitiva; la Pereza metacognitiva es una expresión de origen experimental para una disminución de esa atención. Las cuatro acciones de formular, contrastar, decidir y responder concretan la disciplina cognitiva en una actividad.

Ubicación orientativa: §3.5 y §6.7

Modelo de lenguaje de gran escala (LLM). Tipo de IA generativa especializado en producir lenguaje, entrenado con grandes volúmenes de texto para predecir continuaciones plausibles. Es la base de muchos asistentes conversacionales y de los sistemas agénticos. Producir lenguaje fluido no equivale a comprenderlo ni a verificarlo.

Fuentes públicas y alcance: [Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile](#) (Perfil NIST de riesgos y acciones para IA generativa; definición técnica y marco voluntario, no política educativa obligatoria.)

Modelo de madurez institucional. Instrumento que ordena la revisión de las condiciones de una institución en niveles, de lo emergente a lo consolidado, en dimensiones como quién decide y responde, la formación existente, el apoyo al personal, el aporte educativo, la infraestructura y la protección y recuperación de los datos. Sirve como mapa para decidir por dónde empezar, no para clasificar ni juzgar; un área fuerte no compensa una carencia crítica en otra.

Ubicación orientativa: §6.8

Modelo de pesos abiertos (open-weight). Modelo cuyos parámetros pueden descargarse y, con los conocimientos, la licencia y los recursos adecuados, auditarse, adaptarse o alojarse localmente. Es una opción que una institución puede considerar cuando resulte técnica y pedagógicamente viable. Tener pesos abiertos no equivale a que todo el software sea libre ni reduce automáticamente dependencias, costos, riesgos de privacidad o necesidades de infraestructura; esos aspectos deben evaluarse por separado.

Fuentes públicas y alcance: [Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile](#) (Perfil NIST de riesgos y acciones para IA generativa; definición técnica y marco voluntario, no política educativa obligatoria.)

Ubicación orientativa: §8.3

Modelo SAMR. El **modelo SAMR** distingue cuatro maneras de describir la relación entre tecnología y tarea: **sustitución**, cuando cambia el medio sin cambiar la función principal; **aumento**, cuando se conserva la función y aparece una mejora útil; **modificación**, cuando la tecnología permite reorganizar de forma importante la tarea; y **redefinición**, cuando se vuelve viable una tarea pertinente que antes no podía realizarse de forma comparable.

Estas categorías no son una escala de calidad. Una sustitución puede resolver una necesidad real de acceso; una redefinición puede añadir novedad y carga sin mejorar el aprendizaje. SAMR describe qué cambió en la tarea, pero no demuestra qué comprendió o produjo la

persona. Tampoco existe una correspondencia fija entre SAMR, la Taxonomía de Bloom e ICAP: una actividad con poca transformación tecnológica puede exigir una comparación profunda, y una muy transformada puede limitarse a reconocer información.

Paradoja de productividad. Posible patrón a escala de grupo o de institución en el que se producen más textos en menos tiempo mientras la comprensión no crece al mismo ritmo. Es una forma colectiva de explorar el desfase entre Desempeño y aprendizaje, no un efecto que pueda atribuirse uniformemente a la herramienta. Conviene observar qué preguntas formula la persona, qué comprueba y qué decisiones conserva durante la actividad.

Ubicación orientativa: §7.1

Pensamiento crítico. El **pensamiento crítico** es la capacidad de examinar una afirmación o un problema, distinguir razones de conclusiones, revisar la evidencia disponible y formular un juicio que pueda explicarse y revisarse. No consiste en desconfiar de todo: pide hacer visibles los criterios con los que se acepta, corrige o rechaza una idea.

Al trabajar con IA, esa capacidad se vuelve visible en acciones concretas: pedir la fuente de una afirmación, contrastarla con un material confiable, detectar un supuesto ausente o explicar por qué una respuesta no sirve para el caso. La fluidez de una salida no demuestra que sea verdadera ni que responda a la pregunta. Por eso la persona debe revisar su solidez y decidir qué hacer con ella, en lugar de confundir rapidez de producción con rigor. Este cuidado también ayuda a reconocer cuándo una Descarga cognitiva libera trabajo para juzgar y cuándo deja de hacerlo.

Pereza metacognitiva. Expresión con la que los autores de un experimento sobre escritura académica nombraron una disminución de la atención sobre cómo se está pensando y aprendiendo en ese contexto. Su alcance está delimitado por ese estudio: no diagnostica a quienes usan IA ni demuestra un daño automático. Se relaciona con la Deriva metacognitiva, que nombra un riesgo de pérdida de seguimiento, y con la distinción entre Desempeño y aprendizaje.

Fuentes públicas y alcance: [Beware of metacognitive laziness: Effects of generative artificial intelligence on learning motivation, processes, and performance](#) (Estudio empírico sobre efectos de IA generativa en motivación, procesos y desempeño; usar el término sólo para el contexto y diseño reportados.)

Ubicación orientativa: §2.2

Portafolio iterativo. Conjunto de versiones sucesivas de un trabajo que permite observar qué cambió. Para comprender por qué cambió, puede acompañarse de explicaciones, notas o decisiones que sitúen esas modificaciones. Es una manera de conocer el Recorrido; conviene que la persona pueda guardar, recuperar y exportar sus versiones, de modo que la evidencia le pertenezca y no desaparezca con una sesión.

Ubicación orientativa: §5.4

Pos-plagio. **Pos-plagio** es una expresión usada en algunos debates para reconocer que las definiciones tradicionales de plagio, copiar palabras ajenas sin atribución, no describen por sí solas todos los problemas de un trabajo en el que interviene una IA. El término no sustituye las normas de cada institución ni resuelve por sí mismo cuestiones de autoría, fuentes o derechos; desplaza la atención hacia lo que la persona hizo, comprendió y declara sobre la tarea que realizó la herramienta.

No anuncia el fin de la integridad, sino un cambio de foco: además de preguntar si hubo copia, propone preguntar si la persona explicó si la IA propuso, redactó o revisó algo, si comprende lo que presenta y si puede responder por el resultado. Esa mirada conecta con la Trazabilidad del proceso y desaconseja basar una decisión académica solo en un Detector automático de texto generado por IA.

Ubicación orientativa: sin remisión (cuerpo: §5.6)

Primer escalón laboral (peldaño de entrada). Puestos y tareas de entrada mediante los cuales algunas personas novatas aprendían un oficio. La IA puede modificar o automatizar parte de esas tareas de manera desigual según el campo, la organización y las condiciones de trabajo; no describe una automatización uniforme. Ese riesgo invita a decidir qué prácticas siguen siendo necesarias durante la formación para desarrollar criterio y en qué momento de la trayectoria se enseñan.

Fuentes públicas y alcance: [Generative AI and jobs: A 2025 update](#) (Índice global de exposición ocupacional basado en tareas, expertos y predicciones; estima potencial de transformación, no automatización efectiva de puestos novatos.)

Ubicación orientativa: §2.3 y anexo A

Prompt o instrucción. Petición que una persona proporciona a una IA. Puede incluir el propósito, el contexto, los límites y la forma esperada de la respuesta. Por ejemplo: «Formula una objeción a este argumento; no lo reescribas». Escribir, probar y corregir esas peticiones es la Ingeniería de prompts; en español puede decirse también «instrucción» o «petición».

Ubicación orientativa: sin remisión (cuerpo: §4.4)

Proporcionalidad. Criterio que ajusta la herramienta al tamaño real de la tarea: cuando un sistema más ligero ofrece un resultado suficiente, utilizarlo puede consumir mucha menos energía que un modelo grande de propósito general. Se incorpora a la selección tecnológica comparando el valor educativo con el costo material, junto con los Criterios para elegir tecnología. También aconseja medir aquello que la universidad realmente pueda observar, sin presentar estimaciones generales como si fueran cifras propias.

Fuentes públicas y alcance: [Guidance for generative AI in education and research](#) (Orientación internacional: agencia humana, interacción pedagógica, proporcionalidad, participación y rendición de cuentas; no es una norma universal ni evidencia causal.); [Data centres and data transmission networks](#) (Datos y actualización de IEA sobre consumo eléctrico de centros de datos a junio de 2026; no proporciona un consumo fijo por consulta.)

Ubicación orientativa: §7.4 y §8

Puntos de seguridad. Momentos breves sin IA, ubicados en puntos clave de la trayectoria, reservados para observar una capacidad que la titulación promete. Son una opción de diseño dentro de una Evaluación de dos vías cuando el programa necesita evidencia individual de una capacidad; su ubicación y formato dependen de cuándo se ha formado esa capacidad y de qué correcciones aún son posibles.

Ubicación orientativa: §5.4

Recorrido. El **recorrido** es el conjunto de decisiones, versiones, comprobaciones y explicaciones que produjeron un trabajo. Comprenderlo, en lugar de mirar solo el producto final, es el eje de la evaluación que proponen las Orientaciones, porque un texto terminado puede parecer correcto sin mostrar quién comprendió, comparó o decidió.

Cada actividad puede elegir una muestra distinta sin registrar todo lo que hizo la persona. Cuatro maneras de conocer el recorrido tienen nombre en la literatura: las versiones sucesivas o Portafolio iterativo; la nota sobre una decisión o Bitácora de decisiones (decision log); la conversación sobre el trabajo o Defensa oral; y los criterios explícitos o Rúbrica. Cuando la actividad lo amerita, la evidencia de juicio se pide como Registro reflexivo selectivo. El propósito no es registrar cada paso, sino contar con información suficiente para conversar sobre el aprendizaje y evitar que una señal automática se convierta en veredicto. Un recorrido observable hace posible la Trazabilidad.

Ubicación orientativa: §5

Registro reflexivo selectivo. Evidencia breve del recorrido de una decisión en la que una IA generativa intervino para proponer, redactar o revisar algo. Reúne el propósito de esa intervención, las sugerencias que influyeron, las afirmaciones verificadas, los errores detectados, los cambios introducidos y la justificación final. Se opone al volcado indiscriminado de conversaciones: guardar decenas de páginas de intercambio no demuestra pensamiento crítico. La capacidad de construir un hilo completo sobre el propio proceso de decisión, y de sostenerlo ante personas que no comparten el mismo lenguaje técnico, distingue saber usar una herramienta de saber decidir con ella. Puede pedirse como respuestas breves a unas pocas preguntas: qué problema intentaba resolver, para qué recurrió a la IA, qué verificó, qué rechazó, qué decidió y qué incertidumbres permanecen.

Ubicación orientativa: §5.4

Resiliencia. Capacidad de mantener las funciones importantes y recuperar el trabajo cuando un proveedor falla, aumenta sus precios o cambia sus condiciones. Desaconseja depender de un único servicio para funciones críticas y pide describir el camino de salida antes de entrar. Es el corolario operativo de la Soberanía tecnológica.

Ubicación orientativa: §8.4

Rúbrica. Criterios explícitos acompañados de descriptores de calidad o niveles de logro para valorar un trabajo. Sirve para hacer visible qué se observará y con qué rasgos se distinguirán distintos desempeños; no es solo una lista de requisitos. Puede ayudar a conocer el Recorrido y a valorar cómo se dirigió y comprobó el uso de IA, sin imponer una política universal de calificación.

Ubicación orientativa: §5.4

Sesgo. Inclinação sistemática que favorece, omite o perjudica ciertas interpretaciones, grupos o casos. Puede aparecer en juicios humanos, prácticas institucionales o sistemas técnicos. Esta amplitud obliga a situar el término: no todo sesgo tiene el mismo origen, mecanismo ni forma de corrección. Cuando la inclinación está incorporada en datos, decisiones de diseño o usos de un sistema, se trata de Sesgo algorítmico.

Fuentes públicas y alcance: [Language \(Technology\) is Power: A Critical Survey of 'Bias' in NLP](#) (Revisión crítica de sesgo en tecnologías lingüísticas; respalda el carácter situado y sociotécnico del término, no una definición universal única.)

Ubicación orientativa: §6.2

Sesgo algorítmico. El **sesgo algorítmico** es la tendencia sistemática de un sistema a favorecer, omitir o perjudicar de manera desigual a determinados grupos, perspectivas o casos. Puede provenir de los datos de entrenamiento, de las categorías elegidas al diseñarlo, de los criterios con los que se evalúa o del contexto en que se usa. No es solo un error de datos ni algo que se resuelva automáticamente al añadir más información.

En educación, por ejemplo, una IA puede describir como normal una experiencia escolar que deja fuera a estudiantes con otra lengua, trayectoria o acceso tecnológico. El problema no se corrige pidiéndole al mismo sistema que se declare imparcial: hay que contrastar su respuesta con fuentes y experiencias pertinentes al contexto. El riesgo aumenta cuando se acepta una respuesta como neutral solo porque parece coherente. Examinar el sesgo implica preguntar qué casos aparecen, cuáles faltan, quién definió las categorías y qué consecuencias tendría usar ese resultado; esa práctica enlaza la Alfabetización en IA con el Pensamiento crítico y la Justicia epistémica, sin prometer que una revisión aislada elimine toda desigualdad.

Fuentes públicas y alcance: [AI Risk Management Framework](#) (Marco voluntario de gestión de riesgos de IA; respalda riesgos situados, sesgo, privacidad, homogeneización y asignación de responsabilidades.)

Ubicación orientativa: §6.2

Sistema agéntico. Configuración que organiza uno o varios Agentes de IA, herramientas, fuentes de información, memoria y mecanismos de coordinación para completar una tarea. Puede repartir subtareas, verificar condiciones o encadenar etapas, según los límites que establezcan sus responsables. No es un simple alias de agente: el agente nombra una entidad o capacidad de acción; el sistema agéntico nombra la arquitectura que articula varias de ellas, con sus permisos, reglas y registros.

Soberanía cognitiva. Capacidad de sostener un juicio propio, advertir qué operaciones realiza o sustituye una herramienta y elegir conscientemente la relación que se establece con ella. Las Orientaciones la entienden de manera operativa mediante intencionalidad, comprensión, control crítico y responsabilidad, en diálogo con las preguntas de la Dirección epistémica. No es auto-suficiencia: nadie piensa sin lenguajes, fuentes, docentes, comunidades y dispositivos; se trata de examinar las mediaciones que intervienen en el propio pensamiento. Se refiere a la persona; la Soberanía tecnológica y la Soberanía de datos se refieren a la institución.

Ubicación orientativa: §6.7

Soberanía de datos. Control sobre la información de la comunidad universitaria: qué se procesa fuera, dónde, bajo qué condiciones y garantías, y cómo se recupera o elimina. Los datos de estudiantes, docentes, investigación y gestión requieren control explícito antes de compartirse con terceros; si el proveedor no puede responder con claridad dónde se procesan y quién accede, la universidad tampoco podrá responder por ellos ante su comunidad. La literatura crítica llama colonialismo de datos a la normalización de su apropiación por servicios externos.

Fuentes públicas y alcance: [Guidance for generative AI in education and research](#) (Orientación internacional: agencia humana, interacción pedagógica, proporcionalidad, participación y rendición de cuentas; no es una norma universal ni evidencia causal.)

Ubicación orientativa: §8.4

Soberanía tecnológica. La **soberanía tecnológica** es la capacidad de una institución para conservar opciones sobre las tecnologías que adopta: poder recuperar sus datos, cambiar de servicio, revisar condiciones y evitar dependencias que comprometan su misión. No exige que una universidad construya todas sus herramientas; exige que ninguna compra o integración le quite decisiones indispensables.

Una parte de esa capacidad es la Soberanía de datos. Otra es la Resiliencia: mantener las funciones importantes cuando un proveedor falla o cambia sus condiciones. Los modelos de pesos abiertos y el software abierto son opciones distintas que también necesitan evaluarse. El riesgo que todo esto busca evitar tiene nombre propio: la Dependencia de proveedor (vendor lock-in). Conservar opciones incluye la escala de lo que se adopta: antes de contratar el servicio más grande disponible cabe preguntar si el beneficio justifica sus recursos y su Huella ambiental.

Fuentes públicas y alcance: [Guidance for generative AI in education and research](#) (Orientación internacional: agencia humana, interacción pedagógica, proporcionalidad, participación y rendición de cuentas; no es una norma universal ni evidencia causal.)

Ubicación orientativa: §8

Subyugación algorítmica. Forma extrema de pérdida de la Dirección epistémica en la que la persona deja de dirigir el proceso y se limita a recibir y validar lo que el sistema produce. Entre la cocreación dirigida y esta cesión existe un abanico de relaciones intermedias que la investigación ha comenzado a tipificar; la otra forma extrema es la Alienación epistémica.

Ubicación orientativa: §3.1

Taxonomía de Bloom. Una clase puede pedir que recuerdes una definición, que la expliques con un ejemplo o que la uses para resolver un problema. Aunque trabaje el mismo tema, cada consigna exige algo distinto. La **taxonomía de Bloom** ayuda a nombrar esas diferencias y a formular objetivos de aprendizaje.

La revisión publicada en 2001 por Anderson y Krathwohl distingue seis categorías de procesos cognitivos, recordar, comprender, aplicar, analizar, evaluar y crear, y considera también el tipo de conocimiento sobre el que se trabaja: factual, conceptual, procedimental y metacognitivo. La representación habitual en forma de pirámide es una simplificación didáctica; no obliga a leer las categorías como una secuencia ni atribuir ese dibujo a Bloom. Las categorías pueden solaparse y la revisión relaja una jerarquía estricta: explicar puede resultar más complejo que ejecutar según la tarea. Precisar el proceso y el contenido ayuda a evitar objetivos demasiado generales, como «comprender el tema». Un objetivo completo dice quién actúa, qué hace, sobre qué contenido, en qué condiciones y qué permitirá revisar su trabajo.

La taxonomía no obliga a seguir una secuencia fija ni convierte «crear» en el mejor objetivo para cualquier situación, y un verbo aislado no basta para clasificar la exigencia de una tarea: «analizar dos fuentes» puede significar marcar diferencias ya señaladas o exigir definir criterios y justificar una conclusión. Cuando interviene la IA, el verbo de la consigna o el nombre del producto no sustituyen la comprobación de qué hizo la persona para comprender, comprobar y decidir. Se relaciona sin correspondencia fija con ICAP, el Modelo SAMR y las Literacidades operativa, crítica y de cocreación.

Fuentes públicas y alcance: [A Revision of Bloom's Taxonomy: An Overview](#) (Krathwohl (2002) explica la revisión publicada en 2001, sus seis categorías de procesos y la dimensión de conocimiento; pp. 212–218, especialmente p. 215 sobre solapamiento y jerarquía no estricta.); [A Taxonomy for Learning, Teaching, and Assessing](#) (Catálogo de la Library of Congress para el libro de Anderson y Krathwohl (2001), fuente primaria bibliográfica de la revisión.)

Ubicación orientativa: §4.2

Trazabilidad. La **trazabilidad** es la posibilidad de reconstruir y mostrar cómo se produjo un trabajo: qué versiones existieron, qué decisiones se tomaron y qué papel tuvo la IA en cada momento. Cuando el proceso es trazable, docente y estudiante pueden conversar sobre él; cuando es opaco, solo queda juzgar el producto terminado.

Trazable no significa vigilado. No hace falta guardar cada conversación con una herramienta ni registrar cada tecla: basta una nota breve donde la persona explica qué tarea pidió a la IA —por ejemplo, proponer opciones, esbozar un texto o comentar un borrador—, qué comprobó fuera

de la herramienta y qué conservó o descartó. Quien documenta así su Recorrido puede sostener una Integridad académica centrada en comprenderlo, más justa y más útil que la detección automática de texto generado.

Fuentes públicas y alcance: [C2PA Technical Specification](#) (Especificación técnica de procedencia, manifiestos y marcas visibles/invisibles; señala procedencia verificable, no autoría ni obligación regulatoria.)

Ubicación orientativa: §5.6

Tutor inteligente. Un **tutor inteligente** (*Intelligent Tutoring System*, ITS) es un sistema diseñado para ofrecer orientación, práctica o retroalimentación adaptada a una tarea de aprendizaje. Puede usar reglas, modelos de datos o IA generativa para ajustar sus preguntas y respuestas. No equivale a un docente humano ni garantiza que comprenda la situación de quien aprende.

Para orientar bien, el sistema necesita alguna representación del contenido, información sobre las respuestas o decisiones de la persona y criterios para elegir la siguiente ayuda. Esas representaciones pueden ser incompletas o erróneas; también plantean preguntas sobre privacidad y sobre quién puede revisar la recomendación que recibe cada estudiante. Por ejemplo, un tutor puede pedir que una estudiante explique un paso de su solución y ofrecer una pista si detecta una dificultad: una pista señala un posible siguiente paso, no entrega la solución completa. Los experimentos que comparan asistentes que dan soluciones con asistentes que dosifican pistas muestran por qué esa diferencia importa para el Desempeño y aprendizaje. Antes de adoptarlo, conviene revisar qué datos recoge, qué tipo de retroalimentación ofrece, cuándo interviene una persona docente y qué alternativa existe si la recomendación falla.

Validez del título. Correspondencia formativa entre la credencial que otorga la universidad y las capacidades que afirma. Una duda sobre esa correspondencia requiere revisar qué oportunidades de aprendizaje, evidencias y condiciones de evaluación existieron; no produce por sí misma una invalidación jurídica automática. Una respuesta institucional combina formación y evaluación del Recorrido, en lugar de perseguir señales automáticas.

Ubicación orientativa: §7.3

Virtual (lo virtual). En una sesión de 1984, Deleuze distingue lo virtual de lo posible y lo actual de lo realizado: lo virtual es real aunque todavía no esté actualizado y se concreta en procesos que producen algo singular. En este glosario, esa distinción se usa como una interpretación situada de las Orientaciones para preguntar cómo las decisiones humanas (la instrucción, la selección y la comprobación) dan forma a un resultado de IA; no describe técnicamente a los modelos ni supone que contengan respuestas preexistentes.

Fuentes públicas y alcance: [Cinema 3: sesión del 22 de mayo de 1984](#) (Deleuze distingue virtual/actual de posible/real y reconoce la realidad de lo virtual (p. 15); el uso aplicado a IA es una interpretación de Orientaciones, no una afirmación del autor sobre modelos de lenguaje.)

Ubicación orientativa: §3.1

Taxonomía revisada de Bloom. Véase Taxonomía de Bloom.

APARATO

Referencias del anexo C

Las entradas del glosario explican el alcance de cada fuente en su contexto. Esta relación reúne los enlaces públicos citados, sin duplicar los que aparecen en varias entradas.

- [Guidance for generative AI in education and research](#)
- [Mil mesetas: capitalismo y esquizofrenia](#)
- [Active-Constructive-Interactive: A Conceptual Framework for Differentiating Learning Activities](#)
- [The ICAP framework: Linking cognitive engagement to active learning outcomes](#)
- [Blended learning: Uncovering its transformative potential in higher education](#)
- [Pedagogical partnerships with generative AI in higher education: how dual cognitive pathways paradoxically enable transformative learning](#)
- [Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile](#)
- [Measuring the Impact of Early-2025 AI on Experienced Open-Source Developer Productivity](#)
- [Generative AI and jobs: A 2025 update](#)
- [Beware of metacognitive laziness: Effects of generative artificial intelligence on learning motivation, processes, and performance](#)
- [Cognitive Offloading](#)
- [Generative AI without guardrails can harm learning: Evidence from high school mathematics](#)
- [Liang y colaboradores, 2023](#)
- [GPT detectors are biased against non-native English writers](#)
- [Your Brain on ChatGPT: Accumulation of Cognitive Debt when Using an AI Assistant for Essay Writing Task](#)
- [AI Adoption May Reduce Productivity Before Improving It, New DEC Report Finds](#)
- [Synthetic Replacements for Human Survey Data? The Perils of Large Language Models](#)
- [Unthought: The Power of the Cognitive Nonconscious](#)
- [Authentic assessment: from panacea to criticality](#)
- [Data centres and data transmission networks](#)
- [The Fundamental Values of Academic Integrity](#)

- Reducing Risks Posed by Synthetic Content: An Overview of Technical Approaches to Digital Content Transparency
- Language (Technology) is Power: A Critical Survey of 'Bias' in NLP
- AI Risk Management Framework
- A Revision of Bloom's Taxonomy: An Overview
- A Taxonomy for Learning, Teaching, and Assessing
- C2PA Technical Specification
- Cinema 3: sesión del 22 de mayo de 1984

Licencia y reutilización

El contenido original de esta edición se ofrece bajo [Creative Commons Atribución/Reconocimiento-CompartirIgual 4.0 Internacional](#), salvo las excepciones indicadas abajo.



Puedes compartir y adaptar

Copiar, redistribuir y transformar el material en cualquier formato, incluso con fines comerciales.



Reconoce la procedencia

Identifica la obra original y enlázala cuando sea posible. Conserva los créditos proporcionados, enlaza la licencia e indica los cambios. La atribución no debe sugerir que se respalda tu uso.



Comparte las adaptaciones con la misma licencia

Al distribuir una adaptación, usa CC BY-SA 4.0 o una licencia compatible. No impongas restricciones legales ni medidas tecnológicas que impidan lo permitido.

Excepciones y otros derechos

Los materiales de terceros conservan sus licencias, condiciones de uso y derechos reservados. Los logotipos, escudos y marcas quedan fuera de esta licencia. Pueden requerirse permisos por derechos de imagen, privacidad o derechos morales; la licencia no los sustituye.

El aviso no implica aprobación institucional, patrocinio, aval ni estatus oficial de esta edición o sus adaptaciones. Tampoco concede garantías. Consulta el [texto legal](#); el resumen no lo reemplaza.

Al reutilizar: identifica el título del documento o anexo, la edición de 2026 y los créditos que incluya; añade «CC BY-SA 4.0», el enlace a la licencia y los cambios realizados.

Edición 0.19-2026-09-08

creativecommons.org/licenses/by-sa/4.0/