

La pregunta:

¿El documento “*Achieving Effective Internal Control Over Generative AI (GenAI)*” [Cfr. https://ab5bd07f-8066-4a33-869e-6f82d80ee59a.filesusr.com/ugd/719ba0_08f358f2c8f946fa9d26bd51d37b7117.pdf] permite realmente lograr un control interno efectivo sobre la inteligencia artificial generativa?

Las respuestas:

COPILOT:

Sí **puede ayudarte de forma realista a lograr un control interno “efectivo” sobre GenAI**, pero **no lo “garantiza” por sí mismo**. En términos de control interno, *ningún* documento (por bueno que sea) hace que el sistema sea efectivo automáticamente: la efectividad depende de **diseño + implementación + operación sostenida + evidencia**. Lo valioso del documento de COSO es que funciona como **hoja de ruta estructurada** para construir ese sistema, usando un marco ya aceptado por auditoría y gobierno corporativo. [coso.org], [journalofaccountancy.com]

1) Qué promete (y por qué es relevante para “control interno efectivo”)

Según la propia presentación de COSO, *Achieving Effective Internal Control Over Generative AI (GenAI)* se plantea como una **guía práctica y alineada a COSO** para gestionar riesgos y oportunidades de GenAI, precisamente porque la adopción está ocurriendo “más rápido” que los modelos de gobernanza tradicionales, y aparecen riesgos que pueden dañar la **integridad** de operaciones, reportes y cumplimiento si no hay controles robustos.

En la nota del *Journal of Accountancy* se refuerza que el texto **aplica los cinco componentes** del COSO ICIF (Control Environment, Risk Assessment, Control Activities, Information & Communication, Monitoring Activities) al contexto GenAI, lo cual es clave porque “control interno efectivo” en COSO no es una lista de chequeo tecnológica, sino un **sistema integrado**. [coso.org], [journalofaccountancy.com], [journalofaccountancy.com], [finance.yahoo.com]



SAMUEL MANTILLA

<https://samuelmantilla.substack.com>
<https://samuelmantilla.com>

Interpretación práctica: el documento sí está diseñado para “permitir” la efectividad en el sentido COSO: te da un **andamiaje completo** para diseñar controles que luego puedan operar y auditarse bajo un lenguaje que auditoría/ICFR/SOX ya entiende. [[journalofaccountancy.com](#)], [[finance.yahoo.com](#)]

2) Lo que aporta “de verdad” para operacionalizar (más allá de teoría)

Una crítica frecuente a guías de IA es que se quedan en principios abstractos. Aquí, la cobertura pública destaca varios elementos **operacionales**:

a) Taxonomía por capacidades (8 tipos) en vez de por herramientas.

El *Journal of Accountancy* indica que los casos de uso se organizan en ocho “capability types” (por ejemplo: ingestión, transformación, orquestación, juicio, monitoreo, interacción humano-IA, etc.), con consideraciones de control específicas según dónde se manifiestan los riesgos “en el ciclo de datos a decisión”.

Esto es importante porque en la empresa real hay muchos proveedores/modelos/capas, pero los riesgos se parecen más por **función** que por marca. Esa forma de estructurar ayuda a que el control interno

sea **replicable** y **escalable**. [[journalofaccountancy.com](#)], [[radicalcompliance.com](#)] [[radicalcompliance.com](#)], [[journalofaccountancy.com](#)]

b) Mapeo de controles “audit-ready” y artefactos de implementación.

La nota del *Journal of Accountancy* menciona “audit-ready control mapping” con **ejemplos**, **expectativas mínimas**, y **métricas ilustrativas** para monitoreo y evidencia; además de **plantillas** (matrices de riesgo, procedimientos de pruebas de control, tableros de métricas) para acelerar implementación.

Esto es precisamente lo que suele faltar para que una organización pase de “principios” a “control interno efectivo”: **evidencia y testeabilidad**. [[journalofaccountancy.com](#)], [[finance.yahoo.com](#)]

3) Lo que reconoce como riesgos (y por qué eso afecta la “efectividad”)

COSO subraya que GenAI trae riesgos como mayor exposición ciber, manipulación por prompts, razonamiento opaco, *model drift* y cambios frecuentes de configuración—todos factores que pueden romper controles tradicionales si no se adaptan.

El análisis de *Radical Compliance* refuerza la idea: la guía nace porque los riesgos evolucionan tan rápido como la tecnología; por eso insiste en **mecanismos de gobernanza durables** y en que el monitoreo continuo es crítico “dado el ritmo de cambio”. [[coso.org](#)], [[journalofaccountancy.com](#)] [[radicalcompliance.com](#)], [[coso.org](#)]

Conclusión parcial: el documento sí se toma en serio el motivo por el cual GenAI es difícil de controlar internamente: no es sólo “automatización”; es variabilidad, opacidad y velocidad de cambio.

Eso mejora la probabilidad de que, bien aplicado, logres un sistema efectivo. [[coso.org](#)], [[radicalcompliance.com](#)]

4) La parte clave: “permitir” no es “garantizar”

Aquí está la diferencia importante:

- La guía se presenta como **roadmap**/enfoque práctico para gobernar riesgos y controles de GenAI con COSO, y para que sea **audit-ready**. [coso.org], [journalofa...ntancy.com]
- Pero *Radical Compliance* describe que el propósito es “organizar tu pensamiento de forma disciplinada” y aplicar los **17 principios** del marco COSO (2013) a issues de IA; eso implica que el resultado depende de **cómo tu organización** diseñe, implemente, pruebe y monitoree. [radicalcom...liance.com], [journalofa...ntancy.com]

En otras palabras: **sí permite** (porque da estructura, ejemplos, mínimos y artefactos), pero **no sustituye**:

1. el gobierno corporativo real (roles, accountability),
2. el inventario y priorización de casos de uso,
3. la gestión de datos, ciberseguridad, terceros, y
4. la operación y monitoreo continuo con evidencia. [coso.org], [radicalcom...liance.com]

Un punto útil que aparece en un análisis externo (ITGRC Advisory) es la noción de “AI reliance” en ICFR: si la gerencia depende de salidas de IA como evidencia para un control, entonces debes tener estándares de evidencia (prompts/config/versiones, etc.) comparables a controles automatizados tradicionales. (Ojo: esto lo afirma ese análisis, no estoy citando el PDF directamente). [itgrcadvisory.com], [journalofa...ntancy.com]

5) Señales de que sí puede darte “control interno efectivo” (si lo implementas bien)

Con base en lo que publican COSO y fuentes que reseñan el contenido, hay varias señales “fuertes” de utilidad práctica:

1. **Cobertura completa del modelo COSO** (5 componentes + 17 principios aplicados) → eso es lo más cercano a “criterio de efectividad” reconocido por auditoría. [journalofa...ntancy.com], [radicalcom...liance.com]
2. **Descomposición por capacidades** → facilita inventario, riesgo y control a escala. [journalofa...ntancy.com], [radicalcom...liance.com]
3. **Plantillas y procedimientos** → acelera implementación y, sobre todo, habilita **pruebas** y evidencia (condición práctica para que sea “efectivo” y defendible). [journalofa...ntancy.com], [finance.yahoo.com]
4. **Enfoque en monitoreo por cambio constante** (*drift*, configuraciones) → sin eso, los controles “se vencen” rápido en GenAI. [coso.org], [radicalcom...liance.com]

6) Qué encontré (y no encontré) en tu entorno de trabajo

Busqué en tus archivos corporativos una copia o referencia directa a *Achieving Effective Internal Control Over Generative AI (GenAI)* y **no apareció**; lo más cercano que encontré fue un PDF de COSO sobre **RPA**: COSO RPA Paper Dec 2024.pdf.

Eso no invalida tu pregunta—solo significa que, en tu repositorio actual, no tengo el PDF de GenAI para citar su texto exacto; por eso estoy basándome en las descripciones públicas y reseñas disponibles. [\[COSO RPA P...r Dec 2024 | PDF\]](#) [\[coso.org\]](#), [\[journalofa...ntancy.com\]](#)

Veredicto (claro y accionable)

Sí, el documento “permite” lograr un control interno efectivo en el sentido de que provee un marco COSO aplicado a GenAI con componentes auditables (taxonomía, mapeo de controles, métricas y plantillas).

Pero no lo hace automáticamente: la efectividad real depende de que tu organización lo convierta en controles operativos con gobierno, evidencia y monitoreo continuo, justamente porque GenAI cambia rápido (drift, prompts, configuraciones) y puede afectar operaciones, reporting y compliance si no hay controles robustos.



SAMUEL MANTILLA

<https://samuelmantilla.substack.com>
<https://samuelmantilla.com>

GEMINI:

¡Es una excelente pregunta! El documento "**Achieving Effective Internal Control Over Generative AI (GenAI)**", publicado por **COSO** en febrero de 2026, no es una "solución mágica" que garantice el control por sí sola, pero sí proporciona la **hoja de ruta práctica** necesaria para lograrlo.

La respuesta corta es: **sí, permite lograrlo**, siempre y cuando la organización no lo vea como un check-list estático, sino como un marco adaptable a una tecnología que evoluciona casi a diario.

Aquí te detallo por qué este documento es clave para un control real:

1. No inventa la rueda, adapta una estructura probada

En lugar de crear un marco nuevo desde cero, el documento traduce los **cinco componentes del Marco Integrado de COSO** al contexto específico de la GenAI:

- **Entorno de Control:** Define la gobernanza y la cultura de responsabilidad necesaria antes de desplegar modelos.
- **Evaluación de Riesgos:** Identifica riesgos nuevos como las "alucinaciones", sesgos, deepfakes y la manipulación de modelos mediante prompts.
- **Actividades de Control:** Establece controles específicos para el ciclo de vida del dato (desde la ingesta hasta la decisión final).
- **Información y Comunicación:** Asegura que los resultados de la IA sean trazables y auditables.
- **Monitoreo:** Propone una vigilancia continua de la precisión y equidad de los modelos.

2. Introduce una "Taxonomía de Capacidades"

Para que el control sea efectivo, el documento organiza los usos de la GenAI en **8 categorías** (ingesta, transformación, orquestación, juicio, etc.). Esto permite que los auditores y gestores de riesgos identifiquen exactamente **dónde** se origina el riesgo en un proceso específico, en lugar de tratar a la IA como una "caja negra".

3. Enfoque en la "Auditalidad" (Audit-Ready)

Uno de sus mayores valores es que prepara a las empresas para cumplir con estándares exigentes (como el ICFR para reportes financieros). Si una empresa usa GenAI para generar datos contables, el documento explica qué evidencias debe guardar: configuraciones, versiones del modelo, prompts utilizados y razonamiento del muestreo.

En resumen

El documento es efectivo porque reconoce que la GenAI es "**confiadamente errónea**" y fácil de manipular. Al aplicar los principios de COSO con "rigor renovado", permite que las organizaciones dejen de experimentar a ciegas y comiencen a operar con sistemas adaptables y listos para ser auditados.



SAMUEL MANTILLA

<https://samuelmantilla.substack.com>
<https://samuelmantilla.com>

Nota: El éxito real dependerá de que el monitoreo sea continuo, ya que un control que hoy es efectivo puede dejar de serlo mañana si el modelo de IA sufre una "degradación" (drift).

¿Te gustaría que profundizara en alguno de los 5 componentes específicos de COSO aplicados a la IA?

ChatGPT 5.0

Sí **puede ayudar**, pero **no garantiza por sí mismo lograr un control interno efectivo sobre la inteligencia artificial generativa**. El documento **Committee of Sponsoring Organizations of the Treadway Commission (COSO)** titulado *“Achieving Effective Internal Control Over Generative AI (GenAI)”* es **una guía de gobierno y control**, no una solución automática. Su valor depende de cómo se implemente en la organización. A continuación te explico por qué.

1. Qué pretende realmente el documento

El documento adapta el marco clásico de control interno de **COSO —Internal Control–Integrated Framework—** al uso de la IA generativa. Este marco se basa en cinco componentes:

1. **Ambiente de control**
2. **Evaluación de riesgos**
3. **Actividades de control**
4. **Información y comunicación**
5. **Monitoreo**

La publicación **no crea un nuevo framework**, sino que **traduce esos componentes al contexto de GenAI**.

También introduce una taxonomía de **ocho capacidades de GenAI** (por ejemplo: ingestión, transformación, juicio, interacción humano-IA) para identificar riesgos y diseñar controles específicos en cada etapa del ciclo de datos y decisiones.

2. Lo que el documento sí logra

El documento **sí permite avanzar hacia un control interno efectivo** porque:

a) Integra GenAI en un marco de control ya probado

El enfoque de COSO es ampliamente utilizado en auditoría, control interno y cumplimiento regulatorio.

Esto permite:

- Integrar GenAI en **controles existentes (ICFR, SOX, auditoría interna)**.
- Evitar crear estructuras paralelas de gobernanza.

b) Introduce controles auditables

El documento propone evidencia y trazabilidad, por ejemplo:

- registro de **prompts**
- versiones de modelos
- configuraciones



SAMUEL MANTILLA

<https://samuelmantilla.substack.com>
<https://samuelmantilla.com>

- justificación de muestras
- documentación de resultados

Esto busca que los sistemas de IA puedan **ser auditados como cualquier otro control automatizado**.

c) Ofrece artefactos prácticos

Incluye:

- matrices de riesgo
- procedimientos de prueba de controles
- dashboards de métricas
- plantillas de implementación

Esto reduce la distancia entre teoría y práctica.

3. Lo que el documento NO puede garantizar

Aunque es valioso, **no puede garantizar control efectivo** por tres razones estructurales.

1. La naturaleza probabilística de la GenAI

A diferencia de los sistemas tradicionales:

- las respuestas **no son determinísticas**
- el mismo prompt puede producir resultados distintos

Esto hace que el control tradicional basado en resultados esperados sea más difícil.

2. Los riesgos evolucionan más rápido que los controles

Problemas como:

- **alucinaciones**
- **prompt injection**
- **model drift**
- cambios frecuentes de modelos

pueden aparecer más rápido de lo que los controles se actualizan.

3. La efectividad depende del gobierno organizacional

El documento **solo proporciona orientación**.

El control real depende de:

- madurez de gobierno de datos
- capacidades de auditoría
- supervisión humana

- cultura de control

Sin estos factores, el marco queda en papel.

4. Evaluación crítica (conclusión)

Desde una perspectiva de **gobernanza y auditoría**, el documento:

Fortalezas

- primer intento sistemático de aplicar COSO a GenAI
- orientado a auditoría y cumplimiento
- operacionaliza controles y evidencia

Limitaciones

- depende totalmente de implementación organizacional
- no resuelve riesgos técnicos de los modelos
- puede quedar obsoleto rápidamente frente a nuevas arquitecturas de IA

👉 En síntesis:

El documento no garantiza el control interno efectivo sobre la IA generativa, pero sí proporciona una base estructurada y práctica para construirlo.

💡 Si quieres, también puedo explicarte **3 críticas académicas importantes a este documento (que casi no se mencionan en auditoría)** sobre:

- el problema del **control sobre modelos fundacionales**
- la **responsabilidad distribuida en ecosistemas de IA**
- el límite del **enfoque COSO para sistemas probabilísticos**.



SAMUEL MANTILLA

<https://samuelmantilla.substack.com>
<https://samuelmantilla.com>

PERPLEXITY:

No, el documento “Achieving Effective Internal Control Over Generative AI (GenAI)” de COSO no permite por sí solo lograr un control interno efectivo sobre la IA generativa, ya que es un marco de orientación práctica que debe implementarse y adaptarse en cada organización. Proporciona una hoja de ruta estructurada basada en el marco COSO de control interno, pero su éxito depende de la ejecución real, el monitoreo continuo y la adaptación a riesgos específicos.

Origen del documento

Es una publicación reciente de COSO (Committee of Sponsoring Organizations of the Treadway Commission), lanzada alrededor de febrero de 2026, que adapta su marco clásico de control interno a los riesgos de GenAI. Autores como Scott Emett y David A. Wood contribuyeron, enfocándose en gobernanza, riesgos y assurance para IA generativa.

Contenido clave

Mapea riesgos de GenAI (como generación de insights, colaboración humano-IA y monitoreo automatizado) a los cinco componentes COSO: ambiente de control, evaluación de riesgos, actividades de control, información y comunicación, y monitoreo. Incluye controles comunes como asignación de ownership, revisión humana en lazo, umbrales de confianza y análisis de escenarios.

Hoja de ruta práctica

- Establecer estructura de gobernanza de IA.
- Inventariar casos de uso de GenAI (incluyendo "shadow AI").
- Evaluar riesgos por componente COSO.
- Diseñar y mapear controles.
- Implementar y comunicar.
- Monitorear y adaptar.

Limitaciones reales

Aunque es un puente valioso entre teoría y práctica, no es una solución mágica: requiere compromiso organizacional, entrenamiento y pruebas continuas para ser efectivo, especialmente ante la evolución rápida de GenAI. Expertos sugieren que podría influir en futuros reportes SOC de AICPA, pero el control efectivo surge de su aplicación, no solo de la lectura.

Seguimientos

Cuál es la hoja de ruta práctica para implementar los controles de GenAI



SAMUEL MANTILLA

<https://samuelmantilla.substack.com>
<https://samuelmantilla.com>

Qué riesgos específicos aborda en cada componente COSO

Cuáles son ejemplos de control activities para prompts y outputs AI

Cómo se aplican los controles a capacidades como human-AI collaboration

Qué evolución prevé para frameworks AICPA y SOC reporting con GenAI



SAMUEL MANTILLA

<https://samuelmantilla.substack.com>
<https://samuelmantilla.com>