



OMAR SALINAS SILVA  
Director Ingeniería Civil Informática  
Advance UNAB

## Cuando la filosofía se convirtió en un problema de ingeniería

Que las principales empresas de inteligencia artificial estén contratando filósofos puede parecer una excentricidad o una estrategia para reforzar sus credenciales éticas. La realidad es mucho más interesante: descubrieron un problema que la ingeniería, por sí sola, no puede resolver. Durante décadas, programar consistía en escribir instrucciones para que un computador las ejecutara. La inteligencia artificial cambió esa lógica.

Hoy los modelos ya no siguen un listado de reglas; aprenden un criterio para responder y lo aplican a millones de situaciones que nadie pudo anticipar. El verdadero diseño dejó de estar en el código y pasó a las reglas que determinan cómo debe responder una máquina. Ahí aparece el desafío. Decirle a un sistema "no mientas" ocupa apenas una línea; demostrar que realmente no miente es mucho más difícil. ¿Ocultar información también es mentir? ¿Debe responder cuando no tiene certeza? ¿Es preferible equivocarse o guardar silencio? Antes de programar, alguien debe definir qué significan conceptos como verdad, daño, imparcialidad o responsabilidad. Esa tarea no pertenece a la informática, sino a la filosofía. No porque la inteligencia artificial necesite clases de Platón o Aristóteles, sino porque la ingeniería requiere definiciones precisas para construir sistemas verificables. Si nadie puede establecer qué significa "no mentir", tampoco existe una forma objetiva de comprobar que el modelo cumple esa promesa. Sin definiciones no hay pruebas; sin pruebas, no hay calidad. Esa es la razón por la que los grandes laboratorios tecnológicos comenzaron a contratar filósofos. No para redactar códigos de conducta, sino para convertir conceptos abstractos en criterios verificables que una máquina pueda aplicar de manera consistente millones de veces. Un criterio mal definido no perjudica a un usuario; perjudica a todos. Un sistema que responde de forma distinta frente al mismo dilema pierde confianza, valor comercial y la capacidad de justificar sus decisiones ante un regulador o un tribunal. Las consecuencias trascienden la tecnología. Los criterios que hoy orientan las respuestas de una inteligencia artificial condicionarán las que recibirán millones de personas cuando consulten sobre salud, educación, política o finanzas. En la práctica, terminan moldeando conversaciones cotidianas sin que la mayoría de los usuarios conozca quién escribió esas reglas ni bajo qué principios fueron concebidas.

Mientras Chile avanza en la discusión de una ley sobre inteligencia artificial centrada en la transparencia, la gestión de riesgos y las responsabilidades de quienes desarrollan estos sistemas, permanece abierta una pregunta decisiva: si esos principios influirán en millones de personas, ¿no deberían ser públicos, verificables y auditables?

La noticia nunca fue que las empresas de inteligencia artificial contrataran filósofos. La verdadera noticia es que comprendieron que una máquina nunca responderá mejor que los principios con los que fue diseñada. Y esa decisión, aunque parezca técnica, ya comenzó a definir el tipo de sociedad que estamos construyendo.