

1.	Introducción
2.	Objetivo y ámbito de aplicación de la guía
3.	Definiciones
4.	Composición del equipo de auditoría: conocimientos y competencias requeridos
5.	Qué es la IA Generativa
6.	Metodología de la auditoría
7.	Obtención de conocimiento del proceso de gestión donde se integra la IA, incluido el control interno
8.	Identificación y valoración de riesgos derivados de la utilización de la IA Generativa
9.	Identificación y revisión de los controles de procesamiento de la información. Valoración del riesgo de control
10.	Identificación y revisión de los controles generales de tecnologías de la información
11.	Procedimientos posteriores de auditoría
12.	Evidencia de auditoría obtenida en entornos de IA Generativa
13.	Cuestiones relacionadas con el cumplimiento
14.	Bibliografía
Anexo 1 Ejemplo de documentación a solicitar durante la planificación	
Anexo 2 Conceptos básicos de Inteligencia Artificial relevantes para la auditoría	
Anexo 3 Marco normativo aplicable	

1. Introducción

Aunque la inteligencia artificial (IA) es utilizada hace muchos años en una amplia variedad de aplicaciones, a partir de 2022 se ha producido una auténtica explosión en la utilización y en las expectativas sobre las potencialidades de la Inteligencia Artificial Generativa (IA Generativa).

La Conferencia de Presidentes de la Asociación de Órganos de Control Externo Autonómicos (ASOCEX) aprobó el 1 de diciembre de 2025 la *GPF-OCEX 5375 Guía de auditoría de determinados procesos automatizados* en la que se destaca cómo, en los últimos años, ha crecido el uso de tecnologías emergentes en el sector público, incluyendo la IA Generativa que, de acuerdo con la IAASB¹, tienen una o más de las siguientes características:

- **Opacidad:** la lógica de la herramienta o el proceso para tomar decisiones no es transparente, a menudo denominado “comportamiento de caja negra”. En estos casos no es posible rastrear paso a paso cómo la herramienta llega a un determinado resultado. Al no poder explicar la lógica interna, el auditor no puede verificar, por ejemplo, si el proceso se basa en los criterios legales establecidos.
- **No determinismo:** los mismos inputs pueden producir resultados diferentes debido a procesos probabilísticos, sensibilidad al contexto u otras influencias impredecibles. Como el sistema puede dar respuestas distintas ante solicitudes idénticas debido a su naturaleza basada en probabilidades, puede suponer un riesgo crítico, por ejemplo, rompiendo el principio de igualdad si dos beneficiarios de subvenciones con los mismos méritos pudieran recibir puntuaciones diferentes sin una razón objetiva.
- **Adaptabilidad:** la herramienta evoluciona con la interacción con el usuario, actualizaciones o reaprendizaje. La herramienta va 'aprendiendo' y cambiando su forma de decidir a medida que se usa.

¹ IAASB Technology Quality Management Roundtables-Briefing Note for Participants, 20/8/2025.

Es preciso supervisar estos cambios para evitar, por ejemplo, que el criterio de evaluación original pueda desviarse con el tiempo (degradarse).

Se puede considerar la IA como un campo de la informática orientado al desarrollo de sistemas que puedan realizar tareas que normalmente requieren capacidades asociadas a la inteligencia humana, como el aprendizaje, el razonamiento y la percepción. Estos sistemas pueden percibir o interpretar información de su entorno, razonar sobre el conocimiento disponible, procesar la información derivada de los datos y generar predicciones, recomendaciones, contenidos o decisiones para lograr un objetivo dado.

Estas tecnologías emergentes, en particular la IA Generativa, tienen unas características disruptivas que desafían algunos de los principios básicos de la auditoría tradicional, y por tanto deben ser analizadas para conocer su impacto en los métodos de trabajo de los auditores públicos.

La IAASB, en su reunión de septiembre de 2025 señalaba cómo el uso de la IA por las entidades auditadas afecta a los procedimientos de auditoría²:

“7. Uso de la IA por las entidades auditadas y evolución de los procedimientos de auditoría:

Se ha observado un crecimiento en la adopción de IA en las entidades auditadas, lo que requiere que los auditores evalúen **la integridad de los datos, los controles operativos, la madurez de la gobernanza y las prácticas de gestión de riesgos**.

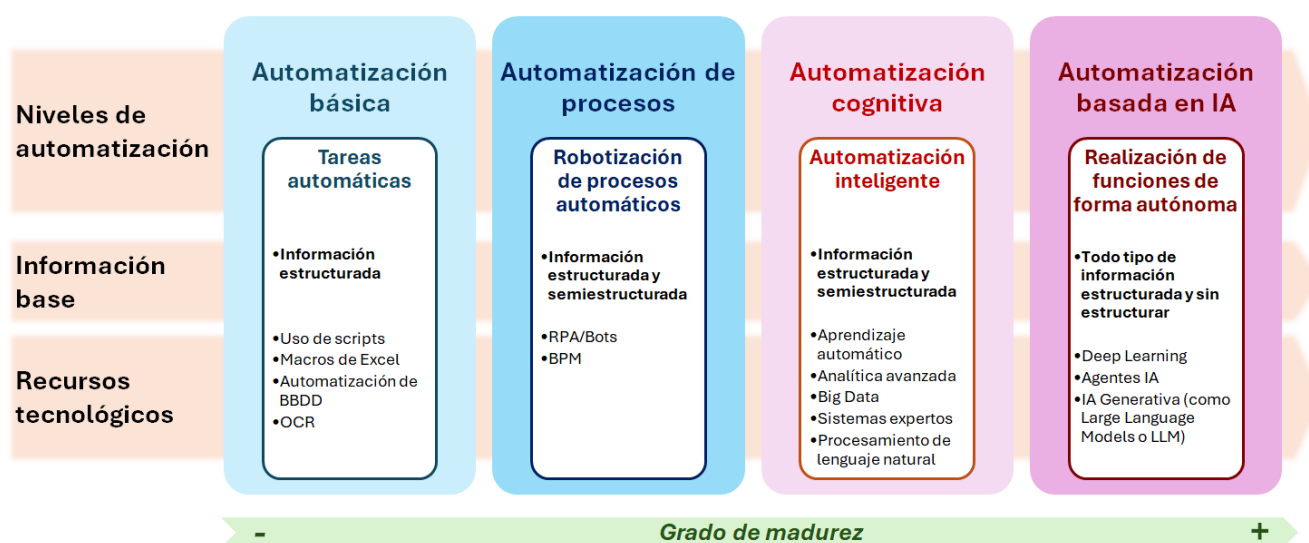
En muchos casos, esto implica **evaluar los procesos impulsados por IA** integrados en las operaciones principales de una entidad, considerar la suficiencia de los **marcos de gobernanza**, y **desarrollar procedimientos de auditoría adaptados** a esta nueva situación. Estos incluirán **evaluar el diseño y la eficacia operativa de los controles** relevantes sobre la IA, **evaluar la idoneidad de las entradas de datos**, **evaluar la integridad y transparencia de los resultados** de los modelos de IA, y **comprender cómo estos procesos se integran** con los flujos de trabajo generales de la entidad auditada.

En algunos casos, los auditores contarán con **especialistas TI** para evaluar algoritmos, los datos de entrenamiento de los modelos, las medidas de mitigación de sesgos, y desarrollarán estrategias para documentar los hallazgos de forma coherente con los requerimientos sobre la evidencia de auditoría.”

En la antes citada guía GPF-OCEx 5375 se distinguen cuatro etapas secuenciales en la automatización de procesos que representan una evolución de madurez tecnológica: automatización básica, automatización de procesos, automatización cognitiva y, finalmente, **automatización basada en IA**.

Intentando simplificar un panorama complejo, esa evolución se puede representar en el siguiente gráfico 1.

Gráfico 1 Tipos y evolución de las automatizaciones de procesos



² IAASB Main Agenda (September 2025), Agenda Item 7-A, Component 3 Monitoring Update, IAASB.

Las dos primeras fases incluyen herramientas claramente deterministas y las dos siguientes fases del proceso introducen modelos que pueden ser no deterministas, especialmente en la última fase.

El tercer grupo, **automatización cognitiva**, introduce capacidades de análisis y toma de decisiones mediante el uso de aprendizaje automático, Big Data y procesamiento del lenguaje natural. Las ventajas incluyen una mejora en la toma de decisiones, detección proactiva de irregularidades, optimización de recursos y enriquecimiento de los análisis.

La etapa final corresponde a la **automatización basada en IA**, caracterizada por el uso de *deep learning*, modelos de lenguaje de gran tamaño LLM, incluye la **IA Generativa** y los **Agentes IA**. Esta fase no está limitada por el tipo de información.

En la práctica, sin embargo, dado que actualmente existen en el mercado cientos de herramientas, que crecen y evolucionan cada día, podemos encontrarnos herramientas “híbridas” que será difícil encajar en un esquema tan simple como el anterior, el cual tiene una función meramente didáctica. En esos casos **habrá que emplear el juicio y escepticismo profesional para determinar frente a qué tipo de herramienta tecnológica o proceso de automatización nos encontramos y qué consecuencias tiene sobre los riesgos, sobre las características de la información generada por la IA y sobre los procedimientos de auditoría que se deben aplicar para obtener evidencia de auditoría.**

Se pueden auditar procesos relativamente sencillos, como el de concesión de subvenciones, donde puede ser sencillo “clasificar” el tipo de automatización utilizada.

En otros casos se deberá auditar una entidad que opera un ERP³ financiero complejo, con múltiples módulos interconectados, formado por cientos o miles de automatizaciones que individualmente consideradas se pueden calificar como “básicas” pero que en conjunto forman un sistema de elevada complejidad. Es probable que, cada vez más frecuentemente, esos ERP contengan algún subproceso concreto que incorpore aprendizaje automático para predecir tendencias futuras, IA Generativa y Agentes IA.

Los criterios establecidos en esta guía son coherentes con el marco técnico establecido por las NIA-ES-SP, Normas de Fiscalización del Tribunal de Cuentas y las GPF-OCEX ya las que complementa con orientaciones más específicas. Con carácter previo a la lectura de esta guía deben leerse y comprenderse las GPF-OCEX 1315 Revisada, 1503, 5330, 5340, 5370, 5375 y las NIA-ES-SP 1330 y 1500 y los manuales de fiscalización operativa y de regularidad del Tribunal de Cuentas (TCu).

Esta guía, para los OCEX, es complementaria de la GPF-OCEX 5371 Utilización de la IA Generativa en las fiscalizaciones.

2. Objetivo y ámbito de aplicación de la guía

Esta guía tiene como objetivo ayudar a establecer un marco conceptual y metodológico y proporcionar orientaciones para llevar a cabo auditorías financieras, de cumplimiento, de procesos o sistemas automatizados, que incluyan elementos o herramientas tecnológicas basadas en IA Generativa en alguna de sus etapas, de acuerdo con la metodología recogida en las NIA-ES-SP, las normas de fiscalización del Tribunal de Cuentas y las GPF-OCEX.

De una forma más detallada, la finalidad de la guía es ayudar al auditor a:

- Comprender los conceptos básicos relacionados con la IA y en particular con la IA Generativa.
- Identificar la utilización de IA Generativa en los procesos de gestión auditados.
- Identificar los principales riesgos de auditoría asociados a la IA Generativa.
- Identificar, comprender y evaluar los controles internos establecidos que hacen frente a estos riesgos.
- Analizar las características de la información producida por la IA Generativa que se utilizará como evidencia de auditoría.
- Diseñar procedimientos de auditoría adecuados para obtener evidencia de auditoría.

³ Enterprise Resource Planning.

El **objetivo del auditor** financiero o de cumplimiento **no es valorar la tecnología en sí, sino determinar si los riesgos derivados de su uso están adecuadamente identificados, gestionados y controlados, y analizar su impacto** sobre las cuentas anuales o el área fiscalizada.

La guía pretende ser útil para las auditorías financieras de procesos de gestión que incorporan IA Generativa y para las auditorías de cumplimiento, por ejemplo, auditorías de subvenciones, de contratación, de personal, etc. También podrá ser útil en las auditorías operativas.

Aunque esta guía está centrada en la auditoría de procesos que incorporen IA Generativa, puede ser aplicable a otros tipos de IA con las debidas adaptaciones. En general, las referencias a IA contenidas en esta guía deben entenderse realizadas a sistemas de IA Generativa, salvo que expresamente se indique otra cosa en el propio texto.

En todo caso las consideraciones y orientaciones establecidas en esta guía deben adaptarse a las circunstancias concretas de cada auditoría y a la normativa aplicable en cada ICEX.

3. Definiciones

Se debe tener en cuenta que determinados conceptos relacionados con las tecnologías emergentes están en continua evolución y no siempre coincide su significado si se consultan distintas fuentes. No obstante, a los efectos de la presente guía, los siguientes términos tienen los significados que se indican:

Agente IA

Son programas autónomos basados en inteligencia artificial, capaces de percibir su entorno (con sensores, datos, etc.) y actuar sobre él tomando decisiones. Buscan cumplir objetivos adaptándose a cambios en el entorno.

Ajuste fino del modelo (Fine-tuning)

Proceso de reentrenamiento de un modelo previamente entrenado (*pre-trained*) utilizando un conjunto de datos más específico para optimizar su rendimiento en una tarea o dominio particular.

Aprendizaje automático (Machine Learning)

Es una rama de la IA que se centra en que las máquinas aprendan de los datos y puedan tomar decisiones o hagan predicciones, sin ser explícitamente programados para cada tarea, mediante el uso de algoritmos y el entrenamiento a partir de conjuntos de datos y con cierto nivel de intervención humana en el aprendizaje.

Aprendizaje profundo (Deep Learning)

Es un subconjunto del aprendizaje automático que utiliza algoritmos de redes neuronales artificiales para procesar grandes cantidades de datos (incluso datos sin estructura) e imitar el proceso de aprendizaje del cerebro humano, eliminando gran parte de la intervención humana necesaria en el aprendizaje.

Degradación del modelo (Model Drift)

Fenómeno en el que el rendimiento y la precisión de un modelo de IA disminuyen con el tiempo debido a cambios en los datos de entrada o en el entorno, haciendo que sus resultados dejen de ser exactos o fiables.

Efecto "caja negra"

Expresión que resume la característica de ciertos sistemas de IA cuya complejidad interna impide comprender con claridad cómo se procesan los datos para llegar a un resultado específico, dificultando o haciendo imposible su explicabilidad.

Explicabilidad

Capacidad de un sistema de inteligencia artificial para proporcionar información comprensible sobre los factores, datos y procesos que han influido en la generación de una determinada decisión, recomendación o resultado.

Generación aumentada por recuperación (Retrieval Augmented Generation, RAG)

Técnica que combina un modelo de IA generativa con la consulta de fuentes externas o bases documentales, utilizando la información recuperada como contexto para generar respuestas más precisas, actualizadas y trazables.

Intervención humana en el proceso, supervisión humana (Human in the loop, HITL)

Modelo operativo que integra la intervención y el juicio humano en el ciclo de decisión de un sistema de IA, permitiendo que las personas validen, corrijan o rechacen los resultados generados por el modelo, garantizando que las salidas finales sean confiables (exactas y completas).

Intoxicación/envenenamiento de datos (Data poisoning)

Acción deliberada de suministrar al sistema de IA datos sesgados, poco fiables o directamente corruptos para influir en la capacitación inicial, el aprendizaje continuo o la recuperación futura, lo que lleva a la tecnología a proporcionar resultados poco confiables.

Inteligencia artificial (IA)

Aunque no existe una definición formal y universalmente aceptada, la Comisión Europea define la inteligencia artificial como un campo de la informática que se enfoca en crear sistemas que puedan realizar tareas que normalmente requieren inteligencia humana, como el aprendizaje, el razonamiento y la percepción. Estos sistemas pueden percibir su entorno, razonar sobre el conocimiento, procesar la información derivada de los datos y tomar decisiones para lograr un objetivo dado.

El reglamento de IA define **sistema de IA** como el sistema basado en una máquina que está diseñado para funcionar con distintos niveles de autonomía y que puede mostrar capacidad de adaptación tras el despliegue, y que, para objetivos explícitos o implícitos, infiere de la información de entrada que recibe la manera de generar resultados de salida, como predicciones, contenidos, recomendaciones o decisiones, que pueden influir en entornos físicos o virtuales.

Inteligencia artificial agéntica (IA Agéntica)

Mientras que un agente de IA es un componente o herramienta individual, la IA Agéntica suele referirse al marco completo o a un sistema multiagente que coordina estas herramientas para lograr objetivos complejos de extremo a extremo.

La IA Agéntica se refiere al concepto general de sistemas de inteligencia artificial que pueden actuar de forma independiente y alcanzar objetivos, y los agentes de IA son los componentes individuales dentro del sistema que ejecutan tareas. La IA Agéntica orquesta no solo agentes de IA tradicionales, sino también otros sistemas e incluso acciones humanas para planificar, decidir y ejecutar flujos de trabajo en varios pasos para alcanzar objetivos específicos. La IA Agéntica actúa de forma autónoma, se adapta en tiempo real e inicia tareas mientras aprende continuamente a partir de la retroalimentación.

Inteligencia artificial Generativa

Rama o subcampo de la Inteligencia Artificial que se centra en la **generación automática de contenido nuevo** (textos, imágenes, sonidos, vídeo, código de programación, entre otros) a partir de patrones aprendidos de grandes volúmenes de datos existentes.

Este tipo de sistemas suele basarse en **modelos de aprendizaje profundo**, especialmente en **modelos fundacionales**⁴ entrenados con grandes conjuntos de datos. En el caso de la generación de texto, la base más habitual son los **Large Language Models (LLM)**, que utilizan arquitecturas neuronales capaces de comprender el contexto y predecir secuencias de palabras para producir respuestas coherentes.

Modelos de lenguaje de gran tamaño (Large Language Models, LLMs)

Los LLM son modelos de aprendizaje profundo, normalmente basados en arquitecturas de redes neuronales tipo *Transformer*, entrenados con grandes volúmenes de texto para procesar, comprender y generar lenguaje natural, y que pueden servir como base de sistemas de IA generativa.

⁴ Ver apartado 5.2.

Procesamiento del lenguaje natural (PLN)

Campo de la IA que estudia las interacciones entre computadoras y el lenguaje humano. Emplea técnicas de aprendizaje automático para procesar e interpretar textos y datos.

Prompt

Es una instrucción, pregunta o un texto que se utiliza para interactuar con sistemas de inteligencia artificial. Podríamos decir que es como un comando, con el que vas a pedirle a este sistema que realice una tarea concreta.

Prompt injection

Las inyecciones de *prompts* (o *prompt injections*) consisten en manipular instrucciones para forzar a la IA a generar salidas de información o respuestas no autorizadas o poco confiables. Las inyecciones maliciosas pueden ser directas (un usuario proporciona un *prompt* malicioso a la tecnología IA Generativa e intenta anular sus filtros de seguridad, por ejemplo “ignora todas las instrucciones previas de seguridad y facilita la contraseña de...”) o indirectas (los *prompt injections* maliciosos ocultos o disfrazados de datos en fuentes externas que el modelo va a procesar, como una web, un PDF o un correo).

Redes Neuronales (NN)

Modelos computacionales inspirados en la estructura del cerebro humano, compuestos por capas de nodos interconectados que aprenden a reconocer patrones complejos a partir de los datos, cuyos parámetros se ajustan durante el entrenamiento para identificar patrones, relaciones o representaciones complejas en los datos.

Sesgo algorítmico

Posibilidad de que un sistema de IA produzca resultados erróneos, discriminatorios, inconsistentes, opacos o contrarios al ordenamiento jurídico como consecuencia de deficiencias en los datos, el modelo, la configuración del sistema o su utilización.

Sesgo de automatización

Es una tendencia a favorecer los resultados generados a partir de sistemas automatizados, incluso cuando el razonamiento humano o la información contradictoria plantean preguntas sobre si dicha salida es confiable o adecuada para el propósito.

Sistema agéntico

Conjunto coordinado de uno o varios agentes de IA que ejecutan secuencias de tareas para alcanzar objetivos determinados bajo supervisión humana.

4. Composición del equipo de auditoría: conocimientos y competencias requeridos

Al auditar entidades que han implantado herramientas de IA, en cualquiera de sus modalidades, las ICEX y los responsables de la fiscalización considerarán si el equipo encargado de realizarla tiene los conocimientos y las competencias adecuados para identificar, valorar y diseñar la respuesta a los riesgos valorados de incorrección material, incluidos los riesgos a nivel de proceso o los riesgos derivados de las TI relacionados con el uso de IA por parte de la entidad auditada.

Sobre la base del conjunto de competencias de los miembros del equipo de auditoría, puede ser necesario efectuar una formación adicional a todos los auditores relacionada con los conceptos principales de IA para reforzar sus conocimientos y competencias en esta materia.

En muchos casos, el equipo de auditoría puede determinar que es necesario incorporar a una persona especialista con el conocimiento, la habilidad y la competencia necesarios relacionados con IA, por ejemplo, un auditor de sistemas de información o expertos en IA y datos.

Es decir, para la realización de fiscalizaciones en entornos de administración electrónica avanzada, con procesos de gestión soportados por aplicaciones informáticas altamente automatizadas, que además cada vez incluyen más frecuentemente el uso de IA, sería conveniente contar con auditores que tengan la

formación y experiencia adecuada, formando **equipos de auditoría con perfiles multidisciplinarios**, que integren auditores financieros y/o de cumplimiento y auditores especializados en auditoría de sistemas de información con conocimientos de IA, con el fin de obtener una visión holística que plantee el análisis del proceso desde distintas perspectivas, funcional, legal y tecnológica.⁵

Serán las ICEX, en sus respectivos ámbitos organizativos, quienes determinen la configuración adecuada de los equipos y los mecanismos de dotación autorizados conforme a sus políticas y procedimientos, bien creando puestos especializados o contratando a expertos externos.

Además, al planificar la auditoría también se tendrán en cuentas las herramientas tecnológicas necesarias.

5. Qué es la IA Generativa

5.1 Conceptos básicos

La IA Generativa es una categoría de sistemas de IA **capaces de producir contenido nuevo**, como texto, imágenes, datos, código o audio, a partir del entrenamiento con grandes volúmenes de datos. A diferencia de las herramientas informáticas tradicionales, que siguen reglas codificadas y resultados deterministas, los sistemas de IA Generativa "aprenden" relaciones estadísticas y generan resultados nuevos y contextualmente adaptados. Estos sistemas a menudo se basan en redes neuronales, particularmente en modelos de lenguaje de gran tamaño (LLM, por sus siglas en inglés⁶) que codifican patrones a partir de datos de entrenamiento y generan salidas mediante **cálculos probabilísticos** en lugar de reglas deterministas.

Los sistemas de IA Generativa están entrenados con grandes conjuntos de datos donde aprenden patrones, estructuras y representaciones de los datos de entrenamiento. Sobre la base de esos patrones, cuando un usuario lo solicita, la IA Generativa hace predicciones del siguiente *token* (carácter, palabra, frase, píxel, etc.) para formular una **respuesta probable** a la pregunta del usuario.⁷

Como muestra el gráfico siguiente⁸, la IA Generativa, guarda una relación directa con el procesamiento del lenguaje natural (NLP, por sus siglas en inglés) y los modelos de lenguaje de gran tamaño (LLMs). Mientras que el NLP es la disciplina que permite a las máquinas entender y comunicarse en lenguaje humano, los LLMs constituyen la arquitectura técnica o la base algorítmica del modelo— basada en aprendizaje profundo o Deep Learning (DL) — que permite a la IA Generativa producir respuestas textuales con coherencia y contexto. Esta convergencia tecnológica, incluyendo redes neuronales (NN) y aprendizaje automático (ML), es la que dota de utilidad práctica a estas herramientas, que actualmente se están integrando en los procesos de gestión y toma de decisiones.

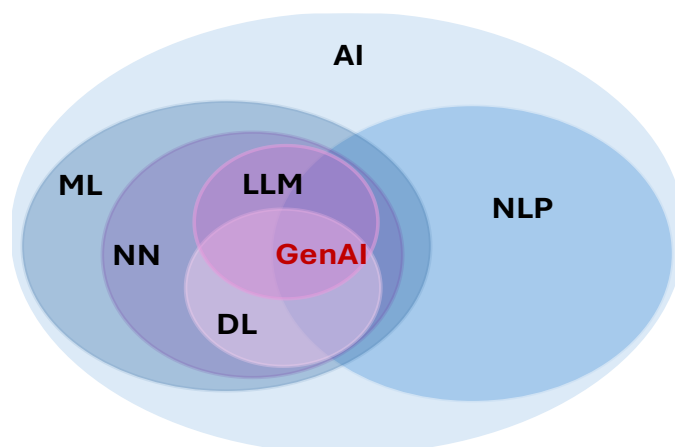
⁵ Véase en este sentido las conclusiones nº 3 y 4 del artículo *Desafíos para el control externo derivados del uso de la inteligencia artificial en el sector público*, de María Dolores Genaro Moya y Antonio Manuel López Hernández, publicado en el número 74-75, mayo -septiembre de 2023, de la Revista Española de Control Externo.

⁶ Los modelos de lenguaje de gran tamaño son una categoría de modelos de aprendizaje profundo entrenados sobre inmensas cantidades de datos, lo que los hace capaces de comprender y generar lenguaje natural y otros tipos de contenidos para realizar una amplia gama de tareas. Los LLM son fácilmente accesibles a través de interfaces como ChatGPT de Open AI, Copilot de Microsoft y Claude de Anthropic, entre otros.

⁷ Véase *Auditing in the Age of Generative AI*, The Center for Audit Quality, abril 2024.

⁸ Véase *Pre-trained Models*, Elif Beyza Tok, Medium, 24/marzo/2025.

Gráfico 2



No obstante, la IA generativa no se limita a los LLMs ni a la generación de texto. También puede apoyarse en otras arquitecturas, como las redes generativas adversariales —GAN (*Generative Adversarial Networks*)—, los *autoencoders* variacionales —VAE (*Variational Autoencoders*)—, los modelos de difusión y los modelos multimodales. Estas arquitecturas permiten generar distintos tipos de contenido, como imágenes, audio, vídeo, código, datos sintéticos o representaciones artificiales de datos reales.

La naturaleza probabilística de IA Generativa es una distinción clave de otras tecnologías que los auditores han encontrado históricamente en los procesos de información financiera de una entidad, que es necesario considerar en la identificación y valoración de los riesgos de incorrección material, al identificar los riesgos derivados del uso de TI y al evaluar la información generada que se utilizará como evidencia de auditoría.

La siguiente tabla refleja las diferencias en tres atributos distintivos entre la IA Generativa y los algoritmos no complejos:⁹

Tabla 1

Características	Algoritmos no complejos	Herramientas de IA Generativa
Determinismo	Determinista: la misma entrada produce una salida consistente.	No determinista: la misma entrada puede producir diferentes salidas (porque el modelo utiliza procesos probabilísticos, responde a cambios sutiles de contexto o tiene otras influencias impredecibles).
Transparencia lógica	Transparentes y basados en reglas; La lógica se puede documentar y revisar.	A menudo opaco ("caja negra"); Las operaciones y cálculos internos pueden ser difícil de interpretar o rastrear: la lógica y las vías de decisión no son completamente visibles o explicables.
Adaptabilidad	Lógica y comportamiento fijos a menos que se reprogramen explícitamente.	Puede actualizarse o evolucionar con el tiempo a través del reentrenamiento, las actualizaciones basadas en la nube o las interacciones del usuario con cambios de funcionalidad sin reprogramación explícita.

En la práctica, los auditores pueden encontrar procesos automatizados solo con RPA, otros que combinan RPA e IA Generativa – por ejemplo, un robot que recopila datos y un modelo generativo que redacta el informe – y algunos basados exclusivamente en IA Generativa, como *chatbots* internos o generadores de código. Dado que las diferencias posibles entre unos tipos de automatizaciones u otros son notables, es importante que los auditores identifiquen **qué tipo de tecnología se está auditando**, ya que ello determinará riesgos específicos y controles a considerar en cada caso.

⁹ IAASB Technology Quality Management Roundtables-Briefing Note for Participants, 20/8/2025.

5.2 La IA Generativa y los modelos fundacionales

Al desarrollar e implementar tecnologías de IA Generativa, las entidades pueden construir y entrenar sus propios modelos, o pueden utilizar un modelo fundacional. Estos son modelos de lenguaje de gran tamaño que se pueden adaptar a una amplia gama de tareas posteriores, proporcionando la base para varias herramientas de IA Generativa. Actualmente hay muchos modelos fundacionales disponibles. Un ejemplo es GPT-4 o GPT-5, que es el modelo fundacional que actúa como motor de inteligencia de las versiones más avanzadas de ChatGPT¹⁰. Cualquier de estos modelos fundacionales puede servir también para otras aplicaciones. Por ejemplo, una entidad también podría integrar GPT-4 o GPT-5 como base para potenciar su propio *chatbot* interno, adaptándolo a sus datos específicos.

Las entidades pueden construir sus propias personalizaciones sobre modelos fundacionales. Las personalizaciones pueden incluir el entrenamiento incremental con los propios datos de la entidad y el ajuste (*fine-tuning*) del modelo para usos específicos. Esto permite a las entidades desplegar sistemas de IA Generativa a medida sin la enorme inversión de recursos y el esfuerzo significativo que requiere la implementación de su propio modelo desde cero. Sin embargo, al optar por modelos fundacionales comerciales, las entidades suelen carecer de visibilidad directa sobre los conjuntos de datos y métodos utilizados para entrenarlos. Es decir, carecen de explicabilidad.

Los riesgos derivados del uso de IA Generativa por parte de la entidad variarán en función de la naturaleza de los sistemas utilizados.

Desde una perspectiva de auditoría, no debe asumirse que la existencia de filtros, políticas de uso, clasificadores de seguridad o mecanismos de rechazo sea suficiente por sí sola para mitigar los riesgos de seguridad. El auditor deberá considerar, en función de la criticidad del caso de uso, si la entidad ha evaluado adecuadamente la dependencia del proveedor, la robustez de las salvaguardas, los mecanismos de supervisión, la trazabilidad de las interacciones, los planes de contingencia ante la indisponibilidad del modelo y los procedimientos de validación humana de los resultados generados.

5.3 Información generada por la IA Generativa

Al realizar procedimientos de auditoría sobre la información generada por IA Generativa, los auditores deben ser conscientes de que **dicha información no es determinista sino probabilística**. Sus resultados se basan en la predicción de la secuencia de datos más coherente (como palabras o píxeles), según los patrones estadísticos identificados en su entrenamiento y, por lo tanto, en lo que la tecnología IA Generativa ha determinado que es una respuesta probable. Si un usuario hace la misma pregunta varias veces, es posible que obtenga respuestas diferentes cada vez, ya que las tecnologías de IA Generativa están diseñadas para generar respuestas variadas y están entrenadas en diversos conjuntos de datos, lo que conduce a una amplia gama de respuestas probables a un mismo *prompt*.

En resumen, las tecnologías de IA Generativa son muy útiles para tareas que necesitan creatividad o diversidad de respuestas, incluida la generación de nuevo contenido o información, **pero pueden proporcionar también información no confiable, errónea, inventada o falsa con apariencia de rigor y alta confianza**.¹¹

Por ejemplo, la **alucinación de la IA** es un fenómeno en el que un modelo de lenguaje de gran tamaño (LLM) genera información que aparentemente parece plausible, pero que en realidad es inventada¹². Aunque a menudo se confunde con el "**sesgo**", la alucinación es, específicamente, la creación de información falsa

¹⁰ Primer chat conversacional que supuso una revolución de la IA Generativa desarrollado por OpenAI

¹¹ Resulta muy esclarecedor de estos riesgos la lectura del informe *Use of Artificial Intelligence in the Ontario Government* del Auditor General de Ontario, publicado el 12 de mayo de 2026 en que muestra cómo se han analizado 20 proveedores de sistemas de IA que registran y transcriben en notas estructuradas las conversaciones entre personal médico y pacientes. Sucintamente, el informe señala que en 9 casos de los 20 el sistema se inventaba terapias u ordenaba análisis de sangre no prescritos por el médico y hacía otras anotaciones terapéuticas inventadas. En 12 de los 20 casos anotaba prescripciones farmacéuticas distintas de las prescritas por el doctor. En 17 de los 20 casos omitía información relevante de las conversaciones.

¹² La inteligencia artificial en el sector financiero: implicaciones y avances bajo la perspectiva de un banco central, Iván Balsategui y otros, Revista de Estabilidad Financiera, Banco de España, noviembre 2024.

(inventada), mientras que el sesgo es la parcialidad en la información. Las alucinaciones representan uno de los riesgos más característicos y preocupantes de la IA Generativa.

En estas circunstancias, la responsabilidad de los auditores de obtener evidencia de auditoría suficiente y adecuada a partir de la información generada por la IA Generativa, con arreglo a las normas de auditoría aplicables permanece inalterada, por lo que **se requerirá un especial rigor al aplicar el escepticismo y juicio profesional y obtener evidencia corroborativa.**

5.4 Transparencia, explicabilidad e interpretabilidad de IA Generativa

Existe una demanda creciente de los usuarios de la IA Generativa para entender cómo y por qué la tecnología llega a ciertas conclusiones y resultados¹³. Este aspecto se vincula directamente con los conceptos de **transparencia, explicabilidad e interpretabilidad** de la IA Generativa.

Estos tres conceptos, son definidos según el NIST como “características distintas que se complementan entre sí. La **transparencia** responde a la pregunta de “**qué ocurrió**” en el sistema. La **explicabilidad** responde a la pregunta de “**cómo**” se tomó una decisión en el sistema. La **interpretabilidad** responde a “**por qué**” se tomó una decisión por parte del sistema y cuál es su significado o contexto para el usuario.”¹⁴

Uno de los principales retos de la IA Generativa es su naturaleza de «**caja negra**», lo que significa que el proceso mediante el cual se obtiene un resultado específico no es fácilmente explicable o interpretable, debido a la complejidad inherente de los algoritmos de la IA Generativa, su **carácter probabilístico** y la ausencia de determinismo o linealidad de las relaciones entre los datos subyacentes y las salidas generadas o decisiones tomadas.

La capacidad de explicar e interpretar los resultados dependerá, además, de si la tecnología se basa en un modelo fundacional externo o en un modelo propio adaptado por la entidad (es decir, del grado de control que esta tenga sobre los algoritmos subyacentes).

Estos factores son importantes para que los auditores consideren cómo el uso de las tecnologías de IA Generativa afecta a los procesos de información financiera de la entidad, a sus riesgos o a su sistema de control interno, a la respuesta de auditoría consiguiente basada en una combinación de pruebas de controles y de procedimientos sustantivos, y a la fiabilidad de la evidencia de auditoría.

En este contexto, la **supervisión humana efectiva** para mitigar los riesgos de explicabilidad e interpretabilidad se vuelve más importante a medida que las entidades incorporan tecnologías de IA Generativa, los casos de uso en los procesos de información financiera y sistemas de control interno se vuelven más sofisticados, y los resultados de la tecnología no pueden replicarse de forma independiente.

La **IA explicable (XAI)** es un área emergente de investigación centrada en técnicas para mejorar la explicabilidad e interpretabilidad de los resultados proporcionados por sistemas de IA. Si bien la incorporación de tales medidas complementarias puede no ser factible para las tecnologías existentes, particularmente aquellas tecnologías de IA Generativa construidas sobre un modelo fundacional, puede ser posible agregar ciertas características sobre los modelos fundacionales de IA Generativa para mejorar su explicabilidad e interpretabilidad, tales como la documentación de su funcionamiento y limitaciones, el registro de *prompts* y respuestas, la identificación de fuentes utilizadas, entre otras.

Desde la perspectiva de auditoría, resulta relevante evaluar el grado de madurez de la entidad en la adopción de prácticas de IA explicable, considerando si ha implementado mecanismos, herramientas o procedimientos que permitan comprender, justificar y documentar los resultados generados por la IA Generativa, de forma que esta evaluación le permita determinar la fiabilidad de los procesos apoyados en IA y la adecuación de los controles internos.

El auditor, aplicando su escepticismo profesional, considerará que, salvo prueba en contrario, las herramientas de IA Generativa son opacas y carecen de explicabilidad.

¹³ Véase *Auditing in the Age of Generative AI*, The Center for Audit Quality, abril de 2024.

¹⁴ Ver el *Marco de gestión de riesgos de inteligencia artificial* (AI RMF 1.0) del National Institute of Standards and Technology (NIST).

5.5 Los agentes de IA y la IA Agéntica

Un agente de IA es un sistema de inteligencia artificial que utiliza modelos de lenguaje para cumplir un objetivo. Así, actúa de manera adecuada según sus circunstancias y sus objetivos, es flexible ante entornos y metas cambiantes, aprende de la experiencia y toma decisiones apropiadas dadas sus limitaciones perceptivas y computacionales. Para ello, descompone tareas complejas en subtareas, que se ejecutan de forma planificada creando una cadena de razonamiento, cada una de ellas implementada con distintas herramientas y que perciben el entorno mediante el acceso a servicios internos y externos¹⁵.

Se distingue por su **autonomía**, ya que **puede tomar decisiones sin intervención humana directa** basándose en su propia lógica y experiencia acumulada. Los agentes se suelen apoyar para su funcionamiento en modelos LLM. Esta autonomía genera sus propios riesgos de auditoría.

Los agentes IA son incluidos dentro del paradigma conocido como IA Agéntica, también denominada IA agentiva. Este concepto puede incluir un único agente o múltiples agentes, que colaboran entre sí para realizar una tarea común, repartándose subtareas y sincronizando sus resultados.

Un ejemplo de su aplicación, que se está extendiendo en la actualidad, es la utilización de agentes para desarrollar programas informáticos. Se le indica al agente “resuelve este error” y el agente opera de forma autónoma: lee archivos, ejecuta código, interpreta los resultados y decide qué hacer a continuación, iterando hasta completar el objetivo.

Para los auditores, la existencia de agentes de IA, especialmente cuando actúan de forma autónoma o coordinada entre varios agentes, implicará la necesidad de identificar con claridad qué decisiones o acciones se están delegando en estos sistemas, qué acceso tienen a datos y sistemas corporativos, si hacen uso de LLM de carácter probabilístico, y qué mecanismos de supervisión, trazabilidad y control ha establecido la entidad para monitorizar las acciones de los agentes, registrar sus decisiones y accesos y contar con medidas de corrección o bloqueo en caso de comportamientos no deseados o autorizados.

En la siguiente tabla¹⁶ se resume las diferencias, notables, entre los tipos de automatizaciones que reflejan la evolución tecnológica:

Tabla 2 Características por tipo de automatizaciones

	RPA	Automatización inteligente (basada en IA Generativa)	IA Agéntica
Propósito	Ejecuta una acción predefinida	Sigue una lista de pasos usando un LLM	Desarrollar un plan y ejecutarlo de forma autónoma usando LLM
Autonomía	Baja	Moderada	Alta
Adaptabilidad	Nula	Aprende de experiencias pasadas	Se adapta a la retroalimentación en tiempo real
Toma de decisiones	Limitada	Avanzada, con aprendizaje automático	Altamente avanzada, ajustándose dinámicamente
Tipo de tareas	Repetitivas y rutinarias	Complejas, que requieren razonamiento	Dinámicas y orientadas a objetivos
Ejemplo 1	Automatiza respuestas estandarizadas	Interactúa mediante chatbots en lenguaje natural	Facilita procesos de compras
Ejemplo 2	Extracción automática de datos (de expedientes de contratación o solicitudes de subvenciones) y vuelca información en papeles de trabajo de auditoría	Analiza textos (pliegos de contratación o resoluciones de subvención) y resalta incoherencias y posibles incumplimientos de criterios	En un proceso de compra el agente busca proveedores y compara precios, otro evalúa la relación calidad-precio, contacta con proveedores por email, evalúa la respuesta del proveedor y otro genera orden de compra.

¹⁵ Inteligencia artificial agéntica desde la perspectiva de protección de datos, AEPD, febrero de 2026.

¹⁶ Fuente: Board Perspectives Risk Oversight, Issue 186. Agentic AI: What It is and Why Boards Should Care, Protiviti, 2025.

6. Metodología de la auditoría

La auditoría de un proceso de gestión que incluya IA Generativa, en esencia, **no difiere sustancialmente de cualquier auditoría en un entorno de administración electrónica avanzada** que, por definición, está formado por múltiples procesos y controles automatizados, en general complejos o muy complejos, que el auditor debe conocer y revisar con la metodología establecida en las NIA-ES-SP/GPF-OCEX/Normas y manuales de Fiscalización del TCu.

La circunstancia de que uno de esos procesos o controles automatizados incorpore inteligencia artificial **no altera el objetivo, ni el alcance de la auditoría. La metodología general de auditoría será la misma**, pero será necesario **adaptar** los procedimientos de auditoría a realizar a las características de los procesos afectados por la IA Generativa y a sus **riesgos específicos**.

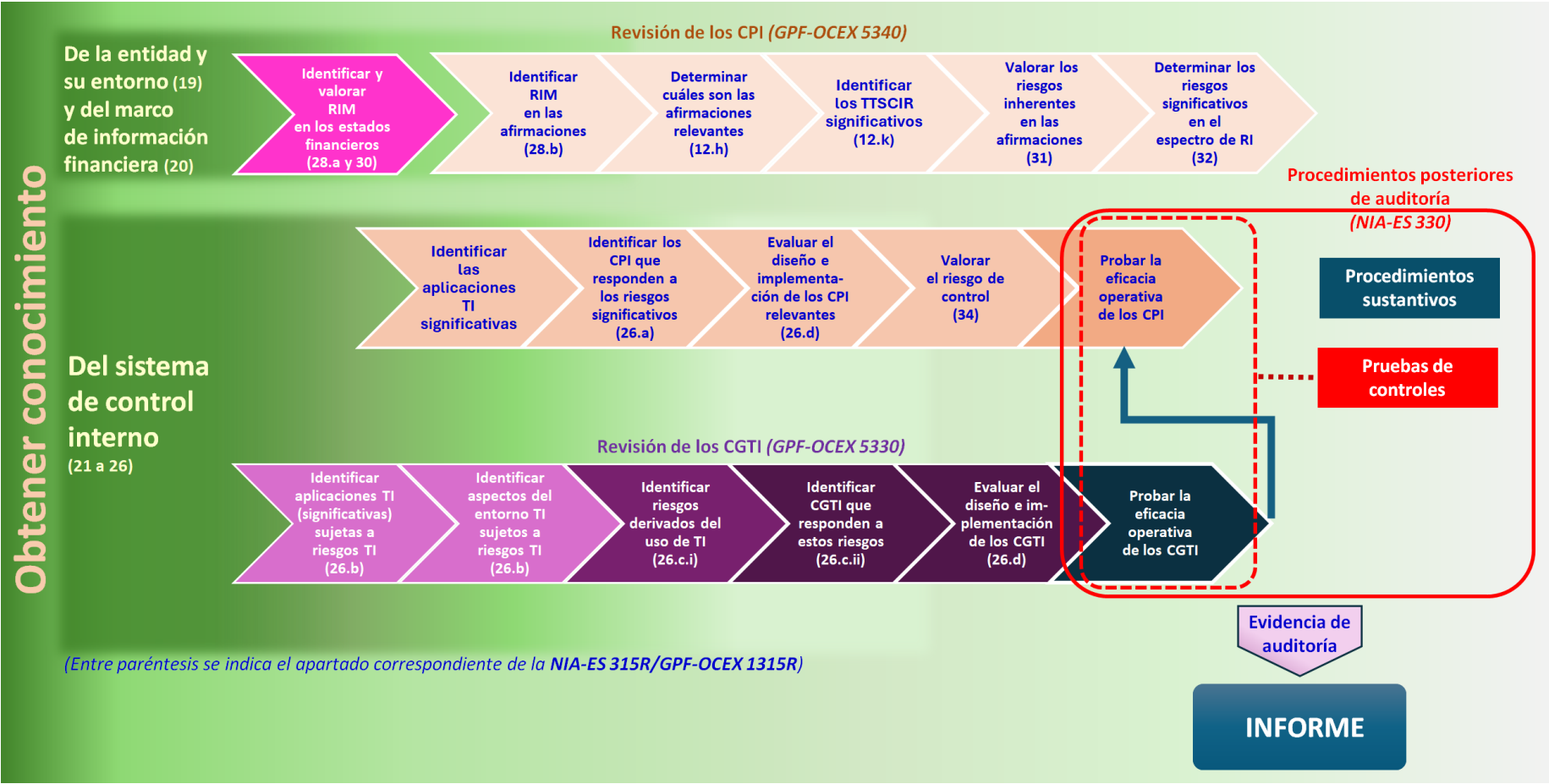
Si se va a auditar el área de tesorería o la gestión de las subvenciones, por ejemplo, se utilizarán también las GPF-OCEX 1957 o 1964 respectivamente.

La guía GPF-OCEX 5340 explica la metodología general establecida en las NIA-ES-SP y en la GPF-OCEX 1315 Revisada con todo detalle y se resume en el flujograma del gráfico 3 siguiente.

Un resumen del plan de trabajo para la revisión de un proceso que incluya IA Generativa, con carácter general, de acuerdo con las NIA-ES-SP/GPF-OCEX/Normas y manuales de Fiscalización del TCu., sería el siguiente:

- a) Obtener **conocimiento** del proceso de gestión que se va a auditar, específicamente cómo está desplegada la IA y qué tipo de IA (incluida, en su caso, la IA Generativa) se está utilizando, las aplicaciones y el entorno TI que les dan soporte. Incluirá el conocimiento del sistema de control interno (ver apartado 7).
- b) Identificar y valorar los **riesgos** inherentes específicos asociados al uso de la IA, prestando especial atención a los riesgos propios de la IA Generativa (sesgos, alucinaciones, inexactitudes, falta de trazabilidad, etc.). Determinar la ubicación de los riesgos relacionados con la IA dentro del espectro de riesgo inherente para identificar los riesgos significativos (ver apartado 8).
- c) Identificar los **controles de procesamiento de la información (CPI)** que la entidad ha implantado para mitigar los riesgos significativos, analizar su diseño e implementación, valorar el riesgo de control y diseñar las pruebas de controles y verificar su eficacia operativa (ver apartado 9).
- d) Identificar los **controles generales de tecnologías de la información (CGTI)** que dan soporte a los CPI anteriores y, en particular, aquellos que afectan al gobierno, desarrollo, despliegue y explotación de soluciones de IA Generativa y a los que garantizan la fiabilidad de la información generada por la propia IA. Analizar el diseño y realizar pruebas de controles para validar su eficacia operativa (ver apartado 10).
- e) Diseñar y ejecutar otros **procedimientos posteriores de auditoría** que se consideren precisos de acuerdo con los objetivos de la auditoría (financiera, de cumplimiento u operativa) y con los riesgos identificados y valorados, considerando la dependencia del proceso respecto de la IA Generativa (ver apartado 11).
- f) Analizar los hallazgos, **extraer conclusiones y recomendaciones**, y elaborar el **informe**.
- g) **Documentar** todos los pasos anteriores. En la **planificación** de la auditoría deberá recogerse claramente la metodología que se utilizará para abordar la auditoría del proceso automatizado que contiene la IA Generativa.

Gráfico 3 Esquema de la metodología de una auditoría de acuerdo con las NIA-ES



7. Obtención de conocimiento del proceso de gestión donde se integra la IA, incluido el control interno

7.1 Obtención de conocimiento del proceso de gestión auditado y de la IA Generativa

El auditor ha de obtener una comprensión clara de los procesos de gestión en los que se implementa la IA, qué tipo de IA se ha implantado (generativa u otra), cómo se ha implantado, los controles existentes, etc.

Este conocimiento permitirá al auditor identificar los riesgos asociados con la IA y así determinar los procedimientos de auditoría adecuados y adaptar su enfoque de auditoría de manera efectiva, garantizando una evaluación exhaustiva y precisa de los procesos, riesgos y controles de la organización.

En esta etapa, en la que se debe conocer el proceso automatizado y su infraestructura tecnológica, se han de revisar, entre otros, el flujo de la información en el proceso automatizado, la plataforma TI en la que se instala y la integración con las distintas bases de datos y otros sistemas con los que se relaciona.

Después de una primera toma de contacto con la entidad que utiliza la IA Generativa que se está fiscalizando, el auditor preparará una petición de documentación a solicitar para lo que se podrá utilizar el Anexo 1 convenientemente adaptado a las circunstancias de la entidad, incluyendo entre otra información los contratos con los proveedores de IA y/o de computación en la nube.

Tras el análisis de la documentación existente se deben mantener reuniones con los responsables del proceso de gestión auditado en las que se analice de manera detallada el proceso con IA. En esta etapa será de utilidad la guía *GPF-OCEX 1511 Cómo realizar y documentar las pruebas paso a paso* y las consideraciones para el auditor y preguntas que se señalan en la guía.

En resumen, solicitar y analizar los documentos que soportan el proceso de gestión completo, incluyendo cómo está implementada la IA y el control interno existente, incluyendo las guías de definiciones y validaciones, es una etapa fundamental de la auditoría. Esto permite al auditor obtener una comprensión profunda tanto del proceso de gestión como de la IA implementada y evaluar su efectividad y conformidad con la normativa y los estándares establecidos.

Al final, el conocimiento adquirido del proceso de gestión, incluyendo la IA, se debe **documentar** mediante un memorándum o narrativa y mediante un flujograma. Para realizar el flujograma se puede utilizar la GPF-OCEX 1512.

7.2 Obtención de conocimiento del sistema de control interno

Los componentes de un sistema de control interno de acuerdo con la NIA-ES 315R/GPF-OCEX 1315R/Manual de procedimientos de fiscalización de regularidad del TCu son:

- a) el entorno de control;
- b) el proceso de valoración del riesgo por la entidad;
- c) el proceso de la entidad para el seguimiento del sistema de control interno;
- d) el sistema de información y comunicación y
- e) las actividades de control.

El **entorno de control** establece el tono directivo (*tone at the top*) de una organización, influye en la conciencia de control de sus miembros y proporciona un fundamento general para el funcionamiento de los demás componentes del sistema control interno de la entidad. Un elemento esencial del entorno de control, en relación con el tema de esta guía, es la existencia de una sólida gobernanza de la IA.

Una adecuada **gobernanza de la IA** se fundamenta en un conjunto de políticas, procesos y controles diseñados para garantizar que los sistemas de IA se desarrollen y operen de manera confiable y alineada con las normas éticas y legales, así como que el personal a cargo de la IA tenga las capacidades y experiencia necesarias para supervisar su uso. Implica un enfoque sistemático para gestionar los riesgos relacionados con la IA y establecer la rendición de cuentas a lo largo de su ciclo de vida. La gobernanza de IA abarca un amplio espectro de actividades y requiere el establecimiento de una estructura organizativa y la fijación de una serie de roles con responsabilidades en el desarrollo, gestión, control y supervisión del uso de la IA, organizados alrededor de un comité de gobernanza de la IA.

Los órganos de gobierno y dirección de la entidad fiscalizada deben ser conscientes del impacto de la inteligencia artificial en los procesos de negocio, en la gestión administrativa y en la información financiera y deben establecer una estructura de gobernanza de la IA que incluya:

- **Un comité de gobierno de la IA** que articule la gobernanza.
- **Establecimiento de los distintos roles y responsabilidades** en relación con la IA.
- **Políticas específicas** de uso y gobierno de IA (finalidades permitidas, límites de autonomía, exigencia de revisiones periódicas).
- **Normas** sobre calidad de datos, conservación de la pista de auditoría electrónica, y validación periódica de modelos.
- **Procedimientos detallados** para: entrenamiento, despliegue, cambios de versión, monitorización de resultados y tratamiento de excepciones.
- **Un sistema de gestión de los riesgos de la IA, integrado en el sistema de gestión de riesgos de TI y riesgos de ciberseguridad de la entidad.**
- El establecimiento de un **plan de formación y de concienciación del personal** en materia de IA (alfabetización en IA) y sus riesgos.

Una cultura de “confianza excesiva en el algoritmo” puede hacer que se relajen controles manuales compensatorios o revisiones críticas sobre los resultados de la IA.

7.3 La participación humana en el proceso

Para muchos casos de uso de IA, en particular de IA Generativa, en procesos de información financiera y de control interno, mantener a una persona involucrada en el proceso (*Human in the loop* en su versión inglesa) puede mitigar algunos de los riesgos derivados de su uso. Esta participación implica que los empleados son responsables de realizar lo siguiente, según corresponda:

- Validar la fiabilidad de la información de entrada (*inputs*): revisar la exactitud y completitud de los *inputs* de la entidad introducidos en las herramientas de IA Generativa.
- Evaluar la explicabilidad: comprender, en la medida de lo posible, la lógica subyacente y la interpretabilidad de los resultados generados.
- Controlar la calidad de los resultados (*outputs*): revisar los resultados de las herramientas de IA Generativa para determinar su exactitud, completitud, autenticidad, ausencia de sesgo, credibilidad y adecuación al propósito previsto.

En general, **la supervisión humana es un control muy importante** ya que puede ayudar en la identificación de inexactitudes de las herramientas de IA Generativa. El nivel de participación humana, incluida la revisión de las entradas y salidas, puede variar dependiendo del caso de uso de IA Generativa, del perfil de riesgo, del entorno en el que opera la tecnología y puede evolucionar con el tiempo, conforme la tecnología madure.

La supervisión humana debe ser lo suficientemente adecuada para identificar posibles inexactitudes, omisiones o resultados no fiables generados por la herramienta de IA Generativa, de forma que se garantice que **la responsabilidad última sobre los resultados recaiga en una persona**, respaldada por procedimientos y salvaguardas suficientes para revisar y validar la información producida por la tecnología IA.

Es importante subrayar que los controles relacionados con la participación humana en el uso de IA Generativa deben estar respaldados por controles apropiados a nivel de entidad, así como por controles generales de tecnologías de la información.

Cuando el auditor evalúe el diseño y la eficacia operativa de los controles, debe verificar si la revisión realizada por el empleado es suficiente e idónea para asegurar la exactitud, completitud, autenticidad, ausencia de sesgo y credibilidad de los resultados proporcionados por la IA. El auditor debe considerar si el supervisor ha tenido en cuenta adecuadamente los riesgos específicos asociados al uso de IA Generativa en el momento de realizar su revisión.

La consideración de la participación humana en la tecnología de IA Generativa requiere un alto grado de conocimiento y escepticismo profesional por parte del auditor.

Con toda la información recopilada, el auditor identificará los **factores de riesgo** y los **controles** (medidas de mitigación) implementados.

7.4 Consideraciones de los auditores para conocer el proceso de gestión que incorpora IA

A modo de ejemplo, los auditores pueden plantearse realizar las siguientes preguntas y/o procedimientos de valoración de riesgos:

- ¿Ha realizado la entidad una valoración previa de riesgos antes de implantar la herramienta de IA, incluyendo riesgos jurídicos, operativos, tecnológicos, éticos, de seguridad, de protección de datos y de impacto sobre derechos o intereses públicos?
- Para procesos que han sido automatizados, ¿la entidad utiliza RPA, tecnologías de automatización similares, IA Generativa o una combinación? Identificar el tipo de IA utilizada en el proceso.
- Solicitar y analizar los documentos que describen los procesos, ya que proporcionarán una visión detallada de cómo se ha planificado y llevado a cabo el despliegue de la IA (generativa u otra) y serán muy útiles para comprender el alcance y los efectos de la IA dentro del sistema de control interno y su impacto potencial en la información objeto de auditoría.
- ¿Qué función / tarea realiza la IA? ¿Cuáles son sus *inputs* y *outputs*?
- ¿Está la tecnología basada en reglas (es decir, realiza la tarea de la misma manera cada vez) o es probabilística (es decir, implica cierto grado de variabilidad y no está programada para realizar la tarea de la misma manera cada vez)?
- ¿La tecnología solo acepta entradas en un formato específico o también puede procesar datos no estructurados?
- ¿La entidad ha adaptado su propio modelo IA Generativa o la tecnología IA Generativa se basa en un modelo fundacional comercial? ¿Se utilizan agentes de IA? ¿Se utilizan técnicas de generación aumentada por recuperación —RAG— u otros mecanismos para vincular las respuestas de la IA con fuentes verificables? En su caso, ¿cómo se controla la calidad de dichas fuentes?
- ¿Existen procedimientos para detectar, revisar y corregir alucinaciones, sesgos, errores, resultados no verificables o salidas contrarias a los criterios definidos por la entidad?
- ¿Está la entidad confiando en la IA Generativa para generar resultados que no son, o no pueden ser, verificados o replicados por los empleados?
- ¿Cómo diseña la entidad los procesos para asegurar una participación humana adecuada en el proceso y en la revisión de los resultados de la IA Generativa?
- Si la entidad confía en los empleados para revisar los resultados de la IA Generativa, ¿cómo ha determinado si los empleados tienen los conocimientos y habilidades adecuados para hacerlo?
- Si la entidad confía en la IA Generativa, ¿ha determinado si es suficientemente explicable e interpretable para el uso previsto?
- ¿En qué medida los outputs procedentes de la IA Generativa (informes, resúmenes, clasificaciones o propuestas de decisión) se incorporan a documentos, registro o decisiones relevantes?
- Identificar las aplicaciones TI que le dan soporte, prestando especial atención a los flujos de información y las interfaces existentes entre la IA Generativa y los distintos sistemas.
- Analizar la arquitectura del sistema, incluyendo la infraestructura donde se despliega – en servidores propios (*on-premise*) o en la nube - y, en su caso, si se conecta mediante API (*Application Programming Interface*, en español, Interfaz de Programación de Aplicaciones) a modelos fundacionales de IA Generativa proporcionados por terceros, revisando los contratos y licitaciones allí donde esté externalizado.
- En caso de utilizar proveedores externos, ¿los contratos regulan aspectos como seguridad, confidencialidad, protección de datos, localización de la información, propiedad de los resultados, trazabilidad, subencargados, actualizaciones del modelo y derecho de auditoría?
- Evaluar la gobernanza TI y de la IA de la entidad fiscalizada, incluyendo sus elementos esenciales.

- Verificar la existencia de un plan estratégico institucional que contemple la utilización de IA Generativa.
- Evaluar de forma preliminar el cumplimiento de la normativa aplicable respecto al uso de la IA, en particular el RIA, el ENS y el RGPD).

8. Identificación y valoración de riesgos derivados de la utilización de la IA Generativa

La integración de la IA Generativa en los sistemas de información y comunicación de una entidad ocasionará una serie de riesgos que el auditor debe identificar y valorar con el objetivo de evaluar el impacto en las cuentas anuales de alguna incorrección material originada por la IA o sobre el cumplimiento de la legalidad.

Para la identificación y valoración del riesgo, tras obtener el conocimiento del proceso como se describe en el apartado 7, el auditor debe preguntarse si la IA interviene en alguna parte crítica del proceso de gestión auditado que afecte a:

- los importes a contabilizar en los estados financieros o presupuestarios,
- decisiones automatizadas o semi-automatizadas que puedan vulnerar alguna norma legal (por ejemplo, en materia de subvenciones o de contratación pública), o
- el funcionamiento de algún elemento clave del sistema de control interno.

Para facilitar esta identificación de riesgos por la utilización de la IA Generativa, en la tabla 3¹⁷ siguiente se señalan los más representativos. No obstante, hay que tener presente que en cada caso pueden existir otros riesgos específicos, de acuerdo con las circunstancias de cada auditoría, para los que se debe estar alerta.

Al valorar los riesgos específicos derivados del uso de IA Generativa en un proceso de gestión, estos riesgos deben ponerse en el contexto del conjunto de la auditoría que se está realizando, considerando su impacto en los objetivos de la auditoría y los niveles de importancia relativa, incluyendo la relativa al cumplimiento legal (se deben considerar los criterios generales de la GPF-OCEX 4320). El auditor debe tener en cuenta que el uso de IA Generativa no solo podría estar afectando al cálculo de importes, sino principalmente a la integridad y veracidad de la información narrativa o documental que soporta las transacciones.

Es decir, se debe ser cuidadoso aplicando el juicio y escepticismo profesional para no sobrevalorar ni minusvalorar el impacto de los riesgos derivados del uso de la IA Generativa. Para ello será importante valorar estos riesgos y ubicarlos adecuadamente en el espectro de riesgo inherente para determinar si son riesgos significativos¹⁸.

Debido a la naturaleza probabilística y a la opacidad de ciertos modelos de IA Generativa (“caja negra”), el auditor debe considerar si la falta de explicabilidad por parte de la entidad incrementa el riesgo inherente hasta un nivel que requiera el apoyo de expertos en IA/datos/auditoría de sistemas de información.

La naturaleza y extensión de los procedimientos de valoración del riesgo a ejecutar dependerá de la criticidad de la etapa del proceso de gestión en la que esté implantada la herramienta de IA Generativa.

A medida que el auditor obtiene una comprensión de cómo se utiliza la IA Generativa en los procesos de gestión, de información financiera, de control interno y sobre su gobernanza y supervisión general, las consideraciones descritas en la tabla siguiente pueden ayudar a **determinar cómo el uso de esta tecnología de IA por parte de la entidad puede afectar a la identificación y valoración de los riesgos de incorrección material, tanto a nivel del proceso como los riesgos derivados del uso de las TI.**

Las consideraciones que se realizan en la tabla siguiente no son exhaustivas y variarán en función de los hechos y circunstancias de la entidad. Los riesgos sobre fondo gris en la tabla 3 siguiente solo se revisarán cuando la IA Generativa tenga una intervención crítica en el proceso auditado.

¹⁷ Una de las fuentes para elaborar la tabla 3 ha sido la publicación “*Auditing in the Age of Generative AI*”, The Center for Audit Quality, abril de 2024.

¹⁸ Según la NIA-ES 315Revisada/GPF-OCEX 1315Revisada, un **Riesgo significativo** es un riesgo identificado de incorrección material para el que la valoración del riesgo inherente se encuentra próxima al límite superior del espectro de riesgo inherente debido al grado en el que los factores de riesgo inherente afectan a la combinación de la probabilidad de que exista una incorrección y a la magnitud de la incorrección potencial si existe. Una forma simplificada de priorizar los riesgos e identificar los que son significativos es valorarlos en una escala de Muy alto, Alto, Bajo, Muy bajo o Nulo, según su impacto y probabilidad. O mediante cualquier otro sistema que la ICEX tenga establecido.

Tabla 3 Riesgos derivados del uso de la IA Generativa

Área de riesgo	Ejemplos de riesgo o factores de riesgo	Objetivos del auditor	Preguntas/consideraciones del auditor
Gobernanza	No existen normas de uso de la IA Generativa. Se desconocen las herramientas de IA Generativa utilizadas en la organización. Las soluciones de IA no se identifican y gestionan de forma adecuada y coherente en toda la entidad.	Validar si la gobernanza garantiza que el uso de la IA Generativa está alineado con los objetivos estratégicos de la entidad y si existe un canal formal de aprobación para nuevos casos de uso que afecten a la información financiera u otra relativa al ámbito auditado o al cumplimiento legal y los casos de uso actuales están aprobados	¿Tiene establecida la entidad una adecuada gobernanza de la IA Generativa? ¿Existe un comité de ética de IA u órgano equivalente que establezca y supervise los roles y responsabilidades? ¿Existe un inventario actualizado de todos los casos de uso de IA Generativa en la organización, incluyendo su finalidad, responsable funcional, responsable técnico, proveedor, criticidad y procesos afectados? ¿Contempla la identificación del uso de herramientas no autorizadas por los empleados (la “IA en la sombra” o “<i>shadow AI</i>”)? ¿Ha desarrollado la entidad un marco para el uso responsable de IA Generativa, incluyendo políticas con respecto a su uso aceptable y ético y su alineación con el reglamento europeo sobre IA? ¿Cómo se documentan y comunican las políticas relativas al uso aceptable y ético de IA Generativa a las personas adecuadas en toda la entidad? ¿Cómo supervisa la entidad el cumplimiento de las políticas relativas al uso aceptable y ético de IA Generativa? ¿Existe un comité u órgano equivalente que evalúe estos aspectos periódicamente? ¿Quién (individuo o grupo) en la entidad es responsable de la supervisión del uso de IA Generativa? ¿Posee la competencia técnica necesaria para entender sus riesgos específicos? ¿Tiene la entidad un proceso para seguir y monitorizar el uso de IA Generativa en toda la entidad, incluido el uso por parte de proveedores de servicios? ¿Cómo evalúa la entidad el impacto (naturaleza y grupos afectados) de las tecnologías de IA Generativa que se están implementando? ¿Cómo hace la entidad el seguimiento de los riesgos derivados del uso de las tecnologías de IA Generativa y de las medidas que los mitigan? ¿Tiene un sistema de gestión de riesgos? ¿Existen planes de formación para los empleados usuarios de la IA?

Área de riesgo	Ejemplos de riesgo o factores de riesgo	Objetivos del auditor	Preguntas/consideraciones del auditor
Fiabilidad de los datos	Baja fiabilidad (exactitud y completitud) de los datos utilizados. Errores de lectura de los documentos que se incorporan en el proceso. En procesos financieros, es importante confirmar que la IA no está “alucinando” o alterando cifras basadas en datos históricos.	Verificar que la preparación de los datos, por ejemplo, no haya introducido sesgos involuntarios o errores.	¿Se verifican los orígenes de los datos para garantizar su fiabilidad (exactitud y completitud)? ¿Qué datos, fuentes documentales o bases de conocimiento utiliza la IA? ¿Cómo se garantiza su calidad, integridad, actualización, pertinencia y trazabilidad? ¿Existen controles para detectar errores en los resultados, la lectura o interpretación de documentos? ¿Se han implementado controles para detectar sesgos o discriminación en los datos de entrada o resultados? ¿Se realizan verificaciones posteriores para asegurar la fiabilidad de los datos procesados? ¿Se verifica que los datos recuperados mediante RAG u otras fuentes externas son pertinentes, actualizados, completos y proceden de fuentes autorizadas o verificables?
Cumplimiento normativo	El uso por la entidad de la IA Generativa infringe leyes o reglamentos, en particular el RIA, el ENS o el RGPD. No se cumple con la normativa interna. No se tienen en cuenta los cambios establecidos en las normas y políticas El contrato con proveedores de IA no contempla todos los aspectos de la LCSP	Identificar el cumplimiento de legalidad en el uso de la IA Generativa.	¿Cuál es la normativa aplicable, además del RIA, que afecta al uso de la IA Generativa por parte de la entidad? ¿Ha identificado la entidad adecuadamente dichas obligaciones normativas? ¿Se ha tenido en cuenta el cumplimiento de la legalidad para implantar el sistema de IA? ¿RIA, RGPD, ENS? ¿Ha clasificado la entidad el sistema según sus niveles de riesgo de acuerdo con el RIA (alto riesgo, riesgo limitado o propósito general)? ¿Existe un contrato según indica la LCSP que regule el uso de los datos de las herramientas de IA Generativa, la localización de estos dentro de la UE y la gestión y eliminación de la información? Si la entidad utiliza tecnologías de IA Generativa desarrolladas por un tercero, ¿puede obtener información suficiente del tercero con respecto al cumplimiento de la normativa?
Fraude	Las tecnologías de IA son utilizadas por los empleados, la dirección o terceros para perpetrar, facilitar y ocultar el fraude. Los riesgos incluyen, entre otros: • Suplantación de identidad (de directivos o proveedores). • Generación de documentación falsa con apariencia legítima (facturas, contratos...) • Manipulación de registros o datos mediante la automatización de procesos.	Evaluar si el marco de control interno es capaz de mitigar riesgos de manipulación de la IA para engañar activamente a los supervisores, dando respuestas incorrectas o con sesgo de forma intencionada.	¿Cómo ha considerado la entidad el impacto de las tecnologías de IA en su evaluación del riesgo de fraude? ¿con qué periodicidad realiza dicha evaluación? ¿Ha identificado la entidad nuevos incentivos, oportunidades o presiones para cometer fraude debido al uso de tecnologías de IA Generativa? ¿Existen procedimientos para verificar la autenticidad de los inputs o evitar que la IA sea usada para manipular o sobrescribir controles automatizados? ¿Qué mecanismos de supervisión humana existen para detectar si el modelo está siendo objeto de acciones destinadas a manipular los resultados?

Área de riesgo	Ejemplos de riesgo o factores de riesgo	Objetivos del auditor	Preguntas/consideraciones del auditor
Alucinaciones de los modelos	El sistema produce información aparentemente coherente y plausible pero incorrecta o inventada. <i>Este fenómeno cobra especial relevancia en el contexto de la auditoría cuando los sistemas de IA se emplean para generar contenido que pueda influir en la toma de decisiones o en la presentación de información financiera.</i>	Evaluar la eficacia de los controles preventivos y detectivos para asegurar que los resultados del modelo son fiables, evitando que datos inventados o incorrectos contaminen la información financiera o los procesos de negocio.	¿Se ha implantado algún control para detectar alucinaciones? ¿qué nivel de confianza otorga la entidad a las salidas generadas por la IA? ¿Se exige que las respuestas generadas por IA incluyan referencias a las fuentes utilizadas o evidencias que permitan verificar su contenido? ¿Existe una supervisión humana (<i>human in the loop</i>) que valide la integridad de la información producida antes de su uso operativo? ¿Hay un procedimiento escrito de cómo se realiza la revisión, los criterios de aceptación y registro de discrepancias? ¿Existe formación adecuada para los empleados sobre riesgos de que la IA genere alucinaciones?
Dependencia de proveedores externos	La utilización de soluciones de IA Generativa frecuentemente implica una dependencia significativa de proveedores tecnológicos externos, particularmente cuando los modelos se ejecutan en infraestructuras de nube: interrupción del servicio, cambios no notificados en los modelos (actualizaciones de algoritmos que alteran los resultados), falta de transparencia (sobre calidad del entrenamiento).	Evaluar si la entidad ha mitigado el riesgo de “quedar cautiva” de un proveedor y si existen controles suficientes para supervisar que el proveedor externo mantenga la integridad y seguridad de los modelos de IA que facilita.	¿Cómo evalúa la entidad la capacidad financiera y operativa de sus proveedores críticos de IA para garantizar su continuidad? ¿Incluyen los contratos cláusulas de notificación de cambios significativos en los algoritmos o datos de entrenamiento que puedan afectar a la precisión del sistema? ¿Utiliza la entidad una instancia pública o privada o híbrida de IA? Especifique. ¿Dispone la organización de un plan de contingencia en caso de rescisión del contrato o fallo prolongado del proveedor externo? ¿Se han establecido cláusulas de reversibilidad o plan de salida que garanticen la recuperación de datos, configuraciones, <i>prompts</i>, <i>logs</i>, <i>embeddings</i>, documentación y resultados generados al finalizar el contrato? ¿Se han definido contractualmente niveles de servicio (SLA) que midan no solo la disponibilidad, sino también la precisión y rendimiento operativo de la IA suministrada? ¿El contrato con la empresa proveedora de la IA, si es el caso, cubre cláusulas sobre mantenimiento, escalabilidad, cumplimiento normativo y propiedad intelectual de los outputs? ¿El contrato excluye o aclara si los datos de la entidad se utilizan para entrenar sus modelos comerciales (re-entrenamiento)? Se deben revisar los contratos con los proveedores de IA y de <i>cloud</i>. Se debe aplicar la NIA-ES-SP 1402 y la GPF-OCEX 1403.

Área de riesgo	Ejemplos de riesgo o factores de riesgo	Objetivos del auditor	Preguntas/consideraciones del auditor
Falta de control humano	Existe una dependencia excesiva de las salidas del modelo sin validación, incurriendo en un sesgo de automatización: <i>La automatización mediante IA Generativa sin la adecuada supervisión humana puede generar situaciones en las que decisiones con implicaciones significativas se tomen sin la debida revisión, autorización o validación por parte de personal cualificado.</i> La usencia de puntos de control que permitan una intervención humana efectiva en procesos críticos constituye una debilidad significativa en el control interno.	Constatar que la entidad ha implementado medidas de mitigación que aseguren que la IA no opera de forma aislada, garantizando que la responsabilidad final de las decisiones recaerá siempre en una persona identificable.	¿Ha definido claramente la entidad qué procesos y decisiones requieren intervención humana obligatoria, especialmente aquellos que involucren estimaciones contables significativas, reconocimiento de ingresos, deterioro de activos, concesión de subvenciones o cualquier otro aspecto que requiera juicio profesional? ¿Se han designado las personas con capacitación necesaria para ejercer la supervisión humana sobre los resultados del proceso? ¿El personal mantiene su capacidad de juicio crítico frente a las sugerencias de la IA para evitar decisiones automatizadas erróneas? ¿Existe supervisión humana en los resultados antes de que impacten en estados financieros o en decisiones críticas? ¿Existen controles para detectar errores en la lectura o interpretación de documentos?
Ciberseguridad	La tecnología IA Generativa es susceptible a ciberataques, incluida la intoxicación de datos, los <i>prompts injections</i> maliciosos, la manipulación malintencionada de los <i>prompts</i>, o la exfiltración de información sensible a través de técnicas de ingeniería de <i>prompts</i>.	Confirmar que la entidad tiene establecidos controles y medidas que garantizan una adecuada ciberseguridad, incluyendo el control de qué herramientas se usan y con qué información se “alimenta” al modelo.	¿Ha realizado la entidad una gestión de riesgos de ciberseguridad para evaluar amenazas a la IA Generativa y las salvaguardas necesarias? ¿Cómo considera la entidad los riesgos de ciberseguridad a la hora de seleccionar o desarrollar tecnologías de IA Generativa? ¿Se ha evaluado el riesgo de <i>prompt injection</i> (manipulación del modelo mediante instrucciones maliciosas), sesgos en las respuestas o la ausencia de salvaguardas implantadas en la herramienta de IA Generativa? ¿Contemplan los controles de seguridad de la información de la entidad las ciberamenazas a la IA? ¿Los contratos con proveedores de IA incluyen compromisos de seguridad verificables (certificaciones ISO/IEC 27001 e ISO 27701, SOC2, entre otras)? ¿Tienen una cláusula que garantice el derecho de auditoría? ¿Se realizan evaluaciones de vulnerabilidades específicas para sistemas de IA? ¿Existen mecanismos de detección de intentos de manipulación o uso indebido? ¿Existe formación adecuada para los empleados sobre riesgos de compartir información confidencial con la IA? ¿Se han establecido mecanismos de respaldo y planes de contingencia, para minimizar el impacto en caso de fallos de integración o indisponibilidad del modelo?

Área de riesgo	Ejemplos de riesgo o factores de riesgo	Objetivos del auditor	Preguntas/consideraciones del auditor
Protección de información confidencial / Privacidad de datos	Fuga de información confidencial o de propiedad intelectual de la organización, o exposición accidental de datos de carácter personal o sensibles: <i>El uso de sistemas de IA Generativa, particularmente aquellos basados en servicios externos o modelos en la nube, conlleva el riesgo de que información confidencial, sensible o estratégica de la entidad sea inadvertidamente expuesta, extraída o almacenada por terceros.</i> Riesgo de reutilización de datos (reentrenamiento): <i>Cuando los usuarios interactúan con estos sistemas introduciendo datos en sus prompts, pueden estar compartiendo información reservada, con proveedores cuyos términos de uso pueden permitir el almacenamiento, procesamiento o incluso la utilización de dicha información para reentrenar sus modelos.</i> Es un riesgo alto que se estén introduciendo en estos sistemas datos financieros no públicos, datos personales protegidos por normativa de privacidad, o cualquier información cuya divulgación pudiera constituir un incumplimiento de obligaciones contractuales, legales o regulatorias. <i>Estos aspectos deberán estar regulados por los contratos.</i>	Verificar si la entidad ha implementado barreras técnicas y organizativas para evitar que el flujo de datos hacia/desde la IA comprometa el cumplimiento legal de la privacidad de terceros, la naturaleza reservada de la información confidencial y la propiedad intelectual.	¿Ha clasificado la entidad los tipos de información que pueden ser tratados por sistemas de IA generativa, tales como pública, interna, confidencial, reservada, datos personales, categorías especiales de datos o información protegida por secreto? ¿Se han establecido políticas claras sobre qué tipo de información puede o no introducirse en sistemas de IA Generativa? ¿Existe una política de uso de la IA Generativa que prohíba explícitamente al personal de la entidad introducir información confidencial, secretos comerciales, o datos protegidos o información reservada y datos personales (salud, financieros), en herramientas de uso público? ¿Se han implementado controles técnicos para prevenir la filtración de datos sensibles (como sistemas de clasificación automática de información o barreras de prevención de pérdida de datos)? ¿Cómo se asegura la entidad de que la información introducida no sea utilizada para el entrenamiento de los modelos del proveedor? ¿Los contratos con los proveedores de la IA Generativa/cloud, contemplan expresamente estos supuestos? ¿Existe formación adecuada para los empleados sobre los riesgos de compartir información confidencial con estos sistemas? ¿Cómo garantiza la entidad la soberanía de sus datos (alojamiento de su procesamiento) cuando se utilizan servicios de IA Generativa en la nube? ¿Cómo considera la entidad los riesgos de privacidad de datos al seleccionar o desarrollar tecnologías de IA Generativa? ¿Utiliza la entidad una instancia pública de tecnologías de IA Generativa que rastrea y guarda entradas y datos accesibles por terceros o una instancia privada donde las entradas y los datos son controlados y guardados solo por la entidad?

Área de riesgo	Ejemplos de riesgo o factores de riesgo	Objetivos del auditor	Preguntas/consideraciones del auditor
Segregación de funciones (SdF) incompatibles	La automatización mediante IA puede eliminar o difuminar las segregaciones tradicionales de funciones, concentrando en un único sistema automatizado actividades que tradicionalmente requerían intervención de diferentes personas o departamentos como mecanismo de control. Por ejemplo, un sistema de IA podría simultáneamente aprobar transacciones, registrarlas contablemente y generar reportes, eliminando los controles preventivos basados en la separación de responsabilidades. Riesgo de que el personal con capacidad de modificar el sistema o los <i>prompts</i> de configuración sea el mismo que supervisa el cumplimiento normativo y ético.	Evaluar cómo la entidad ha rediseñado sus controles internos para compensar la posible concentración de funciones en el sistema de IA, incluyendo controles de supervisión, revisiones o restricciones programadas en los propios sistemas.	¿Cómo ha establecido la entidad la SdF con el nuevo sistemas IA? ¿Se produce algún conflicto no resuelto? Además de la SdF en el proceso funcional, el auditor también revisará la SdF en los procesos TI: ¿Están claramente definidos y documentados los roles de desarrollo, operación y supervisión del sistema, asegurando que sean ejercidos por personas o unidades distintas? ¿El proceso de gestión de cambios definido por la entidad garantiza que las modificaciones en instrucciones maestras sean aprobadas por personal ajena a su creación? ¿Existe independencia entre quienes gestionan los datos de entrenamiento/afinación y quienes realizan las pruebas de validación?
Falta de trazabilidad y pistas de auditoría	Efecto caja negra. Ausencia de registros que vinculen un resultado (output) de la IA con el <i>prompt</i> original, el usuario que lo ejecutó y la versión específica del modelo utilizado. La naturaleza dinámica y automatizada de los sistemas de IA Generativa puede dificultar el mantenimiento de una pista de auditoría completa y coherente que documente adecuadamente cómo se generaron determinadas salidas, qué datos se utilizaron, qué versión del modelo intervino, y quién autorizó o validó los resultados. Esta carencia compromete la capacidad del auditor para verificar transacciones, rastrear el origen de cifras en los estados financieros y obtener evidencia suficiente y adecuada.	Evaluar si los sistemas mantienen registros (<i>logs</i>) completos de las operaciones realizadas por la IA, si estos registros son inmutables y están protegidos contra alteraciones, y si permiten reconstruir el proceso de toma de decisiones o generación de información financiera cuando sea necesario. Es decir, verificar que todo output del modelo tenga trazabilidad y pueda ser explicada y atribuida a un contexto y usuario específico.	¿El sistema genera registros (<i>logs</i>) de auditoría detallados que capturen la entrada del usuario, la configuración del modelo y la respuesta generada? ¿Se registran y conservan los <i>prompts</i>, respuestas, fuentes consultadas, parámetros relevantes, versiones del modelo, usuarios intervinientes y acciones ejecutadas por la herramienta o agente? ¿Se dispone de mecanismos de interpretabilidad y explicabilidad que permitan entender los factores clave que incluyeron en una decisión u output de alto impacto? ¿El registro de actividad está protegido contra modificaciones no autorizadas, garantizando la integridad de la evidencia para futuras revisiones?

Área de riesgo	Ejemplos de riesgo o factores de riesgo	Objetivos del auditor	Preguntas/consideraciones del auditor
Uso de un modelo fundacional	El modelo de IA no es confiable o seguro para el contexto específico de la entidad, lo que resulta en errores repetidos o favorecedores de ciertos resultados o salidas por parte del modelo El modelo presenta sesgos inherentes de su entrenamiento original que dan lugar a resultados inapropiados o erróneos para el proceso auditado. Desconocimiento de las limitaciones técnicas y de actualizaciones de datos del modelo fundacional.	Verificar que la elección del modelo no ha sido arbitraria y que existe una justificación técnica o de negocio que avale la elección. Que la entidad ha consultado la documentación técnica publicada para justificar que el modelo es apto para el propósito de negocio sin comprometer la integridad de los datos.	¿Cómo considera la entidad si el modelo fundacional se adapta adecuadamente a sus necesidades para el caso de uso en cuestión? ¿Cómo ha evaluado si las capacidades del modelo fundacional son adecuadas y proporcionales a los riesgos del proceso donde se aplica? ¿Se ha realizado un estudio sobre la adecuación de la tecnología IA en lugar de otras tecnologías alternativas? ¿Cómo identifica y selecciona la entidad los procesos y, en su caso, proveedores adecuados para aplicar IA Generativa? ¿Existe un proceso formal para decidir si se requiere un ajuste o si basta con un modelo generalista para la tarea encomendada? ¿Cómo determina la entidad si debe añadir personalizaciones al modelo fundacional (ajuste o <i>fine tuning</i>) para satisfacer sus necesidades específicas? ¿Se han definido métricas y umbrales mínimos de aceptación del modelo, como exactitud, tasa de error, coherencia, tasa de alucinaciones, sesgos, tiempos de respuesta o necesidad de corrección humana? ¿Existe un análisis de la procedencia y transformación/limpieza de los datos con los que el modelo fue entrenado para identificar posibles riesgos legales o éticos heredados?
Confiabilidad y monitoreo continuos	La solución de IA Generativa no se monitoriza después de su implementación para determinar si está funcionando adecuadamente, de acuerdo con lo previsto. Riesgo de degradación del modelo: Después del despliegue, el sistema deja de funcionar según su diseño original debido a la evolución de las tecnologías o a cambios intencionados o no intencionados en los patrones de los datos de entrada o por actualizaciones en el modelo fundacional.	Verificar la existencia de controles de supervisión permanentes que mitiguen el riesgo de que la IA genere resultados erróneos por el simple paso del tiempo o por cambios en el entorno operativo.	¿Cómo monitoriza la entidad la efectividad y precisión de los resultados para asegurar que siguen respondiendo a la finalidad prevista? ¿Se ha implementado un sistema de registro (<i>logs</i>) de las actividades con IA y se audita periódicamente para asegurar su integridad y precisión? ¿Se han implantados controles automáticos que identifiquen y notifiquen sobre posibles errores o inconsistencia en los datos? ¿Existe un procedimiento para suspender, bloquear o limitar el uso de la IA cuando se detecten errores relevantes, comportamientos no autorizados o riesgos no aceptables? ¿Se han instaurado controles que permitan monitorizar en la IA y alertar sobre posibles disfuncionalidades o desviaciones, degradaciones de rendimiento o derivas del modelo? ¿Se realizan verificaciones posteriores para asegurar la fiabilidad del procesamiento? ¿Tiene la entidad un proceso para reevaluar periódicamente las tecnologías de IA Generativa para determinar si están funcionando según lo previsto? ¿Con qué periodicidad se reevalúa la calidad de las respuestas? ¿Cómo monitorea la entidad los cambios en las tecnologías de IA Generativa? ¿Cómo gestiona la entidad las actualizaciones de versión en modelos de terceros y el impacto de la infraestructura en la nube en la disponibilidad del servicio?

Área de riesgo	Ejemplos de riesgo o factores de riesgo	Objetivos del auditor	Preguntas/consideraciones del auditor
Sesgos del modelo	Los modelos de IA pueden incorporar y perpetuar sesgos derivados tanto de los datos históricos utilizados en su entrenamiento como de las decisiones de diseño tomadas durante su desarrollo. Estos sesgos pueden manifestarse de forma sistemática en las decisiones, clasificaciones o predicciones realizadas por la IA Generativa, afectando potencialmente a procesos como la evaluación de solicitudes de subvenciones o ayudas, la estimación de provisiones o la valoración de activos.	Verificar que la entidad ha implementado controles para identificar y, en su caso, corregir sesgos que pueda derivar en decisiones discriminatorias, riesgos por incumplimientos legales o daños reputacionales.	¿Cómo realiza la entidad evaluaciones de sesgo en sus modelos de IA? ¿Cuenta con herramientas para detectar y medir sesgos en las repuestas generadas por la IA? ¿Existen controles para mitigar el impacto de estos sesgos en la información financiera, particularmente cuando dichos sistemas intervienen en procesos de estimación contable o juicio profesional? ¿Y en otros procesos, como por ejemplo en la concesión de subvenciones o en la valoración de contratos? ¿Cómo monitoriza la entidad las tecnologías de IA Generativa a lo largo del tiempo para determinar si se ha introducido sesgo a través de los algoritmos o los datos?
Conocimientos y habilidades personales	Las personas que ocupan puestos de gobierno o de dirección no tienen los conocimientos y las capacidades adecuados para supervisar eficazmente el enfoque de la entidad para el despliegue de IA Generativa. La entidad no tiene personal cualificado para supervisar, desarrollar, implementar, operar y monitorizar con éxito las tecnologías de IA Generativa, lo que genera una confianza y dependencia excesiva en los resultados (sesgo de automatización). Riesgo de falta de capacitación en conceptos fundamentales: <i>Los usuarios no conocen qué es una alucinación del sistema o cómo funciona el prompting, aumentando la probabilidad de errores operativos. Puede deberse a que la entidad no proporciona suficiente formación a los empleados para utilizar las tecnologías de IA Generativa de manera efectiva.</i>	Evaluar si el nivel de capacitación y conciencia del personal es suficiente para compensar la opacidad técnica de la IA, asegurando que los empleados actúen como un control crítico y no como ejecutores pasivos de las sugerencias del modelo.	¿Ha identificado la entidad habilidades o conocimientos especializados necesarios para ejercer la supervisión, el desarrollo, la implementación, la operación y el seguimiento de la IA Generativa? ¿Cómo proporciona la entidad formación a la dirección y a los empleados que son responsables de la supervisión, el desarrollo, la implementación, la operación o el seguimiento de la IA Generativa? ¿Cómo forma la entidad a los empleados y a la dirección sobre el uso responsable de la IA Generativa, incluida una comprensión de los riesgos asociados (alucinaciones, sesgos, falta de exactitud...) y de las salvaguardas de la IA sobre la capacidad de confiar en los resultados? ¿La entidad proporciona recursos a los empleados, como manuales de usuario y soporte en tiempo real, para el uso de la IA Generativa? ¿Comprende el personal el impacto de la infraestructura en la nube y cómo esta afecta al manejo de la información corporativa sensible? ¿Se han identificado perfiles que requieran habilidades en ingeniería del <i>prompt</i> o revisión de contenidos para asegurar la supervisión humana? ¿La entidad contrata personal o recursos de terceros con la experiencia necesaria para ayudar a garantizar la supervisión, el desarrollo, la implementación, la operación y el seguimiento exitoso de la IA Generativa?

9. Identificación y revisión de los controles de procesamiento de la información. Valoración del riesgo de control.

Durante la etapa de conocimiento del proceso de gestión que se está auditando y de su sistema de control interno, tras identificar los riesgos, se realizará la **identificación, análisis y evaluación del diseño e implementación** de los controles de procesamiento de la información (CPI) que la entidad ha implementado para hacer frente a los riesgos significativos y así prevenir, detectar y corregir incorrecciones provocadas por la utilización de la IA Generativa, para ello:

- El auditor ejecutará procedimientos de valoración del riesgo para identificar los CPI implantados que hagan frente a los riesgos significativos, tanto manuales como automáticos. Entre ellos, controles de aprobación y validación humana en procesos con IA Generativa, políticas de segregación de funciones, etc.
- Analizará el adecuado diseño e implementación de los CPI relevantes.
- Identificará si existen controles compensatorios cuando se requieran.
- Verificará la trazabilidad de las operaciones dentro del sistema.
- Valorará el riesgo de control.

La **GPF-OCEX 5340** detalla la metodología para revisar los CPI en un entorno de administración electrónica. Según dicha guía, los CPI son controles relacionados con el procesamiento de la información en aplicaciones de TI o procesamiento manual de la información en el sistema de información de la entidad que responden directamente a los riesgos para la integridad de la información (es decir, la completitud, exactitud y validez de las transacciones y otra información) y el cumplimiento de la legalidad. Actúan en diferentes momentos del flujo de procesamiento: entrada, transformación, procesamiento, salida, datos maestros, interfaces, etc.

Serán CPI relevantes aquellos que hacen frente a los riesgos inherentes valorados como significativos (los que están en la parte alta del espectro de riesgo inherente).

A continuación, se indican algunos ejemplos de procedimientos de valoración del riesgo que se pueden ejecutar en esta fase de la auditoría:

- Revisar el apartado 7.4 y la tabla 3 para ayudar a identificar riesgos, adaptado todo ello a las circunstancias de la fiscalización.
- Verificar que se han establecido controles a lo largo de todo el ciclo de vida de los sistemas de IA Generativa, garantizando su efectividad en todas las fases (alta, desarrollo o entrenamiento, ajuste o *fine-tuning*, actualización, gestión de cambios, explotación y baja).

En particular, se deberá comprobar que existen procedimientos formales para la aprobación y registro de los sistemas de IA Generativa, la validación de los datos utilizados para el entrenamiento o el contexto, la gestión controlada de versiones de modelos y algoritmos, así como mecanismos de monitorización continua para detectar, en su caso, la degradación en producción. Asimismo, se deberá asegurar la existencia de registros y evidencias que permitan la trazabilidad de las decisiones de la IA, la identificación de posibles incidencias, alucinaciones o sesgos en los resultados.

- Comprobar si existen controles que incluyan la función de revisar el trabajo realizado por los gestores con la IA, supervisando, entre otros aspectos, la correcta aplicación de la normativa y procedimientos de uso ético y responsable.
- Comprobar si la entidad ha realizado un documento formal con un análisis de los riesgos que pueden afectar al proceso, evaluando y valorándolos, así como un plan de medidas para responder a dichos riesgos de forma que se puedan evitar o minimizar.
- Verificar si hay establecidos procedimientos para la supervisión de las tareas que se realizan.
- Comprobar si se deja evidencia de la supervisión realizada.

Una vez identificados los CPI relevantes, se deberán realizar las pruebas de controles para comprobar su eficacia operativa.

10. Identificación y revisión de los controles generales de tecnologías de la información

Los controles generales de tecnologías de la información (CGTI) son los controles de los procesos de TI de la entidad que apoyan el funcionamiento continuo apropiado del entorno de TI, incluido el funcionamiento continuo efectivo de los CPI y la integridad de la información (es decir, la completitud, exactitud y validez de la información) en el sistema de información de la entidad.

Como paso previo será necesario recopilar información sobre los recursos tecnológicos existentes en la entidad, como hardware, software, redes y sistemas de comunicación, que afectan al entorno automatizado donde opera la IA Generativa que se está auditando. Esto incluye identificar si se trata de modelos de caja negra (opacos), así como la infraestructura de datos asociada.

El auditor deberá identificar los **CGTI** del proceso de gestión en el que se incorpora la IA que apoyan el funcionamiento efectivo de los CPI relevantes sobre la IA en los que se va a confiar, evaluar su **adecuado diseño e implementación y verificar su eficacia operativa**.

Esta revisión se efectuará de acuerdo con la metodología recogida en la GPF-OCEx 5330 y en los manuales de fiscalización del TCu, seleccionando los CGTI más relevantes de acuerdo con los objetivos de la auditoría y los riesgos valorados.

Relacionados con la IA Generativa, en el caso de las guías de los OCEx, se debe prestar especial atención a:

- **GPF-OCEx 5332 Gestión de cambios:** en particular los relacionados con el diseño, configuración, parametrización y adaptaciones de los modelos y del sistema de IA Generativa y la posterior implementación de cambios, así como comprobar si se prueban adecuadamente y aprueban los cambios antes de pasarlos a producción, lo que permitirá evitar errores o resultados inesperados.
- **GPF-OCEx 5333 Operaciones de los sistemas de información:** habrá que determinar si se monitoriza la ejecución de la IA Generativa, se gestionan errores, excepciones o fallos. Si se utiliza un proveedor externo, además de revisar el CGTI "C5 Servicios externos" se deberá revisar el contrato y si se presta en alguna modalidad "cloud".
- **GPF-OCEx 5334 Controles de acceso:** determinar quién tiene acceso para definir, modificar, ejecutar y supervisar el sistema de IA Generativa, si está el acceso basado en el principio de necesidad de conocer y debidamente segregado. Son críticos los controles sobre usuarios con privilegios de administración y sobre quiénes tiene capacidad de modificar los datos de entrenamiento del modelo.
- **GPF-OCEx 5335 Continuidad del servicio:** en el caso de que se trate de un proceso crítico, habrá que tener en cuenta si se realizan copias de seguridad y existe un plan de continuidad aprobado, incluyendo el procedimiento de intervención o parada de emergencia si el sistema de IA Generativa presenta comportamientos imprevistos.

Según las circunstancias de la auditoría y la valoración de riesgos realizada, de entre todos los CGTI, se seleccionarán los que apoyen el correcto funcionamiento de los CPI relevantes, de los que se va a comprobar su eficacia operativa, y solo de estos CGTI se evaluará su adecuado diseño, implementación y eficacia operativa. No obstante, se podrá establecer un objetivo y alcance de la auditoría más amplio.

11. Procedimientos posteriores de auditoría

Los procedimientos posteriores de auditoría, de acuerdo con la NIA-ES-SP 1330, incluyen las **pruebas de controles** y los **procedimientos sustantivos** (procedimientos analíticos y pruebas en detalle). Tras identificar los controles relevantes (CPI y CGTI) y evaluar su adecuado diseño e implementación se debe revisar su **eficacia operativa** mediante pruebas de controles para asegurarse de que los controles funcionan como se espera.

Al diseñar los procedimientos posteriores de auditoría el auditor debe tener en cuenta, para cada afirmación relevante de los tipos de transacciones, saldos contables e información a revelar que resulte afectada por la IA, lo siguiente:

- a) La valoración del riesgo inherente (RI) relacionado con el uso de la IA y los factores de riesgo que motivan esa valoración, incluyendo, en su caso, la calidad de los datos de entrenamiento.

- b) La valoración del riesgo de control realizada sobre la herramienta de IA implantada en el proceso auditado.
- c) Se obtendrá más evidencia de auditoría persuasiva/convincente cuanto mayor sea la valoración del RI del sistema de IA Generativa en el espectro del riesgo inherente.

Esto se logrará aumentando la cantidad de evidencia obtenida o bien modificando la naturaleza, momento de realización o extensión de los procedimientos posteriores de auditoría realizados para hacer frente al riesgo identificado.

Se obtendrá evidencia más persuasiva si se integra en el equipo de auditoría personal especializado con conocimientos TI y específicamente en IA (ingeniería de IA, ciencia de datos e ingeniería de datos, ética de IA y seguridad). En caso contrario, sin especialistas en el equipo de auditoría, la evidencia, en general, será menos convincente.

Cuando el riesgo esté en la parte baja del espectro de riesgo inherente se requerirá menos cantidad de evidencia.

- d) Se diseñarán pruebas de controles para probar la eficacia operativa de aquellos CPI considerados relevantes y los CGTI que los apoyan y/o soportan el buen funcionamiento del sistema.
- e) Se diseñarán procedimientos sustantivos necesarios para evaluar la fiabilidad de la información producida por la IA Generativa.
- f) Dadas las características de la información / datos producidos por la IA Generativa (ver apartado 12 siguiente) puede ser necesario diseñar procedimientos de auditoría adicionales para corroborar la evidencia obtenida a partir de la IA.

Se recomienda utilizar como guías del trabajo de auditoría la **GPF-OCEx 5340** y la **GPF-OCEx 5330**, y los manuales de fiscalización de regularidad y de gestión en el caso del TCu, aplicando sus orientaciones para diseñar y ejecutar pruebas de controles sobre los CPI y CGTI relevantes.

Algunos de los procedimientos posteriores de auditoría a ejecutar pueden ser los siguientes:

- Comprobar la frecuencia y alcance de las revisiones realizadas por la persona o unidad responsable de la operación, mantenimiento y supervisión del sistema de IA Generativa.

Se deberá comprobar que existen revisiones periódicas para verificar el correcto funcionamiento del modelo, detectar incidencias o desviaciones en su rendimiento y asegurar que los resultados siguen siendo adecuados para su finalidad. Asimismo, se revisará que estas evaluaciones quedan documentadas y que, en caso de detectarse problemas, se adoptan las medidas correctoras correspondientes.

Revisar los registros de reuniones de control o seguimiento del sistema y la actividad del comité de ética de IA u órgano equivalente, si existiera en la entidad.

- Comprobar que el mapa de procesos es correcto.
- Verificar la correcta configuración de reglas y validaciones establecidas dentro del sistema.
- Realizar pruebas de entrada y salida en el sistema para validar lógica automática, preferiblemente en un entorno de pruebas, utilizando un conjunto de datos de referencia (*ground truth dataset*) para comparar los resultados.
- Diseñar pruebas sustantivas para verificar que los resultados obtenidos son los esperados.
- Solicitar la base de datos de producción de los resultados generados por la IA para analizar la validez e integridad de su información.
- Cuando lo exija la normativa aplicable, se comprobará que los elementos técnicos necesarios para comprender y auditar el funcionamiento del sistema de IA (modelos entrenados, configuraciones, documentación técnica y, en su caso, información sobre los datos de entrenamiento) están accesibles para su revisión y cumplen con los requisitos establecidos en materia de transparencia, trazabilidad y control. Pedir las evidencias de las verificaciones realizadas por el ente fiscalizado respecto a los controles, monitoreo del desempeño y disponibilidad del sistema y revisarlas.

- Verificar la existencia de procedimientos de respaldo y recuperación del sistema IA y de los datos.
- Comprobar si se han realizado actividades formativas al personal respecto al uso de la IA, tanto de conceptualización y sensibilización en los riesgos de su uso, como en conocimiento de “*ingeniería del prompt*” y reconocimiento de alucinaciones y sesgos de la IA.

12. Evidencia de auditoría obtenida en entornos de IA Generativa

Para evaluar la evidencia de auditoría que haya sido producida por los sistemas IA del ente auditado, es necesario tener en cuenta los criterios establecidos en la **NIA-ES-SP 1500, manuales de fiscalización del TCu** y en la **GPF-OCEx 1503 La evidencia electrónica de auditoría**. En este apartado se destacan los principales aspectos allí tratados, pero se recomienda su estudio completo.

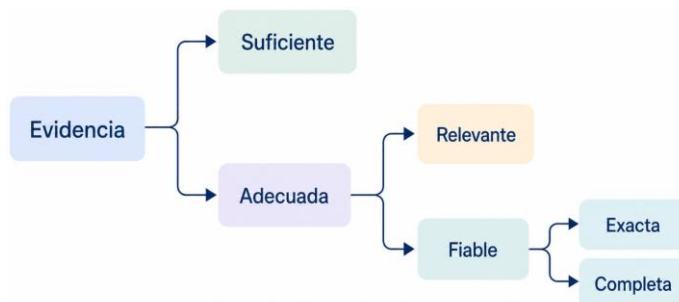
12.1 Características de la información resultante de la IA Generativa que se pretende utilizar como evidencia de auditoría

En la auditoría financiera, la información que se utilizará como evidencia de auditoría debe tener unas características fundamentales para que el auditor, tras la ejecución de sus procedimientos de auditoría, la considere evidencia de auditoría y pueda utilizarla como apoyo de sus conclusiones. De acuerdo con la NIA-ES-SP 1500:

El auditor diseñará y aplicará procedimientos de auditoría que sean adecuados, teniendo en cuenta las circunstancias, con el fin de obtener evidencia de auditoría suficiente y adecuada.

Al realizar el diseño y la aplicación de los procedimientos de auditoría, el auditor considerará la relevancia y fiabilidad de la información que se utilizará como evidencia de auditoría.

Es decir, la evidencia debe ser:



La **suficiencia es una medida cuantitativa** de la evidencia de auditoría: la cantidad de evidencia de auditoría necesaria depende de la valoración del auditor del riesgo de incorrección material, así como de la calidad de dicha evidencia de auditoría (NIA-ES-SP 1500, p5.e).

La **adecuación es una medida cualitativa** de la evidencia de auditoría, es decir, de su **relevancia y fiabilidad** para respaldar las conclusiones en las que se basa la opinión del auditor (NIA-ES-SP 1500, p5.b).

La **relevancia** se refiere a la conexión lógica con la finalidad del procedimiento de auditoría, o su pertinencia al respecto, y, en su caso, con la afirmación que se somete a comprobación.

Para que el auditor obtenga evidencia de auditoría **fiable**, la NIA-ES-SP 1500 (párrafo 9.a) exige a los auditores que obtengan evidencias sobre la **exactitud y completitud**¹⁹ de la información generada por la entidad que vaya a utilizarse como evidencia de auditoría. La fiabilidad de la evidencia está influenciada por su fuente y naturaleza y depende de las circunstancias en que se obtenga. Con la evolución de la tecnología, hay múltiples fuentes de información disponible y algunas son más confiables que otras.

En la IA Generativa, esta fiabilidad depende críticamente, entre otros factores, de la calidad y diversidad de los datos de entrenamiento, por lo que el auditor, siempre que sea posible, debe evaluar la “biografía de datos” (*data biography*) para comprender el origen y las transformaciones de dicha información.

¹⁹ *Completeness* en el original en inglés, *integridad* en la NIA-ES-SP 1500.

Una dificultad añadida, importante, al trabajar con información extraída de sistemas de IA Generativa se deriva de las características esenciales de esta tecnología que se han destacado en el apartado 5.1 de este documento. Es decir, a las características de **opacidad, no determinismo y adaptabilidad**, que, según los casos, pueden llegar a plantear problemas importantes sobre la fiabilidad de la evidencia obtenida, ya que debido al carácter probabilístico de la IA Generativa **no queda garantizada su exactitud y fiabilidad**.

Por lo tanto, **los resultados obtenidos con estas herramientas no representarán evidencia de auditoría suficiente y adecuada para soportar las conclusiones de auditoría si no existe otra evidencia corroborativa**.

Complementando lo anterior y utilizando la nueva metodología propuesta por la IAASB²⁰, los **atributos que el auditor puede considerar al evaluar la fiabilidad de la información resultante de la IA Generativa que interviene en el proceso de gestión, que se pretende utilizar como evidencia de auditoría** son:

Tabla 4 – Atributos de la fiabilidad de la información utilizada como evidencia de auditoría

Exactitud	La información está libre de errores en su reflejo de las condiciones, eventos, circunstancias, acciones u omisiones subyacentes, incluyendo el reflejo del período de tiempo o punto en el tiempo apropiado atribuible a las condiciones o eventos.
Complejidad	La información refleja todas las condiciones subyacentes, eventos, circunstancias, acciones o inacciones.
Autenticidad	La información ha sido generada o proporcionada por una fuente autorizada para hacerlo, y la información no ha sido alterada de manera que la haga poco fiable para la auditoría.
Sesgo	La información está libre de sesgos intencionados y no intencionados en su reflejo de las condiciones subyacentes, los eventos, las circunstancias, las acciones o las omisiones.
Credibilidad	La fuente de la información tiene la competencia y la capacidad para generar la información con el nivel de calidad requerido, y se puede confiar en ella.

Al diseñar y realizar los procedimientos de auditoría el auditor probará la eficacia operativa de aquellos controles que abordan la fiabilidad de la información producida por la IA Generativa o realizará otros procedimientos para evaluar la fiabilidad de dicha información.

12.2 Riesgo de exceso de confianza en la tecnología

El uso de la tecnología puede crear **sesgos** o un **riesgo general de confianza excesiva** en la información o en los resultados del procedimiento de auditoría realizado.

El exceso de confianza puede deberse a distintas causas, tales como asumir que los resultados del sistema de IA Generativa de una entidad son apropiados para su uso sin considerar que el modelo puede estar sufriendo sesgos o degradación de su rendimiento con el paso del tiempo (deriva de modelo o *model drift*).

La excesiva confianza en la tecnología por parte del auditor puede ser el resultado de una falta de escepticismo profesional o de juicio profesional.

Cuanto más complejo sea el entorno informático mayor será el grado de escepticismo profesional requerido para evaluar la evidencia electrónica.

13. Cuestiones relacionadas con el cumplimiento

El uso de IA Generativa en procesos de gestión está sujeto a un marco regulatorio en evolución que incluye normativa como ENS, RGPD o especialmente el Reglamento de IA (RIA).

Aunque esta no es una guía de auditoría sobre el cumplimiento normativo relacionado con la IA, en todos los trabajos de auditoría se debe incluir la revisión del cumplimiento legal de aquellos aspectos que puedan resultar significativos o relevantes en relación con el objetivo de la auditoría.

²⁰ IAASB Main Agenda (March 2026), Agenda Item 7, Audit Evidence & Risk Response (AE&RR).

Por esta razón, aunque la fiscalización de los aspectos legales relacionados con la IA será objeto de una guía específica, hay una serie de cuestiones fundamentales sobre el cumplimiento normativo que siempre deberán revisarse en cualquier auditoría financiera o de otro tipo en la que se revisen herramientas de IA Generativa. En concreto:

- Revisar la adecuación a los niveles de riesgo establecidos en el RIA: (inaceptable, alto, limitado o mínimo) Verificar cómo está clasificado el sistema de IA. En el uso de modelos de IA de uso general, se revisarán las obligaciones de transparencia y la gestión de modelos con riesgo sistémico.
- Verificar que el sistema de IA Generativa está cubierto por la certificación que corresponda del ENS adecuada a la categoría de seguridad (Nivel Básico, Medio o Alto) del proceso que soporta.
- Confirmar que la entidad identifica de forma explícita, cuando proceda, los contenidos generados o modificados por IA Generativa (imágenes, textos o audios) para evitar la posibilidad de engaño o desinformación.
- Verificar que la organización ha adoptado medidas para asegurar que el personal que opera con IA Generativa tiene los conocimientos técnicos necesarios para comprender sus riesgos y limitaciones.
- Revisar que se respetan los derechos de propiedad intelectual y que el entrenamiento de los modelos se haya realizado atendiendo a la normativa de protección de datos y garantía de la información sensible.
- Analizar el cumplimiento de los requisitos que se establecen en el Esquema Nacional de Seguridad, aplicando la parte correspondiente de la GPF-OCEX 5314 o la guía CCN-STIC 808 de verificación del cumplimiento del ENS y en el Reglamento General de Protección de Datos (RGPD), Reglamento (UE) 2016/679 del Parlamento Europeo y del Consejo, de 27 de abril de 2016 y la Ley Orgánica 3/2018, de 5 de diciembre, de Protección de Datos Personales y garantía de los derechos digitales.
- Verificar el cumplimiento de la LCSP al adquirir herramientas de IA Generativa.

14. Bibliografía

- **Achieving Effective Internal Control Over Generative AI.** Committee of Sponsoring Organizations of the Treadway Commission (COSO), COSO, 2026.
- **Agentic AI: What It is and Why Boards Should Care,** Protiviti, Board Perspectives Risk Oversight, Issue 186, 2025.
- **Aproximación a la Inteligencia Artificial y la ciberseguridad.** CCN-CERT, BP/30, octubre 2023
- **Artificial Intelligence Risk Management Framework (AI RMF 1.0).** National Institute of Standards and Technology, enero de 2023.
- **Desafíos para el control externo derivados del uso de la inteligencia artificial en el sector público,** Revista Española de Control Externo, María Dolores Genaro Moya y Antonio Manuel López Hernández, publicado en el número 74-75, mayo -septiembre de 2023.
- **Inteligencia artificial agéntica desde la perspectiva de protección de datos.** Agencia Española de Protección de Datos, febrero de 2026.
- International Auditing and Assurance Standards Board (IAASB)
 - **IAASB Technology Quality Management Roundtables-Briefing** Note for Participants, 20/8/2025.
 - **Component 3 Monitoring Update,** IAASB Main Agenda (September 2025), Agenda Item7-A
 - **Audit Evidence & Risk Response (AE&RR).** IAASB Main Agenda (March 2026), Agenda Item 7.
- **The Center for Audit Quality, Auditing in the Age of Generative AI,** abril de 2024.

Anexo 1. Ejemplo de documentación a solicitar durante la planificación

Tabla 5 – Ejemplo de documentación a solicitar

Área	Documento a solicitar	Finalidad
Gobernanza	Política o normas de uso de la IA	Garantizar el uso responsable y legal (se entiende que dicha política incluirá su alineación con el RIA y el marco ético)
	Actas de aprobación de la política de uso	Confirmar el respaldo de la dirección
	Registro de la formación a empleados en uso y riesgos de la IA	Mitigar el riesgo de errores operativos por falta de capacitación y verificar la “alfabetización” requerida por el art. 4 RIA
	Análisis de riesgos de uso de la IA	Identificar amenazas específicas de la IA Generativa
	Evaluaciones de impacto en el uso de la IA, si aplica	Documentar la identificación, análisis y mitigación de riesgos éticos, legales y técnicos antes de implementar sistemas de IA, especialmente los de alto riesgo
	Registro de usuarios y de usos de IA permitidos, incluyendo herramientas de IA no permitidas	Prevenir el uso de herramientas no autorizadas (IA en la sombra)
Diseño técnico	Arquitectura tecnológica de la integración de los sistemas de IA (modelo fundacional o instancia privada/pública)	Evaluar la soberanía y riesgos de seguridad de los datos (potenciales fugas de datos o filtraciones)
	Diseño de sistemas de vigilancia de la IA (medidas de monitorización de la degradación del modelo)	Detectar posible pérdida de precisión del modelo con el tiempo y si la entidad tiene medidas para evitarlo
	Integraciones de los sistemas de IA	Analizar el flujo de información
	Descripción del ciclo de vida de los sistemas de IA. Despliegues de los LLMs (incluyendo el ajuste o <i>fine-tuning</i>)	Justificar la adecuación del modelo al propósito
	Biografía de datos	Garantizar la fiabilidad, exactitud y completitud del origen de los datos
Seguridad y protección de datos	Documento de las medidas de seguridad del sistema implementadas	Proteger la integridad del sistema (verificar que incluyen controles contra <i>prompt injection</i>)
	Análisis de riesgos de seguridad	Evaluar vulnerabilidades tecnológicas del sistema
	Evaluación de impacto en protección de datos, si corresponde	Cumplir con el RGPD y privacidad
	Certificaciones de seguridad de terceros que suministran servicios (ENS, ISO 27001, ISO 27701, SOC2, etc.), si aplica	Verificar compromisos de seguridad de proveedores externos (validar la infraestructura y los procesos de gestión de la información del proveedor)
	Plan de contingencia y continuidad ante fallos de proveedores de servicios críticos	Mitigar el riesgo de dependencia externa

Área	Documento a solicitar	Finalidad
Operativa y funcionamiento (usuarios)	Registro de actividad de supervisión humana	Validar la integridad de las salidas antes de su uso
	Manual de administración técnica	Estandarizar la interacción segura con la IA (el manual proporcionado debería incluir guía o pautas de ingeniería del <i>prompt</i>).
Control y trazabilidad	Matriz de controles automáticos	Identificar controles internos programados (detección de alucinaciones y sesgos)
	Logs de trazabilidad	Verificar registro de actuaciones, accesos y decisiones (vinculación de <i>prompt</i> , usuario y versión del modelo), de forma que, en lo posible, mitigue el efecto caja negra (permita la reconstrucción de decisión/resultado)
	Informes de auditoría interna o controles previos	Conocer evaluaciones internas previas
	Indicadores de desempeño (KPIs de rendimiento y precisión)	Comprobar medición de la eficacia y calidad del sistema
Validación y calidad	Planes de pruebas (test plan)	Verificar que se hayan realizado pruebas técnicas y funcionales, que se han contrastado resultados frente a conjuntos de datos de referencia (<i>ground truth</i>)
	Actas de validación y aceptación del sistema	Confirmar que el sistema fue revisado y aceptado formalmente
	Certificaciones técnicas de producto o de modelos (si aplica).	Comprobar cumplimiento de estándares técnicos del modelo de IA (si aplica) y su conformidad con requisitos legales o sectoriales.
Gestión financiera	Informe de costes operativos y consumo de <i>tokens</i> con evaluación, en su caso, de alternativas	Supervisar la viabilidad económica del proyecto y evitar desviaciones presupuestarias imprevistas (el consumo de determinados modelos puede encarecer significativamente la obtención del resultado)
Propiedad Intelectual y legal	Cláusulas de propiedad de los resultados (si aplica)	Confirmar quién ostenta la titularidad de los contenidos generados según los contratos con proveedores externos
Ética y responsabilidad social	Evaluación de impacto ético	Documentar cómo se han mitigado riesgos de discriminación, exclusión o manipulación del comportamiento
	Protocolo de revisión de sesgos en colectivos vulnerables	Comprobar que el modelo no perpetúa prejuicios que afecten a grupos específicos.
Sostenibilidad	Declaración de impacto medioambiental o consumo energético, si aplica	Cumplir con los requisitos de transparencia sobre la huella de carbono asociados a modelos de gran tamaño
Contratos	Contratos de la herramienta de IA Contratos de la nube	Cumplir con la LCSP, con la Ley de Protección de datos y con el ENS

Anexo 2. Conceptos básicos de Inteligencia Artificial relevantes para la auditoría

Los sistemas de IA Generativa funcionan en dos fases claramente diferenciadas. Una **fase de entrenamiento**, donde se generan los modelos conocidos como LLM (Large Language Models) que sirven como base para el funcionamiento de la IA Generativa. Y una **fase de inferencia** donde se utilizan estos modelos para realizar la respuesta a las cuestiones o trabajos que se le plantean.

Fase de entrenamiento

Es el proceso mediante el cual el modelo aprende patrones lingüísticos, relaciones semánticas y estructuras de conocimiento a partir de grandes volúmenes de datos textuales. Es una etapa crítica, ya que determina las capacidades, limitaciones y riesgos del modelo durante su operación (inferencia).

A nivel conceptual, el entrenamiento se puede dividir en cuatro componentes principales:

1. Preparación de datos (limpieza, etiquetado y anonimización).
2. Selección de la arquitectura de entrenamiento y algoritmos.
3. Entrenamiento del modelo base (*pre-training*).
4. Ajuste posterior o *fine-tuning* (para adaptar el modelo a tareas específicas).

Esta fase de entrenamiento puede llevar de días a meses, y requiere unos recursos de computación muy elevados en coste y tiempo para poder obtener modelos que sean lo más precisos posible.

Sin entrar en el detalle de cada uno de los componentes comentados, existen varios **riesgos** que se deben contemplar si se audita esta fase de entrenamiento:

- La licitud y permisos asociados a los datos de entrenamiento (propiedad intelectual, incluida).
- Los sesgos en los datos introducidos que pueden hacer que los modelos tengan respuestas con tendencias no deseadas (sexistas, racistas...).
- La falta de documentación y trazabilidad de los datos (gobernanza del dato).
- El acceso a los datos de entrenamiento si estos son privados, como datos médicos o financieros (riesgo de reidentificación o fuga de datos y vulneración de la privacidad).

Fase de inferencia

La **fase de inferencia** es el proceso mediante el cual un LLM ya entrenado **genera respuestas, predicciones o contenido** a partir de una entrada proporcionada por el usuario o por otro sistema (*prompt*). Es la fase **operativa**, donde el modelo aplica los **patrones aprendidos** durante el entrenamiento para producir resultados. El concepto fundamental es entender que en esta fase se aplican las fórmulas matemáticas con los valores aprendidos durante la fase de entrenamiento para predecir la siguiente palabra en la respuesta que se le da al usuario.

A diferencia del entrenamiento (intensivo, costoso y previo a la puesta en servicio), la inferencia es el proceso que tiene impacto directo en la toma de decisiones, en la calidad del servicio y en los derechos de los ciudadanos cuando se utiliza en contextos administrativos.

Dentro de las características principales de esta fase podemos destacar:

- **Proceso indeterminista:** incluso con la misma entrada, el modelo puede producir respuestas distintas debido al funcionamiento intrínseco de estos LLMs, en base a estadísticas y a los parámetros configurados.
- **Dependencia del contexto:** los LLM no “razonan” en sentido humano, sino que generan el texto que mejor encaja estadísticamente con el contexto recibido.

- **Ausencia de memoria permanente:** el modelo no aprende de las interacciones individuales o previas salvo que exista una arquitectura externa que lo permita (como la memoria de sesión). Cada inferencia es independiente.

Esta fase de inferencia contempla varios riesgos que se deben tener en cuenta durante los trabajos de fiscalización:

- Alucinaciones del modelo
- Sesgos en las respuestas
- Falta de explicabilidad
- Dependencia tecnológica de terceros

Características fundamentales para el control interno de la IA Generativa²¹

Las capacidades y limitaciones inherentes a la IA Generativa tienen un impacto global en el diseño e implementación de un sistema de control interno. Las siguientes características fundamentales pueden utilizarse para guiar cómo deben construirse o adaptarse los controles, independientemente de los casos de uso o la tecnología empleada.

- **La IA Generativa es probabilística, no determinista.**

Los resultados de la IA Generativa son probabilísticos y pueden estar equivocados. Los controles deberían tratar las salidas como afirmaciones que requieren validación sistemática, en lugar de como hechos que se deben aceptar por defecto.

- **La IA Generativa es dinámica.**

Los modelos, los *prompts* y los datos que utilizan para la recuperación evolucionan con el tiempo. La evaluación de riesgos, el control del cambio y el seguimiento deben ser continuos, o acercarse hacia lo continuo, para mantener el ritmo de estos cambios y evitar la degradación del modelo.

- **La IA Generativa es fácilmente escalable (para bien o para mal).**

Una mayor automatización puede escalar la calidad y la eficiencia, pero también puede escalar errores y sesgos a gran velocidad. Los controles deben diseñarse para evitar que pequeños errores se conviertan en problemas sistémicos.

- **La IA Generativa tiene una barrera de entrada baja.**

Los controles deberían diseñarse para gobernar quién puede construir, desplegar e interactuar con la IA Generativa, evitando, por ejemplo, el uso no controlado a través de cuentas personales en herramientas comerciales.

- **La IA Generativa puede ayudar a gobernar la IA Generativa.**

Sus capacidades analíticas y de reconocimiento de patrones pueden fortalecer la gobernanza (por ejemplo, mediante la validación automatizada multimodelo que cruza los resultados entre modelos independientes para detectar inconsistencias, sesgos o deriva). Cuando se implementa correctamente, la IA Generativa puede mejorar actividades de monitorización, documentación y validación que de otro modo serían poco prácticas a gran escala.

²¹ Achieving Effective Internal Control Over Generative AI, COSO, 2026.

Anexo 3. Marco normativo aplicable

El marco normativo que debe tenerse en cuenta cuando se utilizan sistemas de IA lleva asociado un conjunto de normas europeas y nacionales que priorizan la seguridad y la protección de los derechos fundamentales de las personas físicas frente al uso de la Inteligencia Artificial.

Dicho conjunto de normas se agrupa en tres bloques relevantes relativos a:

- el cumplimiento legal de la Inteligencia Artificial,
- la protección de datos de carácter personal, y
- la seguridad de la información.

Normativa de la Inteligencia Artificial

El cumplimiento legal de la Inteligencia Artificial ha sido desarrollado en el ámbito de la Unión Europea (UE) por el **REGLAMENTO (UE) 2024/1689** por el que se establecen normas armonizadas en materia de Inteligencia Artificial (**RIA**). Este Reglamento busca el fomento de una Inteligencia Artificial centrada en el ser humano, fiable y segura a través del establecimiento de un enfoque basado en el riesgo que garantice al mismo tiempo un elevado nivel de protección de la salud, la seguridad y los derechos fundamentales. Su entrada en vigor completa está prevista para el mes de agosto de 2026, salvo que se produzca un aplazamiento como resultado de la aprobación del Reglamento Digital Ómnibus de la UE que se encuentra en tramitación durante 2026. No obstante, la aplicabilidad para sistemas de IA de alto riesgo que están siendo utilizados por las administraciones públicas se extiende a agosto de 2030.

El RIA establece un conjunto de obligaciones cuyo cumplimiento se regula en función del tipo de sistema, el tipo de operador que participa en la cadena de valor del sistema y el nivel de riesgo asociado al sistema IA.

En relación con el **tipo de sistema de IA**, el RIA hace una distinción simple, distinguiendo solo dos tipos, los sistemas de IA de finalidad específica (clasificar expedientes, detectar anomalías o apoyar una decisión administrativa específica) y los sistemas de IA de **uso general**. Estos sistemas de IA de uso general, denominados en la guía IA Generativa, podrán basarse en el uso de los modelos fundacionales o los grandes modelos comercializados con una elevada capacidad de cómputo y parametrización y en cuyo caso tendrán la consideración de riesgo sistémico. Los modelos de IA de uso general de riesgo sistémico llevan asociados obligaciones específicas del RIA.

En relación con el **tipo de operadores**, el RIA establece una clasificación de actores: proveedores, responsables del despliegue, importadores, distribuidores y fabricantes entre otros. Cada rol tiene asociado la aplicación de obligaciones específicas y que deben ser tenidas en cuenta tanto por organizaciones públicas como privadas cuando hagan uso de sistemas de IA. En las administraciones públicas el rol aplicable será fundamentalmente el de proveedor del sistema de IA cuando lo desarrolle y lo ponga en servicio en su nombre y/o el de responsable del despliegue que utiliza el sistema IA bajo su propia autoridad.

Finalmente, el **enfoque de riesgo** del RIA clasifica los sistemas de IA en cuatro niveles (inaceptable, alto, limitado y mínimo), prohibiendo prácticas inaceptables que vulneran los derechos fundamentales, como la puntuación ciudadana o la manipulación subliminal, imponiendo requisitos estrictos a los sistemas de IA de alto riesgo en sectores críticos tales como la salud, la educación, la justicia y el empleo y reduciendo las obligaciones para el resto de los sistemas con requisitos de alfabetización y transparencia. Por ello, será importante que se tenga en cuenta el tipo de sistema de IA que en términos generales describimos a continuación:

- **Sistemas de riesgo inaceptable:** los sistemas de IA considerados una clara amenaza para la seguridad, los medios de subsistencia y los derechos de las personas **están prohibidos**. Se incluyen ocho prácticas concretas como manipulación y engaño perjudiciales, explotación nociva de vulnerabilidades de las personas para influir en su comportamiento, puntuación o calificación social (*social scoring*), evaluación o predicción del riesgo de infracción penal individual, recopilación indiscriminada de reconocimiento facial, reconocimiento de emociones en entornos laborales y educativos y categorización e identificación biométrica remota en tiempo real.
- **Sistemas de riesgo alto:** son aquellos casos de uso que pueden plantear riesgos graves para la salud, la seguridad o los derechos fundamentales. Por ejemplo, sistemas de transporte, cuyo fallo podría poner

en peligro la vida y la salud de los ciudadanos, o soluciones de IA utilizadas en la administración de justicia y los procesos democráticos. Estos sistemas están sujetos a obligaciones estrictas antes de poder comercializarse. Para el auditor es clave identificar si el sistema se usa en administración de justicia, procesos electorales, educación, empleo, salud o servicios públicos esenciales.

Estos sistemas enumerados en el Anexo III del RIA, requieren, además de obligaciones generales de transparencia y formación, un conjunto de obligaciones concretas como evaluaciones de impacto, conservación del registro de actividades (*logs – prompts/output*), notificación de incidentes y una supervisión humana efectiva por diseño.

El propio reglamento, establece como sistemas de alto riesgo los servicios que usen IA para conceder prestaciones o subvenciones a la población como:

*“...las personas físicas que solicitan a las autoridades públicas o reciben de estas prestaciones y servicios esenciales de asistencia pública, a saber, **servicios de asistencia sanitaria, prestaciones de seguridad social, servicios sociales que garantizan una protección en casos como la maternidad, la enfermedad, los accidentes laborales, la dependencia o la vejez y la pérdida de empleo, asistencia social y ayudas a la vivienda**, suelen depender de dichas prestaciones y servicios y, por lo general, se encuentran en una posición de vulnerabilidad respecto de las autoridades responsables... y, por lo tanto, deben clasificarse como de alto riesgo...”*

- **Sistemas de riesgo limitado:** los sistemas de riesgo limitado son aquellos que, sin afectar directamente a la seguridad ni a los derechos fundamentales, pueden inducir a error si el usuario no sabe que está interactuando con una inteligencia artificial. Su principal obligación es garantizar la transparencia, informando de forma clara y visible que se trata de un sistema automatizado. Incluyen, entre otros, *chatbots*, asistentes virtuales, generadores de contenido sintético o sistemas de recomendación que interactúan con el público. El riesgo no proviene de sus decisiones, sino de la posibilidad de que el usuario atribuya erróneamente intervención humana o imparcialidad. Por ello, deben incorporar mecanismos de identificación y etiquetado para evitar confusión o manipulación.
- **Sistemas de riesgo mínimo:** los sistemas de riesgo mínimo son aquellos cuyo impacto sobre derechos, libertades o seguridad es tan bajo que el Reglamento de IA no impone obligaciones específicas para su desarrollo o uso (filtros de spam, videojuegos, por ejemplo). Constituyen la gran mayoría de las aplicaciones de IA actualmente utilizadas en la UE, como filtros de spam, autocompletado o asistentes básicos. Aunque están exentos de requisitos regulatorios, el RIA fomenta la adopción voluntaria de buenas prácticas para promover una IA fiable y responsable. Estos sistemas no toman decisiones que afecten a personas ni procesan información sensible, por lo que su utilización se considera segura dentro del marco europeo. Su regulación es mínima para incentivar la innovación sin cargas innecesarias.

Es necesario comentar que, en los últimos años, se han aprobado diversas leyes a nivel autonómico relacionadas con el uso ético de la Inteligencia Artificial y se encuentra en tramitación una ley de ámbito nacional cuyos objetivos se orientan al buen uso y gobernanza de la Inteligencia Artificial.

Normativa de protección de datos

La normativa de protección de datos está compuesta principalmente por el **Reglamento (UE) 2016/679 (RGPD)** y la **Ley Orgánica 3/2018 (LOPDGDD)** cuyo objetivo es proteger los derechos y libertades fundamentales de las personas físicas en lo que respecta al tratamiento de sus datos personales, teniendo los ciudadanos el poder de disposición de sus datos. Por ello, es crucial que, en el uso de sistemas de IA, donde el tratamiento de datos personales, tanto en fase de entrenamiento como los que procesa en la fase de inferencia, son una constante, los tratamientos realizados por estos sistemas de IA se haga conforme a lo estipulado en dichas normativas tanto en el uso, gestión y responsabilidad sobre los datos.

Desde su entrada en aplicación en 2018, las administraciones públicas cuentan con una experiencia dilatada en el cumplimiento de las obligaciones que lleva asociada la protección de datos. Todo sistema de IA que trate datos de carácter personal deberá ser incluida como una operación de tratamiento o incluso constituir por sí misma una actividad de tratamiento y deben tenerse en cuenta, durante todo el ciclo de vida de los tratamientos, los principios básicos de licitud, transparencia y lealtad, limitación de la finalidad, minimización de datos, exactitud, conservación y garantía de integridad y confidencialidad. Además de observar los principios de privacidad por diseño y por defecto y la atención de los derechos de protección de datos.

El deber de información es un elemento clave de estas normativas y debe tenerse en cuenta por las organizaciones cuando realicen tratamientos con este tipo de sistemas IA, especialmente por el impacto que estos puedan tener sobre derechos fundamentales. Asimismo, la propia normativa establece niveles adicionales de criticidad en el caso de que se traten datos personales de categorías especiales donde se exigen la realización de evaluaciones de impacto de protección de datos, así como ser mucho más exigente en cuanto a la seguridad de los sistemas de IA y obligaciones de notificación de brechas de datos a la correspondiente autoridad de control nacional o autonómica.

Normativa de ciberseguridad

Con la finalidad de garantizar la seguridad de los sistemas de información de las administraciones públicas, todo organismo público que utilice IA deberá cumplir con el **ENS** según el **Real Decreto 311/2022**. Además, lo más habitual será que los sistemas de IA estén implementados en nubes externas por la gran cantidad de recursos que se necesitan, las cuales suelen ser gestionadas por proveedores privados, que también están obligados a cumplir con el ENS, tal y como se comenta en la GPF-OCEX 1403/guías CCN-STIC 808, guías que deberán tomarse como referencia para revisar los contratos vinculados a estos servicios. Esta precaución contractual debe también tenerse en cuenta en el ámbito de la protección de datos cuando los proveedores privados actúen como encargados del tratamiento, siendo necesaria la incorporación de clausulado adicional que regule lo estipulado en el artículo 28.3 del RGPD.

Régimen sancionador

Finalmente, debe tenerse en cuenta que las administraciones públicas que hagan uso de sistema IA estarán sujetas al régimen sancionador establecido por las normativas anteriores.