

AI en acción en la AAPP: Ejemplos reales de utilización y marco ético

Jesús Aguilera Olarte



Retos actuales de la IA



Costes vs ROI

Control de riesgos

Operacionalización

Madurez IA

High Cancellation Rate (Gartner)

“Over 40% of Agentic AI projects will be canceled by the end of 2027, due to **escalating costs, unclear business value** or **inadequate risk controls.**”

Limited Investment and Hype (Gartner)

“Only 19% of organizations invest significantly in Agentic AI, while many projects are **early hype-driven experiments.**”

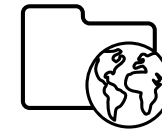
Integration and Workflow Challenges (McKinsey)

“About 90% of vertical (function-specific) use cases **remain stuck in pilot mode.**”

Maturity and Technical Barriers (Forrester)

“69% of organizations are still at **beginner or lower-intermediate levels of GenAI maturity.**”

Porqué y Cómo tenemos que afrontar los retos



Productividad/ Madurez IA

¿Por qué?

- **Time-to-Value**
- Adopción generalizada IA
- Reducción TCO



¿Cómo?

- Interfaces alta productividad
- Reutilización componentes
- Compartición de recursos
- Alfabetización IA

Intervención humana

¿Por qué?

- Decisiones más seguras
- Requerimiento regulatorio
- Responsabilidad y trazabilidad



¿Cómo?

- Asegurar balance entre determinista y probabilístico
- Basado en reglas alineadas con negocio
- Monitorización continua

Precisión decisiones

¿Por qué?

- Asegurar confianza y credibilidad
- Alineamiento con el negocio
- Impacto reputacional



¿Cómo?

- Añadir capacidades como RAG (Aumento de LLMs)
- Sacar partido de NLP
- Establecer controles y validación continua

Operacionalización

¿Por qué?

- Time-to-Market
- Ventaja competitiva
- Demostrar ROI



¿Cómo?

- Clouds publicas y privadas
- Containers & APIs
- Integración sistemas
- Test y validación
- Gobierno centralizado

Control de riesgos

¿Por qué?

- Asegurar adopción segura
- Cumplimiento regulatorio
- Prevenir impacto sistémico
- Protección de derechos



¿Cómo?

- Ethics "by design"
- Gestión y mitigación de riesgos
- Monitorización y vigilancia continua

Propuesta SAS para una IA en Acción y Confiable



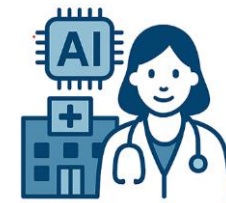
Algunos ejemplos de utilización...



Gobernanza de la IA y alineamiento con EU AI Act

Garantizar la precisión de IA aplicada a la medicina y el cumplimiento con las regulaciones

Gobernanza IA y AI Act para el ámbito Sanitario

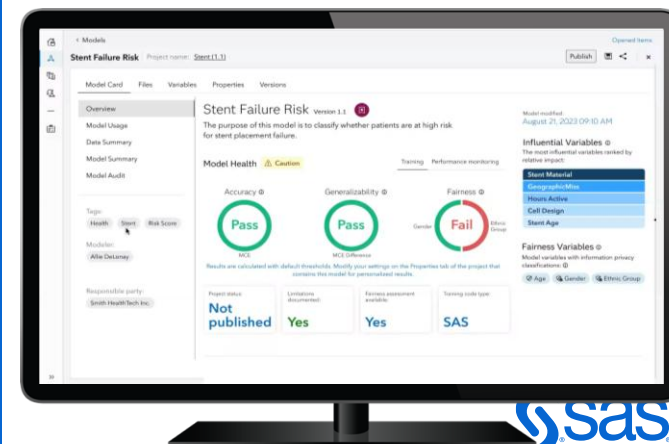
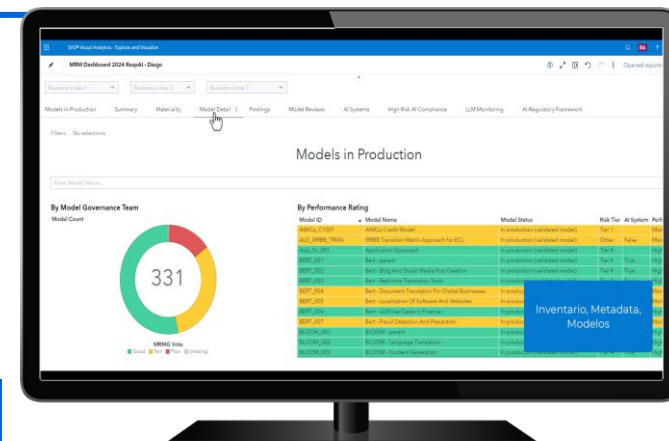


Alineamiento con las regulaciones existentes

Implementación de una metodología de gobernanza de acuerdo con el *EU AI Act* y otras políticas internas y regulaciones, con una **gestión continua de los riesgos** de la IA y asegurando **total transparencia y auditabilidad**, para dar una respuesta eficiente caso de auditorías..

Asegura el correcto rendimiento y gobernanza de los modelos de IA y la intervención humana

Monitorización continua de los sistemas de IA utilizados por el personal sanitario respecto a la precisión y equidad de las respuestas, alertando de posibles de incidencias



Beneficios



Asegura el cumplimiento regulatorio, evitando sanciones e impacto reputacional



Explicabilidad, transparencia y equidad combinada con precisión de los sistemas de IA



Máxima efectividad para la anticipación continua a posibles problemáticas

Sistema de Gestión de Incidencias

Mejora de la productividad y satisfacción de usuarios y ciudadanos

Al basada en asistentes para la mejora de la experiencia del ciudadano



Decisiones Inteligentes SAS como núcleo del sistema
para **orquestrar** modelos analíticos y LLMs basados en reglas de negocio con búsqueda avanzada en repositorios documentales (Retrieval Augemented Generation_RAG).

Priorización de solicitudes de forma avanzada y acelerada
Clasificación automática de la información con un control basado en reglas de negocio y con asistentes digitales para **aumentar la productividad de los agentes**, la efectividad, la **satisfacción de usuarios** y reducción de riesgos

Gobernanza, explicabilidad y auditabilidad de las decisiones
Transparencia y trazabilidad total en cada respuesta propuesta, con “guardar- railes” y **supervisión humana** (human in the loop).



Beneficios



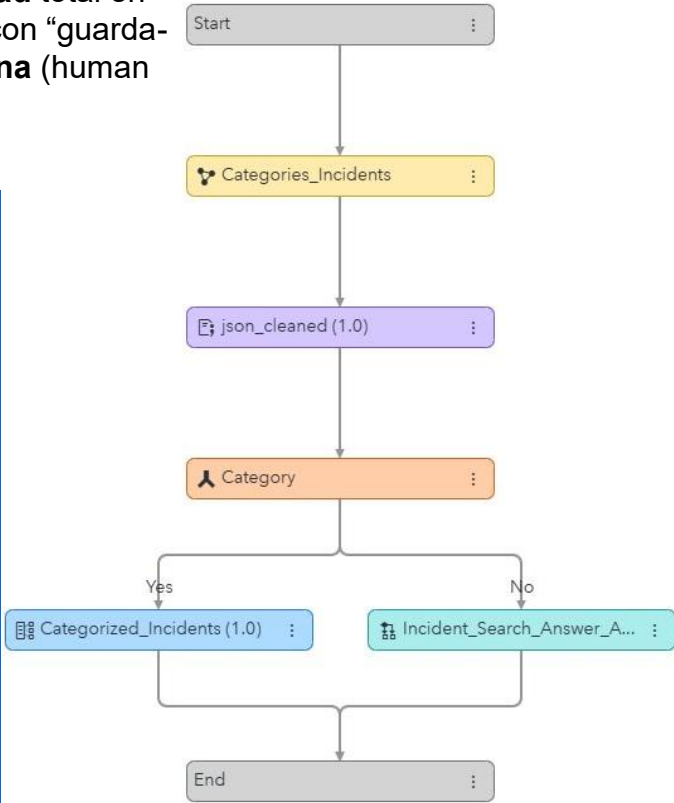
Efectividad y Productividad



Satisfacción usuarios y ciudadanos



Decisiones explicables con supervisión humana



Asistente de reglas para la lucha contra el Fraude

Agentic AI para Fraude



Hay una nueva tendencia de fraude mediante estafas. ¿Cómo puedo encontrar información para crear una regla que me ayude a detectar pagos autorizados?

Generación en base a tendencias

Genera automáticamente reglas antifraude inteligentes aprovechando las tendencias emergentes detectadas en la web para anticipar nuevos esquemas y reducir riesgos.

Ya no detecto fraude como antes, ¿puedes sugerirme nuevas reglas usando IA?

Sugerencia de reglas basadas en modelos de IA y Machine Learning

Genera reglas antifraude a partir de modelos de ML entrenado, combinando IA Generativa con IA “tradicional”.

Beneficios



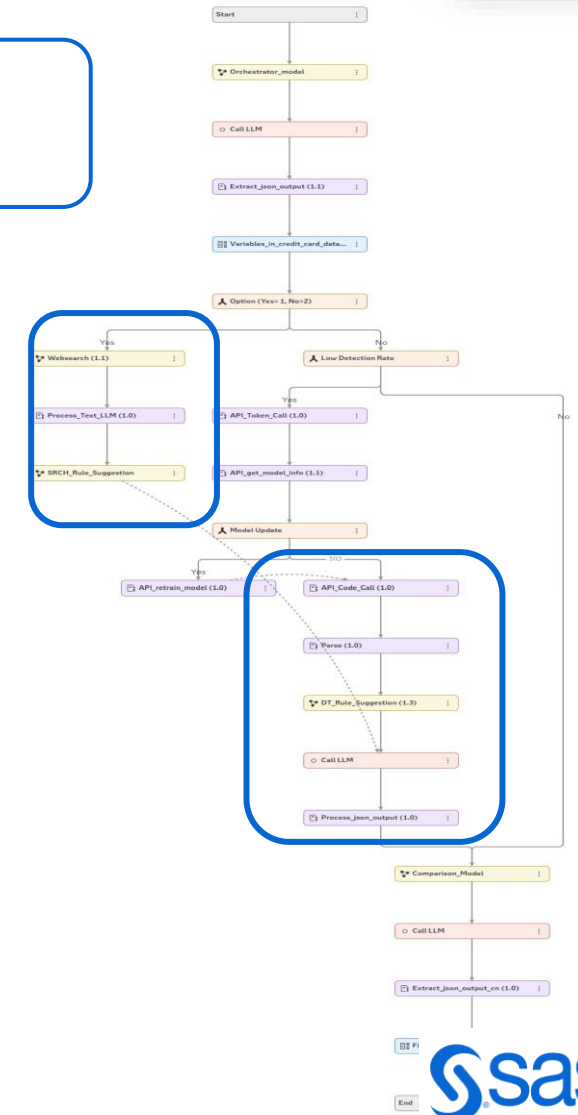
Máxima flexibilidad para la investigación, combinando tendencias emergentes, LLMs y modelos ML



Anticipación de nuevos patrones, basados en tendencias emergentes



Sugerencias basadas en IA y ML, para actualizar reglas existentes y mejorar efectividad



Muchas Gracias

