

2DO INFORME TRIMESTRAL 2026

El estado de la inteligencia artificial

Cuatro miradas sobre qué funciona,
qué cuesta y qué conviene supervisar.

Índice

Cuatro artículos y una presentación de la Comisión.

Editorial Lo que vas a encontrar en este informe

Ricardo Díez · Analytics Town

01 El oráculo, auditado: ¿qué tanto acertaron las predicciones del Mundial de Fútbol 2026?

Ricardo Díez · Analytics Town

02 Agentes de IA en marketing: casos reales, señales claras y por dónde empezar

María Laura Barreto · Doppler

03 Human-in-the-loop: el equilibrio entre escala y criterio

Natalia Borrazás · Growlat

04 El ROI real de la inteligencia artificial

Fernando Cuscuela · Wego Digital Hub

Cierre Sobre la Comisión de Inteligencia Artificial de la DMA

Integrantes y propósito

Lo que vas a encontrar en este informe

La inteligencia artificial dejó de ser una promesa hace rato. El problema, ahora, es separar lo que efectivamente funciona de lo que apenas suena bien.

Ricardo Díez

CEO de Analytics Town y líder de la Comisión de Inteligencia Artificial de la DMA Argentina

Este informe reúne cuatro trabajos de miembros de la Comisión que, cada uno desde su ángulo, apuntan a la misma pregunta: cómo distinguir el valor real del entusiasmo. No son cuatro textos sueltos sobre un tema de moda. Leídos en orden, se responden entre sí.

Abrimos con un experimento natural. Durante el Mundial 2026, bancos de inversión, universidades, supercomputadoras deportivas, mercados de predicción y modelos de IA publicaron sus probabilidades sobre quién sería campeón. Hoy conocemos el resultado, y eso habilita algo que en nuestro trabajo casi nunca se puede hacer: auditar pronósticos contra la realidad. Lo que aparece no es la historia de un ganador, sino una lección incómoda sobre cómo se evalúa —y cómo no— cualquier modelo predictivo que llegue a nuestra mesa.

De ahí pasamos a la capacidad nueva. María Laura Barreto explica qué son los agentes de IA autónomos, en qué se diferencian de lo que veníamos usando y qué está ocurriendo hoy en

equipos de marketing con objetivos y presupuestos reales. Cada caso de uso viene con un punto de entrada concreto, que es exactamente lo que suele faltar en este tipo de conversaciones.

Natalia Borrazás pone el contrapeso. Si los agentes ejecutan solos, ¿qué se supervisa y cuándo? Su artículo sobre *human-in-the-loop* propone pensar la supervisión como un espectro que se calibra según el costo del error, y no como un interruptor de encendido y apagado.

Y cerramos con la evidencia. Fernando Cuscuela desarma la comparación más frecuente en los comités ejecutivos —el precio de una licencia contra un salario— y la reemplaza por seis casos medidos que contrastan tres modelos operativos: humano puro, IA pura e híbrido. Los números confirman la intuición de los dos artículos anteriores. Con una excepción que vale por sí sola, y que conviene leer completa antes de aprobar el próximo presupuesto de IA.

Cuatro artículos, un hilo común: la IA rinde cuando alguien decidió con precisión qué delegar, qué supervisar y cómo medirlo.

Que lo disfruten.

◆ PRONÓSTICOS, PROBABILIDAD Y DECISIONES

El oráculo, auditado

¿Qué tanto acertaron las predicciones del Mundial de Fútbol 2026? Una auditoría a los modelos, los mercados y las inteligencias artificiales que intentaron adivinar al campeón.

Ricardo Díez

CEO de Analytics Town

01

En síntesis

Durante meses, bancos de inversión, universidades, supercomputadoras deportivas, mercados de apuestas y modelos de inteligencia artificial

publicaron sus probabilidades sobre quién levantaría la Copa del Mundo 2026. Hoy conocemos el resultado. Esta es la auditoría.

1-0

España campeón, a los 106 minutos

48

selecciones en el nuevo formato

- España se consagró campeón al vencer 1-0 a Argentina en tiempo suplementario. Inglaterra terminó tercera y Francia cuarta.
- Los cuatro semifinalistas fueron, exactamente, las cuatro selecciones que los principales modelos habían ubicado arriba antes del torneo. En un certamen de 48 equipos, eso no es poca cosa.
- **Goldman Sachs** fue el único que publicó la final exacta —España vence a Argentina— y además le asignó a España la probabilidad más alta de todo el tablero: 26%.
- El modelo **PELE de Silver Bulletin** (Nate Silver) rankeaba antes del primer partido a España, Argentina, Inglaterra y Francia. Terminaron en ese orden.

9

pronósticos públicos auditados

12,5×

mejor que el azar, el más certero

- Los **mercados de predicción** no fueron más certeros que los mejores modelos antes del torneo: fueron más rápidos después. Y sobrereaccionaron a la fase de grupos.
 - Los **modelos de lenguaje generalistas** —ChatGPT, Gemini, Claude— consultados tras la fase de grupos fallaron el campeón. Razonar con narrativa no es lo mismo que estimar probabilidades.
 - La lección de fondo: casi todos identificaron bien el grupo de candidatos; ninguno pudo separar con confianza al primero del cuarto. Con una sola Copa del Mundo de evidencia, “quién acertó” dice mucho menos de lo que parece.
-

1. Lo que finalmente pasó

El 19 de julio de 2026, en el MetLife Stadium de East Rutherford, Nueva Jersey, España venció 1-0 a Argentina con gol de Ferran Torres a los 106 minutos, ya en tiempo suplementario. Fue su segundo título mundial, dieciséis años después del de Sudáfrica 2010, y lo consiguió con un registro difícil de discutir: ocho partidos, siete victorias, un empate y un solo gol en contra, la valla menos vencida de un campeón del mundo en la historia del torneo. Del otro lado, Emiliano Martínez firmó once atajadas, la mayor cantidad jamás registrada en una final.

Argentina había llegado a la final tras remontar sobre el cierre ante Inglaterra en Atlanta: Anthony Gordon la puso en ventaja a los 55 minutos y la Albiceleste dio vuelta el partido con un remate lejano de Enzo Fernández a los 85 y un cabezazo de Lautaro Martínez a los 92, ambos asistidos por Lionel Messi. Se fue del Mundial con la sensación —justa— de haberlo dejado todo.

Inglaterra se quedó con el tercer puesto luego de un partido desquiciado frente a Francia en Miami: 6-4, con tres goles de Bukayo Saka. Fue el encuentro con más goles de un Mundial desde 1982 y el más goleador de la historia en un partido por el tercer lugar. Para Inglaterra es el primer tercer puesto de su historia y su mejor resultado desde el título de 1966.

Los premios individuales se repartieron entre campeones y eliminados. El Balón de Oro al mejor jugador del torneo fue para **Rodri**, el mediocampista español, con Lionel Messi segundo y Kylian Mbappé tercero. España completó la cosecha con el Guante de Oro para Unai Simón y el premio al mejor futbolista joven para Pau Cubarsí, por delante de Lamine Yamal. Mbappé se llevó la Bota de Oro —10 goles, contra 8 de Messi y 7 de Jude Bellingham— y se convirtió en el primer jugador en ganarla en dos Mundiales consecutivos. Con los dos tantos que

marcó en el partido por el tercer puesto llegó a 22 goles en Copas del Mundo y superó a Messi (21) como máximo artillero histórico de la competencia.

Pero el Mundial no fue solo la parte alta del cuadro. En dieciseisavos de final —la ronda nueva que estrenó el formato de 48 equipos— Paraguay eliminó a Alemania por penales y Marruecos hizo lo mismo

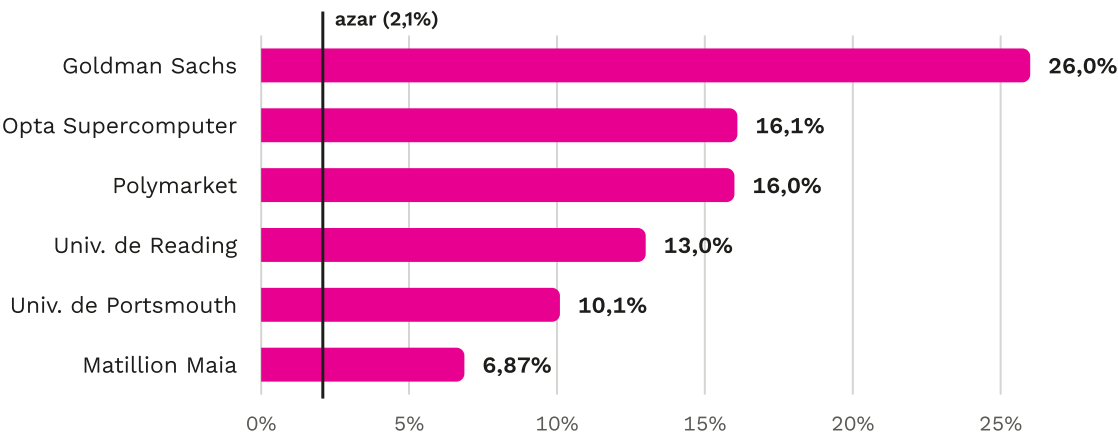
con Países Bajos. En octavos, Brasil cayó 2-1 ante la Noruega de Erling Haaland y quedó fuera de los cuartos por primera vez desde 1990; Portugal se despidió ante España con un gol de Mikel Merino en tiempo de descuento. De los ocho que alcanzaron cuartos de final, tres —Noruega, Marruecos y Suiza— no figuraban en el radar de casi nadie.

2. El tablero completo: qué se había dicho antes

La columna que importa no es la del campeón elegido, sino la probabilidad que cada uno le asignó a España antes de que rodara la pelota. Ahí se ve

quién vio venir el resultado y quién no.

Probabilidad asignada a España antes del torneo



Modelos con distribución publicada. La línea marca el 2,1% que correspondería a cada selección si el torneo se decidiera por sorteo entre 48 equipos. EA Sports y Klement publican pronósticos determinísticos, sin probabilidades.

Fuente	Enfoque	Campeón pronosticado	P(España)	Final proyectada
Goldman Sachs 29 may	Regresión sobre Elo; ~20.000 partidos desde 1978	España	26,0%	España vs. Argentina, campeón España
Opta Supercomputer 1 jun	25.000 simulaciones de Monte Carlo	España	16,1%	No publicada
Silver Bulletin – PELE 4 jun	100.000 simulaciones; Elo, valor de plantilla y localía	España	No pública	No publicada
Polymarket fines de mayo	Mercado de predicción con dinero real	Francia (17%)	16,0%	No aplica
Univ. de Reading 8 jun	10.000 simulaciones sobre goles esperados desde 2023	Argentina (24%)	13,0%	Campeón Argentina
Univ. de Portsmouth jun	1.000.000 de simulaciones; base histórica desde 1930	Inglaterra (15,9%)	10,1%	Campeón Inglaterra

Fuente	Enfoque	Campeón pronosticado	P(España)	Final proyectada
Matillion Maia jun	10.000 simulaciones; Elo ajustado por localía y viajes	España	6,87%	España vs. Inglaterra
EA Sports FC 26	Simulación del videojuego sobre los 104 partidos	España	No aplica	No publicada
Joachim Klement	Heurística con variables macroeconómicas y demográficas	Países Bajos	No aplica	P. Bajos vs. Portugal
Azar puro 48 equipos	Referencia de control	—	2,1%	—

Resultado real: España campeón, Argentina subcampeón, Inglaterra tercera, Francia cuarta.

3. El podio del meta-análisis

Goldman Sachs: el pleno

El equipo de Jan Hatzius publicó a fines de mayo un modelo de regresión construido sobre puntajes Elo, alimentado con unos 20.000 partidos internacionales desde 1978 y corregido por variables como la altitud de Ciudad de México y el llamado *winner's slump*, la caída estadística que suelen sufrir los campeones defensores. El resultado: España 26%, Francia 19%, Argentina 14%, Brasil 8%, Inglaterra 5%. Y una final: España vence a Argentina.

Es el único pronóstico público que combinó las tres cosas: nombró al campeón, nombró al subcampeón y acertó el ganador del cruce. Merece el crédito.

Vale recordar, eso sí, el historial. En 2014 el mismo banco había anunciado a Brasil campeón, y Brasil cayó en semifinales por 7-1 ante Alemania. En 2018 volvió a elegir a Brasil, esta vez proyectando una final contra Alemania: Brasil se fue en cuartos, Alemania ni siquiera superó la fase de grupos y el título quedó en manos de Francia. Y hay un detalle que conviene subrayar: aquel modelo de 2018 corría 200.000 modelos estadísticos y un millón de simulaciones, muchísimo más poder de cómputo que el de 2026, que acertó. **Más simulaciones no es lo mismo que mejor pronóstico.**

Silver Bulletin: el orden exacto

El modelo PELE de Nate Silver corrió 100.000 simulaciones del torneo combinando calificaciones Elo, valor de mercado de los planteles y ventaja de localía ajustada equipo por equipo. Sus cuatro selecciones mejor rankeadas al inicio fueron España,

Argentina, Inglaterra y Francia. La tabla final del Mundial: España, Argentina, Inglaterra, Francia. En ese orden.

Es probablemente el resultado más llamativo de todo el ejercicio, y también el que más invita a la prudencia. Acertar cuatro posiciones consecutivas es tanto la señal de un buen modelo como un golpe de suerte considerable. Las dos cosas a la vez.

Opta: el grupo correcto y una lección sobre actualizar de más

La supercomputadora de Opta simuló el torneo 25.000 veces antes del inicio y colocó a España primera con 16,1%, seguida por Francia (13,0%), Inglaterra (11,2%) y Argentina (10,4%). Esas cuatro fueron las semifinalistas.

Lo interesante vino después. Tras una fase de grupos impecable de Francia, la actualización de Opta del 28 de junio invirtió el orden: Francia primera con 18,7%, Argentina segunda con 16,3% y España relegada al tercer lugar con 13,5%. El modelo pre-torneo tenía razón; la actualización, movida por el rendimiento reciente, se equivocó. Es un caso de manual de sobreajuste a datos recientes, y lo veremos repetirse en el mercado.

EA Sports FC 26: cinco de cinco

El videojuego simuló el torneo completo y dio a España campeona. Con eso, EA acumula cinco Mundiales masculinos consecutivos acertando al campeón: 2010, 2014, 2018, 2022 y 2026. Es una racha extraordinaria y merece ser tomada en serio.

También merece el mismo escepticismo con el que miraríamos el caso de éxito de cualquier proveedor. EA publica su simulación del Mundial masculino con enorme despliegue, pero sus otros pronósticos fallan con normalidad: para la Eurocopa 2024 anunció una victoria de Inglaterra sobre Alemania en la final y ganó España; para el Mundial Femenino

2023 dio campeón a Estados Unidos frente a Alemania, y Estados Unidos quedó eliminado temprano —también ganó España—. Esos resultados circularon muchísimo menos. La racha en el torneo masculino es real. El recorte de la muestra, también.

La tarjeta de puntuación

Fuente	¿Nombró a España campeona?	¿Acertó la final?	Semifinalistas reales en su top 4	P(España) frente al azar
Goldman Sachs	Sí	Sí	3 de 4	12,5×
Silver Bulletin – PELE	Sí	No publicada	4 de 4, en orden exacto	No pública
Opta (pre-torneo)	Sí	No publicada	4 de 4	7,7×
Polymarket (pre-torneo)	No: segunda, a un punto	No aplica	3 de 4	7,7×
Univ. de Reading	No: Argentina	No	3 de 4	6,2×
Univ. de Portsmouth	No: Inglaterra	No	4 de 4	4,8×
Matillion Maia	Sí	No	No publicado	3,3×
EA Sports FC 26	Sí	No publicada	No publicado	No aplica
Joachim Klement	No: Países Bajos	No	0 de 4	No aplica

“P(España) frente al azar” divide la probabilidad asignada al campeón real por el 2,1% que correspondería a cada selección si el torneo se decidiera por sorteo.

4. Los que fallaron (y por qué el fallo es lo interesante)

Reading: un error que no fue tan grande

El modelo del economista James Reade, de la Universidad de Reading, estima la fuerza ofensiva y defensiva de cada selección a partir de los goles esperados de todos los partidos internacionales disputados desde enero de 2023, y corre 10.000 simulaciones del cuadro completo. El 8 de junio publicó a Argentina como favorita destacada con 24%, muy por encima de España (13%) y Francia (12%). Casi nadie más tenía a Argentina arriba, y menos a esa distancia.

Argentina no ganó, pero llegó a la final. El propio equipo de Reading publicó después del torneo un epílogo notablemente honesto: Argentina llegaba al partido decisivo con una ventaja mínima en Elo y terminó siendo claramente inferior a España. El modelo no confundió a un candidato con un equipo

mediocre: se equivocó en el margen más fino que existe. En términos de calidad de pronóstico, esa distinción importa mucho.

Portsmouth: la lógica correcta, la conclusión equivocada

El doctor Sarthak Mondal corrió un millón de simulaciones sobre una base de partidos mundialistas que se remonta a 1930 y publicó a Inglaterra como favorita con 15,9%, por delante de Argentina (10,9%), Francia (10,2%) y España (10,1%). Su razonamiento: el camino más probable de Inglaterra al título pasaba por definiciones desde el punto del penal, donde el modelo asumía una mejora sostenida del rendimiento inglés.

Inglaterra terminó tercera, su mejor resultado en sesenta años. Y el modelo tuvo a las cuatro semifinalistas reales en su top 4. Falló el orden, no

el diagnóstico. La distancia entre 15,9% y 10,1% es, en la práctica, indistinguible del ruido.

Klement: el oráculo que se quedó sin magia

Joachim Klement había acertado tres campeones consecutivos: Alemania 2014, Francia 2018 y Argentina 2022. Su modelo heurístico, que mezcla ranking FIFA con indicadores económicos y demográficos, proyectó para 2026 una final entre Países Bajos y Portugal, con los neerlandeses campeones. Países Bajos quedó eliminado en dieciseisavos por Marruecos, en definición por penales. Portugal cayó ante España en octavos.

El propio Klement había advertido, con humor, que quien apostara dinero siguiendo su pronóstico estaba más allá de toda ayuda. Conviene tomarlo literalmente. Tres aciertos consecutivos en un evento que ocurre cada cuatro años es exactamente el tipo de racha que el azar produce con regularidad, sobre todo si en cada edición se elige a alguno de los cinco o seis candidatos razonables. Una racha no es un método.

Los modelos de lenguaje: narrativa no es probabilidad

Aquí hay un dato incómodo y muy valioso para esta comisión. A fines de junio, terminada la fase de grupos y ya con el cuadro de eliminatorias definido,

varios medios consultaron a modelos de lenguaje generalistas. ChatGPT osciló entre Francia, Argentina y España. Gemini se inclinó por Argentina o Francia. Claude eligió a Francia, con el argumento de que Mbappé estaba en el pico de su carrera y de que Argentina dependía demasiado de un Messi de 38 años.

Ninguno acertó el campeón. Y los argumentos que usaron son reveladores: son razonamientos narrativos, plausibles, contruidos con el vocabulario del análisis deportivo. *Suenan* a análisis. No son estimaciones. Un modelo estadístico entrenado sobre resultados no “cree” que Mbappé esté en su mejor momento: calcula distribuciones a partir de miles de partidos.

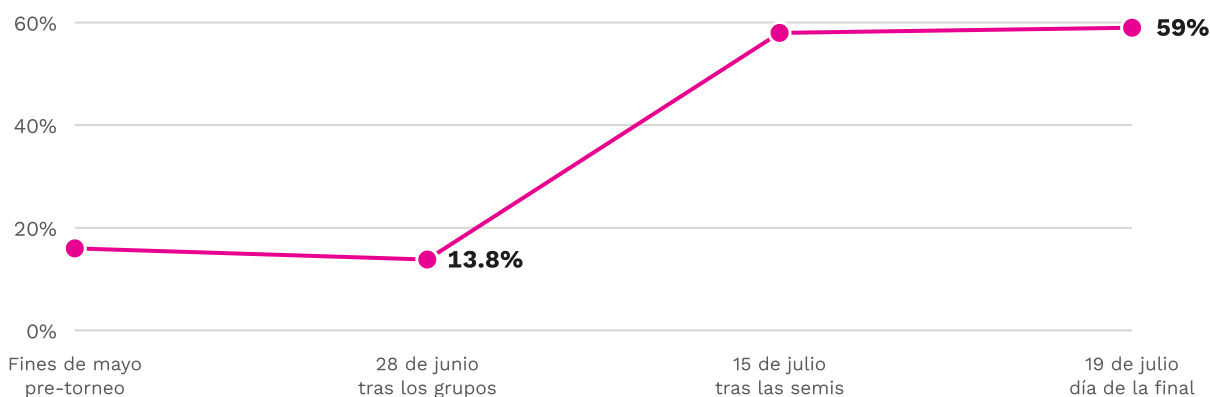
Hay un epílogo que conviene mirar de cerca. Ya en julio, con la final España-Argentina definida, varios medios volvieron a consultar a las IA y varias dieron a España ganadora por 1-0: el marcador exacto. El titular fue espectacular y el mérito, escaso. Para entonces quedaban dos equipos y el marcador más frecuente en una final cerrada es, justamente, 1-0. Eso no es pronosticar: es describir el escenario más común de un universo que ya se había reducido a dos opciones. Es el mismo truco que hace ver geniales a muchos tableros predictivos que, en realidad, solo informan lo que ya ocurrió.

5. Los mercados de predicción: rápidos, pero no clarividentes

Polymarket fue, por lejos, el termómetro más observado del torneo: los contratos vinculados al Mundial movieron miles de millones de dólares,

según los agregadores del sector. Su curva cuenta una historia que vale más que cualquier pronóstico aislado.

Probabilidad de España en Polymarket a lo largo del torneo



El mercado castigó a España tras el 0-0 del debut y solo la reconoció favorita cuando ya quedaban dos equipos en carrera.

El movimiento del 28 de junio es el más instructivo de todo el torneo. España debutó el 15 de junio con un 0-0 ante Cabo Verde, que jugaba su primer Mundial; Lamine Yamal, que venía de una lesión muscular sufrida el 22 de abril, entró recién en el segundo tiempo. El mercado castigó a España. Francia, en cambio, goleó 3-1 a Senegal con Mbappé encendido, y el mercado la premió.

Los dos movimientos resultaron equivocados: España fue campeona con la valla menos vencida de la historia y Francia terminó cuarta. Y hay un detalle que vuelve el episodio todavía más elocuente: en ese 0-0, España remató 27 veces contra 6 de Cabo Verde. El resultado dijo “empate”; el juego dijo otra cosa. Quien miró el marcador se equivocó; quien miró el volumen de juego, no.

Los mercados de predicción son extraordinarios agregando información. No son inmunes al sesgo de lo reciente: incorporan las noticias rápido, a veces más rápido de lo que la noticia merece.

Y hay un segundo punto, más sutil. En la mañana de la final, con dos equipos en cancha, el mercado daba 59% a España. Ganó España. ¿Fue un buen pronóstico? Sí, pero apenas: 59% contra 41% es poco más que una moneda levemente cargada. Un

mercado que dice 59% y acierta no demostró clarividencia; demostró que sabía que el partido estaba parejo. Y lo estaba: se definió recién a los 106 minutos.

6. Caos en el medio, orden en la cima

El formato de 48 selecciones agregó 104 partidos, una ronda nueva —los dieciseisavos de final— y la exigencia de encadenar cinco eliminaciones directas consecutivas. La hipótesis dominante antes del torneo era que ese diseño multiplicaría la aleatoriedad y devoraría a los favoritos. La aritmética la respaldaba: un equipo con 80% de probabilidad de ganar cada partido tiene apenas un 33% de ganar cinco seguidos.

La aleatoriedad efectivamente se multiplicó. Pero no donde se esperaba. La ronda nueva se llevó puestas a Alemania (penales ante Paraguay) y a Países Bajos (penales ante Marruecos). Brasil cayó en octavos ante Noruega y se quedó sin cuartos por primera vez desde 1990. Tres de los ocho cuartofinalistas — Noruega, Marruecos y Suiza— no aparecían en el top 8 de casi ningún modelo. Con 307 goles, fue además el Mundial más goleador de la historia.

Y sin embargo, arriba no pasó nada. Los cuatro semifinalistas fueron los cuatro favoritos pre-torneo de Opta y de PELE. El Mundial produjo más sorpresas que nunca en las rondas intermedias y ninguna en la cúspide.

La explicación no es misteriosa. Cada ronda adicional es un sorteo más, y cada sorteo castiga más a quien tiene menos margen. Alemania y Países

Bajos son equipos de élite con márgenes finos sobre sus rivales: les alcanza un mal día. España, Argentina, Francia e Inglaterra llegaban con márgenes más amplios en las primeras rondas. La entropía del formato se comió a los candidatos de segunda línea antes de llegar a los de primera.

Para quien modela, la lección es transferible: agregar etapas aumenta la varianza del resultado, pero no de manera uniforme. Aumenta más donde las diferencias de calidad son menores. Es lo mismo que ocurre en un embudo de conversión o en un test A/B con múltiples variantes: el ruido no golpea igual a todas.

Las tres explicaciones en disputa, resueltas

- **La ruta manda.** Sostenía que, con diferencias de calidad marginales entre potencias, el cuadro decide. *Parcialmente cierta:* Argentina tuvo el camino más despejado y llegó a la final. Pero España atravesó el cuadrante más duro — Portugal en octavos, Francia en semifinales— y ganó igual.
- **El plantel manda.** Sostenía que la calidad individual acumulada termina imponiéndose. *La más respaldada:* los cuatro planteles mejor valuados terminaron en los cuatro primeros puestos, y en el orden que anticipaba PELE.

- **El caos manda.** Sostenía que penales, expulsiones y VAR devorarían al favoritismo. Se *cumplió en el medio del cuadro, no en la cima:*

dos potencias murieron por penales en la ronda nueva, pero la final la ganó el favorito, en tiempo suplementario y sin llegar a los once metros.

7. Cómo se evalúa de verdad un pronóstico

Todo lo anterior es entretenido. Esta sección es la que importa para nuestro trabajo.

Un pronóstico no es una predicción

Cuando Goldman Sachs dice “España 26%”, no está diciendo “España va a ganar”. Está diciendo que en 100 mundiales como este, España ganaría 26. El mismo enunciado implica que en los otros 74 no gana. Que España haya ganado no prueba que el 26% fuera correcto, del mismo modo que si hubiera perdido no probaría que fuera incorrecto.

No es una sutileza académica. Es el error más común y más caro que se comete al leer un tablero de control: confundir la estimación con el resultado, y evaluar a quien estima por lo que pasó una sola vez.

Las reglas de puntuación

La disciplina que evalúa pronosticadores tiene herramientas específicas. La más usada es el **Brier score**, que promedia el cuadrado de la diferencia entre la probabilidad asignada y lo que efectivamente ocurrió (1 o 0). Va de 0 —pronóstico perfecto— a 1. Su primo, el **log score**, castiga mucho más duro a quien asigna probabilidad casi nula a algo que termina ocurriendo.

Ambas comparten una propiedad decisiva: son *reglas de puntuación propias*. El pronosticador maximiza su puntaje diciendo exactamente lo que cree, sin exagerar para llamar la atención ni suavizar para cubrirse. Es la única forma conocida de evitar que alguien juegue al titular.

Y ambas comparten un requisito: necesitan muchas observaciones. Un Brier score calculado sobre un único evento no dice absolutamente nada.

Por qué este ranking es provisorio

Todo lo que leyeron hasta acá —incluido el podio de este artículo— descansa sobre una sola observación. Con $n=1$ no se puede ranquear pronosticadores. Goldman Sachs pudo haber tenido el mejor modelo, o el más afortunado.

Por eso el hallazgo más sólido de esta auditoría no es el podio, sino la última columna de la tarjeta de puntuación: **todos los modelos serios superaron ampliamente al azar**. Entre 3,3 y 12,5 veces mejor que tirar un dardo sobre 48 equipos. Esa es la señal real. La diferencia entre 26% y 10% es la diferencia entre dos buenos modelos, no entre uno bueno y uno malo.

La forma correcta de evaluar a un pronosticador es su historial completo, publicado de antemano y puntuado con una regla propia. Cualquier otra cosa es contar la historia del ganador.

8. El dato que sí se dejó predecir

Hay una asimetría que conviene tener presente. Mientras los modelos peleaban por distinguir un 16% de un 13% en el resultado del torneo, las

proyecciones sobre el comportamiento de la audiencia se cumplieron con una precisión muchísimo mayor.

136 M

hinchas de fútbol en Norteamérica

+10,9%

de crecimiento en cinco años

76%

de la afición de EE.UU. es millennial o Gen Z

43%

de los hinchas norteamericanos son mujeres

El informe de Nielsen previo al Mundial cifraba en más de 136 millones los hinchas de fútbol en Norteamérica, un crecimiento del 10,9% en cinco años. Estados Unidos aporta 62,5 millones, la cuarta afición más grande del mundo. En México el fútbol es el deporte dominante, con 63% de interés declarado. Y el perfil es exactamente el que a nuestra industria le interesa: la edad promedio es de 33 años, las mujeres representan el 43% de los hinchas norteamericanos —contra 36% en Europa— y el 72% consume fútbol por televisión o *streaming*.

Ninguna de esas cifras es una adivinanza, y todas se sostuvieron. La diferencia con el pronóstico deportivo es la ley de los grandes números: estimar

qué hará una población de millones es muchísimo más fácil que estimar el desenlace de un evento único e irrepetible.

La bajada para el marketing es directa. La inteligencia artificial y el modelado predictivo rinden mucho mejor aplicados a agregados —qué porcentaje de una audiencia va a convertir, qué segmento va a rotar, cómo va a escalar una campaña— que cuando se les pide adivinar el desenlace de una negociación puntual o el éxito de un lanzamiento único. No es una limitación de la tecnología: es una propiedad de la realidad.

9. Siete lecciones para quien decide con datos

- 1. Una probabilidad no es una promesa.** Antes de discutir si el modelo “acertó”, verifiquen qué dijo exactamente. Un 26% que se cumple y un 26% que no se cumple son, ambos, resultados compatibles con un buen modelo.
- 2. Juzguen la trayectoria, no el titular.** Pidan el historial completo de pronósticos de cualquier proveedor, no la selección de los que salieron bien. Si no lo publica de antemano, no hay historial: hay marketing.
- 3. Pregunten siempre por la referencia.** “Nuestro modelo acierta el 70%” no significa nada si el evento ocurre naturalmente el 70% de las veces. La pregunta correcta es: ¿cuánto mejor que la línea de base?
- 4. Desconfíen de las rachas.** Klement llevaba tres de tres y falló. EA lleva cinco de cinco y puede seguir. En ninguno de los dos casos la racha, por sí sola, prueba que el método funcione.
- 5. Cuidado con sobrereaccionar a lo reciente.** Opta bajó a España después de un 0-0 y el mercado hizo lo mismo. Los dos se equivocaron. Su tablero semanal tiene el mismo problema cada vez que una campaña rinde mal siete días seguidos.
- 6. La convergencia entre modelos independientes vale más que cualquier modelo aislado.** Nueve metodologías distintas coincidieron en el mismo grupo de cinco o seis candidatos, y el campeón salió de ahí. Esa coincidencia era la señal; el ranking interno era casi ruido.
- 7. Un modelo de lenguaje generalista no es un modelo predictivo.** Para estimar probabilidades sobre datos cuantitativos, gana un modelo construido para eso. Usen la IA generativa para lo que sí hace bien: estructurar el problema, escribir el código, explicar el resultado. El caso de Matillion es ilustrativo: su IA no pronosticó el Mundial, *ayudó a construir en horas* el modelo estadístico que antes llevaba meses.

10. Cierre

Antes de que rodara la pelota teníamos nueve pronósticos serios y casi nueve respuestas distintas. Ahora tenemos una sola verdad y la tentación de repartir medallas.

Conviene resistirla. España ganó y Goldman Sachs lo había dicho con nombre, apellido y rival: es un resultado notable y merece el crédito. Pero el

hallazgo más útil de este ejercicio no es quién ganó la carrera de los pronósticos. Es que todos los que la corrieron con método llegaron cerca, y que la distancia entre ellos era, casi siempre, más chica que su propio margen de error.

En marketing tomamos decisiones con esa misma materia prima todos los días: estimaciones con incertidumbre, presentadas con más seguridad de la que tienen. El Mundial 2026 nos dejó un experimento natural inusualmente limpio para observar cómo se comportan esos modelos cuando la realidad finalmente llega. Vale la pena guardarlo.

Y para quienes prefieren la versión corta: el banco tenía razón, el mercado tenía razón, el videojuego tenía razón. Ninguno estaba seguro. Ese es, exactamente, el punto.

Nota metodológica y fuentes

Todas las probabilidades citadas corresponden a publicaciones públicas de cada entidad, fechadas antes del partido inaugural salvo donde se indica lo contrario. La columna “P(España) frente al azar” se calcula dividiendo la probabilidad asignada a España por 2,083% (1/48). Los pronósticos determinísticos —EA Sports, Klement— no publican distribuciones de probabilidad y por eso no admiten esa comparación.

RESULTADOS DEL TORNEO

FIFA — España 1-0 Argentina, crónica de la final
ESPN — Spain edge Argentina after extra time to win 2nd World Cup title
FIFA — Rodri, Cubarsí, Simón y Mbappé, ganadores de los premios individuales
NBC Sports — Lista completa de premios del Mundial 2026
ESPN — Mbappé supera a Messi como máximo goleador histórico
ESPN — Argentina remonta ante Inglaterra en la semifinal
FIFA — España 0-0 Cabo Verde, debut en la fase de grupos
Al Jazeera — Las mayores sorpresas de las eliminatorias

Resultados, marcadores y premios individuales se verificaron contra fuentes primarias —FIFA y las publicaciones de las propias instituciones que emitieron cada pronóstico— y contra medios de referencia. La aclaración no es un formalismo: la cobertura secundaria de premios y probabilidades contiene errores con frecuencia, y buena parte de lo que circula sobre “qué modelo acertó” proviene de repeticiones no verificadas. Es, en pequeño, exactamente el problema que este artículo intenta describir.

MODELOS Y PRONÓSTICOS

Bloomberg — El modelo de Goldman Sachs da 26% a España
CNBC (2018) — Goldman Sachs proyectaba a Brasil venciendo a Alemania
Opta Analyst — Predicciones pre-torneo del Opta Supercomputer
Silver Bulletin — Pronósticos del modelo PELE
Universidad de Reading — Argentina favorita según la supercomputadora
Economics Observatory — Epílogo del equipo de Reading tras el torneo
Universidad de Portsmouth — Inglaterra como favorita
Matillion — Cómo se construyó el modelo Maia en horas
ESPN — Klement elige Países Bajos
Sports Illustrated — Qué predijeron ChatGPT, Gemini y Claude
Polymarket — Mercado “World Cup Winner”
Nielsen — The Fans Behind The Game: FIFA World Cup 2026 Edition

◆ Enlaces completos a todas las fuentes en ia.dma.org.ar

◆ AUTOMATIZACIÓN Y AGENTES

Agentes de IA en marketing

Casos reales, señales claras y por dónde empezar. Durante años la IA en marketing fue asistencia; los agentes autónomos cambian la arquitectura: reciben un objetivo y ejecutan de principio a fin.

María Laura Barreto

Doppler

02

Durante algún tiempo, la inteligencia artificial en marketing fue sinónimo de asistencia. Útil, sin dudas, pero requería que una persona tomara las riendas en cada paso, validara cada salida, tomara cada decisión. La IA hacía el borrador; el humano, el trabajo real.

Ese modelo no desapareció, pero está siendo desplazado. Los agentes de IA autónomos no esperan instrucciones para cada acción: reciben un objetivo, evalúan un contexto, acceden a las

herramientas disponibles y ejecutan de principio a fin. No es el mismo modelo que conocíamos con otro nombre: es una arquitectura diferente, con implicancias distintas para los equipos de marketing.

Lo interesante es que esto ya no es un escenario futuro. Está pasando ahora, en organizaciones reales y con resultados medibles. Entender qué está ocurriendo —y qué no— es clave para no quedar fuera de la conversación.

¿Qué es exactamente un agente de IA?

La confusión terminológica no ayuda. Un chatbot no es un agente. Un modelo que genera texto, tampoco. Un agente de IA es un sistema que puede planificar, decidir y actuar de forma autónoma para alcanzar un resultado definido. La diferencia está en el nivel de autonomía: tiene acceso a herramientas — navegadores, APIs, bases de datos, plataformas— y puede encadenar múltiples acciones sin que nadie le indique cada paso.

Una analogía que ayuda: contratar a alguien para que escriba un email es distinto de contratar a alguien para que gestione toda la estrategia de email marketing. El primero ejecuta una tarea. El segundo

toma decisiones, ajusta según los resultados y actúa dentro de un marco de objetivos más amplio. Los agentes funcionan más parecido al segundo caso.

En la práctica, esto significa que un equipo de marketing puede darle a un agente un objetivo como “lanzar la campaña de reactivación para contactos inactivos del último trimestre”, y el sistema se encargará de segmentar la audiencia, armar el email, seleccionar el horario óptimo, configurar el flujo de seguimiento y disparar la campaña, tomando decisiones intermedias en base a los datos disponibles. Sin que nadie tenga que intervenir entre pasos.

Casos de uso concretos

Más allá de la teoría, vale la pena detenerse en qué está pasando hoy. No en laboratorios, sino en equipos de marketing con objetivos y presupuestos reales.

Generación y distribución de contenido

Los agentes pueden encargarse de redactar, adaptar y publicar contenido en múltiples canales, coordinando calendarios editoriales. No solo generan el texto: verifican el tono de marca, ajustan el formato según el canal y programan los tiempos de publicación. Lo que antes requería la coordinación de tres roles distintos empieza a poder orquestarse desde un único agente con acceso a las herramientas correctas.

POR DÓNDE EMPEZAR Mapeá un proceso editorial repetitivo —los posts de redes de una campaña, las novedades de email marketing semanales— y evaluá qué pasos son puramente de ejecución. Esos son los candidatos para delegar primero.

Email, automatización y segmentación dinámica

Los flujos de automatización ya eran potentes. Con agentes, se vuelven adaptativos. En lugar de seguir una secuencia fija, el sistema puede rediseñar el recorrido del suscriptor en tiempo real según su comportamiento, ajustando asuntos, frecuencias y variantes sin que nadie tenga que intervenir manualmente.

Lo mismo aplica a la segmentación: los agentes pueden mantener los segmentos de contactos actualizados combinando señales de

comportamiento, datos de compra y actividad en múltiples canales, algo que antes exigía trabajo manual constante o integraciones complejas.

POR DÓNDE EMPEZAR Revisá cuántos segmentos de tu base se actualizan manualmente o con criterios estáticos. Si la respuesta es “la mayoría”, ese es el punto de entrada más directo: automatizar la lógica de segmentación libera tiempo y mejora la relevancia de cada campaña de inmediato.

Optimización de presupuesto en tiempo real

Un agente puede monitorear el rendimiento de una campaña y redistribuir el presupuesto entre canales, creativos y audiencias de forma autónoma. Si los anuncios en un canal convierten al doble que en otro, el sistema reasigna sin esperar la reunión de seguimiento semanal. Es la diferencia entre reaccionar y anticipar, y entre reportar lo que pasó y actuar mientras pasa.

POR DÓNDE EMPEZAR Antes de automatizar decisiones de presupuesto, el prerequisite es tener las métricas de conversión unificadas en un solo lugar. Si los datos de rendimiento viven en silos —una plataforma por canal, sin visión integrada—, ningún agente puede tomar buenas decisiones. Resolvé eso primero.

Investigación competitiva continua

Agentes que monitorean precios, mensajes y lanzamientos de la competencia de forma permanente y generan reportes accionables sin que nadie se lo pida. Lo que antes era una tarea trimestral que alguien hacía manualmente puede convertirse en un flujo continuo de inteligencia competitiva.

POR DÓNDE EMPEZAR Definí explícitamente qué competidores querés monitorear y qué señales importan —precios, textos de anuncios, lanzamientos de producto, cambios en la web—. Sin ese criterio claro, el agente va a generar volumen de información, pero no inteligencia útil.

Hacia dónde vamos

La tendencia apunta hacia equipos de marketing más pequeños operando con mayor escala. No porque los equipos grandes vayan a desaparecer, sino porque la proporción entre personas y tareas ejecutadas está cambiando estructuralmente. Un equipo que antes necesitaba diez personas para gestionar cinco canales puede empezar a operar de forma diferente cuando parte de esa coordinación se delega a agentes.

El desafío real no es tecnológico. Es organizacional: aprender a trabajar con agentes, definir qué procesos se pueden delegar de forma segura y construir los sistemas de control para que esa delegación sea predecible. Quienes empiecen a desarrollar esa capacidad ahora van a tener una ventaja difícil de alcanzar después.

Hay señales claras de hacia dónde se mueve la industria. Algunos casos concretos que vale la pena seguir de cerca:

A**Comercio agéntico****El consumidor que nunca llega a tu sitio**

Los consumidores están empezando a delegar sus compras a agentes de IA: le piden a ChatGPT, Gemini o Perplexity que comparen productos, evalúen opciones y, en algunos casos, directamente ejecuten la transacción. Para los equipos de marketing, la pregunta ya no es solo cómo aparece la marca en Google. Es cómo la percibe y la representa un agente de terceros que toma decisiones en nombre del consumidor.

B**AEO****La próxima batalla del SEO ya empezó**

El AEO —*Answer Engine Optimization*— es la práctica de estructurar el contenido para que los motores de respuesta impulsados por IA lo citen como fuente confiable. No se trata de rankear: se trata de ser la marca que el modelo menciona cuando alguien pregunta. Los primeros en establecer esa presencia van a tener una ventaja que será muy difícil de alcanzar después.

C**Agentes embebidos****El que ya está en la herramienta que usás hoy**

Los CRM, las plataformas de email, los sistemas de automatización y los gestores de contenido están integrando agentes directamente en sus flujos de trabajo. Cuando el agente llega embebido en una plataforma ya adoptada, la pregunta deja de ser “¿deberíamos explorar agentes?” para convertirse en “¿estamos aprovechando lo que ya tenemos?”. Ese cambio de pregunta suele ser el que acelera la adopción real.

Tres preguntas para empezar hoy

1. ¿Qué tareas de mi equipo son puro volumen de ejecución? No las más estratégicas ni las más creativas: las que se repiten, siguen una lógica predecible y consumen tiempo desproporcionado. Esos son los candidatos más claros para delegar a un agente.

2. ¿Qué tan consolidados están los datos con los que trabajo? Los agentes amplifican lo que tienen disponible. Si los datos de contactos, comportamiento y rendimiento están fragmentados, el agente va a tomar decisiones fragmentadas. La calidad del dato es el prerequisite de cualquier automatización seria.

3. ¿Qué necesita aprobación humana antes de ejecutarse? No todo tiene que ser autónomo del todo. Definir los puntos de control —qué puede ejecutar el agente solo y qué necesita que un humano apruebe— es lo que hace que la delegación sea sostenible y no una fuente de errores costosos.

Los equipos que van a sacar ventaja de esta tecnología no son necesariamente los que tienen más presupuesto ni los que adoptan más rápido. Son los que empiezan con claridad: saben qué quieren delegar, tienen los datos para respaldarlo y construyen los controles para que funcione bien. Eso no requiere un proyecto de seis meses. Requiere hacerse las preguntas correctas esta semana.

◆ SUPERVISIÓN Y CRITERIO

Human-in-the-loop

El equilibrio entre escala y criterio. Automatizar no significa delegar todo: la IA ejecuta, propone y automatiza; la persona supervisa, valida y aporta criterio.

Natalia Borrazás

Growlat

03

Automatizar no significa delegar todo

La inteligencia artificial está transformando la forma en que trabajamos. Hoy puede analizar datos, redactar contenido, segmentar audiencias, responder consultas y ejecutar tareas que hace poco requerían intervención humana constante. Frente a esta velocidad de evolución, muchas organizaciones se hacen la misma pregunta: si la IA puede hacerlo sola, ¿por qué seguir supervisándola?

La respuesta es simple: porque hacer algo rápido no siempre significa hacerlo bien. La IA puede procesar enormes volúmenes de información y automatizar decisiones en segundos. Pero todavía tiene limitaciones importantes cuando necesita interpretar contexto, comprender matices culturales o anticipar consecuencias estratégicas.

Ahí entra un concepto que se está volviendo central en el debate sobre adopción responsable de IA: *human-in-the-loop*.

¿Qué significa human-in-the-loop?

Human-in-the-loop (HITL) es un modelo de trabajo colaborativo entre personas y sistemas de inteligencia artificial. La lógica es simple: la IA ejecuta, propone o automatiza; la persona supervisa,

valida y aporta criterio. No se trata de frenar la automatización, sino de asegurarse de que las decisiones importantes sigan teniendo contexto humano.

El verdadero potencial de la inteligencia artificial no está en reemplazar criterio humano. Está en amplificarlo.

Un espectro, no una decisión binaria

Este modelo existe en un espectro continuo. No hay una única forma de aplicarlo: el nivel de supervisión debe calibrarse según el tipo de tarea, el costo del

error y la madurez del sistema.

01

IA asistente

La IA propone, el humano decide y ejecuta.

Para tareas nuevas o de alto impacto.

02

Revisión selectiva

La IA ejecuta en casos estándar y escala al humano en excepciones.

Cuando la confianza del modelo es baja.

03

Auditoría periódica

La IA opera de forma autónoma. El humano audita resultados.

Y ajusta el modelo con lo aprendido.

El grado de autonomía crece de 01 a 03; la intensidad del color acompaña esa progresión.

¿Por qué sigue siendo necesario?

La IA es poderosa, pero no infalible. Puede equivocarse, puede interpretar mal una instrucción, puede ofrecer respuestas técnicamente correctas pero completamente desacertadas en contexto.

EJEMPLO Un sistema puede redactar una campaña gramaticalmente impecable, pero desconectada del tono real de la marca. Puede segmentar audiencias de forma eficiente, pero sin detectar sensibilidad cultural o reputacional. Puede recomendar acciones optimizadas para métricas de corto plazo que dañen vínculos de largo plazo con clientes.

La IA detecta patrones. Las personas interpretan significado. Y esa diferencia importa.

◆ Guías de adopción responsable de IA en ia.dma.org.ar

◆ ECONOMÍA DE LA IA

El ROI real de la inteligencia artificial

Por qué comparar licencias con salarios es la peor forma de medir el impacto de la IA en la empresa. Seis casos medidos que contrastan tres modelos operativos.

Fernando Cuscuela

Co-fundador de Wego Digital Hub

04

En el último año, prácticamente todas las conversaciones de dirección sobre inteligencia artificial empiezan en el mismo lugar: el precio mensual de una licencia de ChatGPT, Claude o Gemini, comparado con el costo cargado de un profesional senior, parece anunciar una revolución inmediata. La realidad, sin embargo, es más

interesante y más exigente. La IA está reescribiendo la economía de la productividad, pero el verdadero retorno no se encuentra en la diferencia entre una licencia y un salario, sino en el rediseño de los procesos que conectan a las personas con las máquinas.

Una nueva economía de la productividad

Durante décadas, la productividad del trabajo intelectual estuvo limitada por la velocidad con que una persona podía leer, escribir, analizar y decidir. La IA generativa rompe ese techo. Hoy, un mismo profesional puede investigar un mercado, redactar un informe, generar variantes creativas y validar hipótesis en una fracción del tiempo que requería hace tres años. La materia prima ya no es el tiempo, sino el criterio: qué preguntar, qué validar, qué descartar.

Las organizaciones más competitivas serán aquellas capaces de combinar talento humano e inteligencia artificial en flujos diseñados a propósito. La pregunta estratégica deja de ser “¿reemplazo personas con IA?” y pasa a ser “¿qué capacidades amplifico, qué capacidades protejo y qué procesos rediseño?”. En esa transición, medir mal el retorno es la forma más rápida de tomar decisiones equivocadas.

El error de comparar licencias con salarios

La comparación más extendida en los comités ejecutivos es también la más engañosa. Tomar el costo mensual de una licencia de IA y enfrentarlo al salario de un profesional sugiere ahorros de dos órdenes de magnitud. Pero el costo verdadero de un profesional incluye experiencia acumulada, juicio contextual, responsabilidad legal y capacidad de coordinar con otros humanos. El costo verdadero de una herramienta de IA, a su vez, va mucho más allá del abono mensual.

Un análisis correcto debe contemplar al menos cinco capas adicionales: capacitación del equipo, seguridad de la información, supervisión y validación de proyectos, gestión del cambio cultural y gobierno de datos. Cuando estos costos se incorporan, la ecuación deja de ser “IA versus humano” y se convierte en “qué combinación maximiza valor por unidad de tiempo y riesgo”.

El ahorro no nace de reemplazar al profesional con un modelo; nace de rediseñar el proceso para que el profesional pueda producir el doble.

Los costos ocultos que nadie presupuesta

Toda implementación seria de IA tiene un iceberg debajo de la línea de flotación. La capacitación es el primer costo o, mejor dicho, la mejor inversión: no basta con contratar las licencias, hay que enseñar a usar los modelos con criterio, con instrucciones efectivas y con expectativas calibradas. La seguridad es el segundo: información confidencial, propiedad intelectual y datos personales requieren acuerdos

contractuales, controles técnicos y políticas de uso. La adopción cultural es el tercer costo y, paradójicamente, el más subestimado. Implementar IA en una organización que no la entiende ni la integra a sus rituales de trabajo genera más fricción que productividad.

El cuarto costo es el control de calidad. Las herramientas generativas producen resultados verosímiles incluso cuando son incorrectos. Sin un proceso de validación humana, los errores se acumulan en silencio. Y el quinto, transversal a los anteriores, es el gobierno de datos: trazabilidad,

versionado de instrucciones, registro de decisiones y auditoría. Sumados, estos rubros representan una parte significativa del costo total de propiedad y son la diferencia entre un piloto exitoso y una iniciativa estancada.

ChatGPT, Claude y Gemini: tres motores complementarios

El mercado tiende a presentar a los tres grandes asistentes como sustitutos. La práctica demuestra que son herramientas con personalidades distintas. ChatGPT brilla por la amplitud de su ecosistema y por la madurez de sus integraciones, lo que lo convierte en una plataforma de propósito general particularmente fuerte cuando se requieren complementos, agentes y conexiones con sistemas externos. Claude se distingue por su capacidad de programación y de análisis profundo de documentos largos, por la claridad de su redacción y por la prudencia de sus respuestas en contextos sensibles.

Gemini, finalmente, se beneficia de su integración nativa con Google Workspace y de un acceso fluido a la búsqueda en tiempo real.

En la mayoría de los proyectos que hemos observado, el mayor valor no surge de elegir un único modelo, sino de orquestarlos. Investigar con Claude, estructurar con ChatGPT y validar con Gemini suele producir mejores resultados que insistir en una sola herramienta. Según lo que hemos observado en nuestras implementaciones, quien decide por la estandarización absoluta tiende a resignar alrededor del 30% del beneficio potencial de la IA.

La paradoja de la productividad

Hay una trampa frecuente en la adopción de IA: aplicarla sobre procesos mal diseñados. Automatizar un flujo ineficiente solo permite cometer errores más rápido. Si una empresa produce contenidos sin un plan editorial, sumar IA acelerará la producción de contenidos irrelevantes. Si la atención al cliente carece de criterios de resolución, sumar un bot duplicará el ruido. La IA no compensa la falta de diseño de procesos: lo expone.

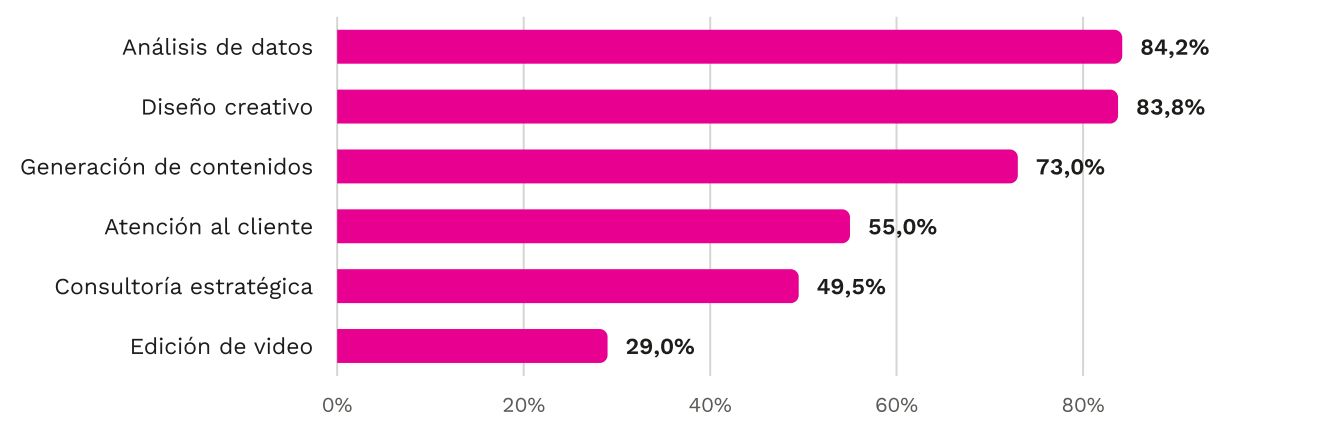
Captar valor real exige rediseñar el flujo antes de automatizarlo. Mapear etapas, identificar dónde el juicio humano es irremplazable, dónde la generación masiva agrega valor y dónde la validación humana es indispensable para evitar errores costosos. El rediseño es el paso que la mayoría de las organizaciones saltea y, también, el que explica por qué pilotos prometedores no se convierten en transformaciones sostenidas.

El modelo híbrido: humano más IA

La evidencia empírica de los últimos dos años es consistente: los mejores resultados aparecen cuando las personas utilizan la IA como amplificador de capacidades, no como sustituto. Los profesionales que combinan su criterio con la velocidad de los modelos producen mejor calidad, en menos tiempo y a menor costo que cualquiera de los extremos. El modelo híbrido es, hoy, la frontera de eficiencia para casi todas las funciones del trabajo intelectual.

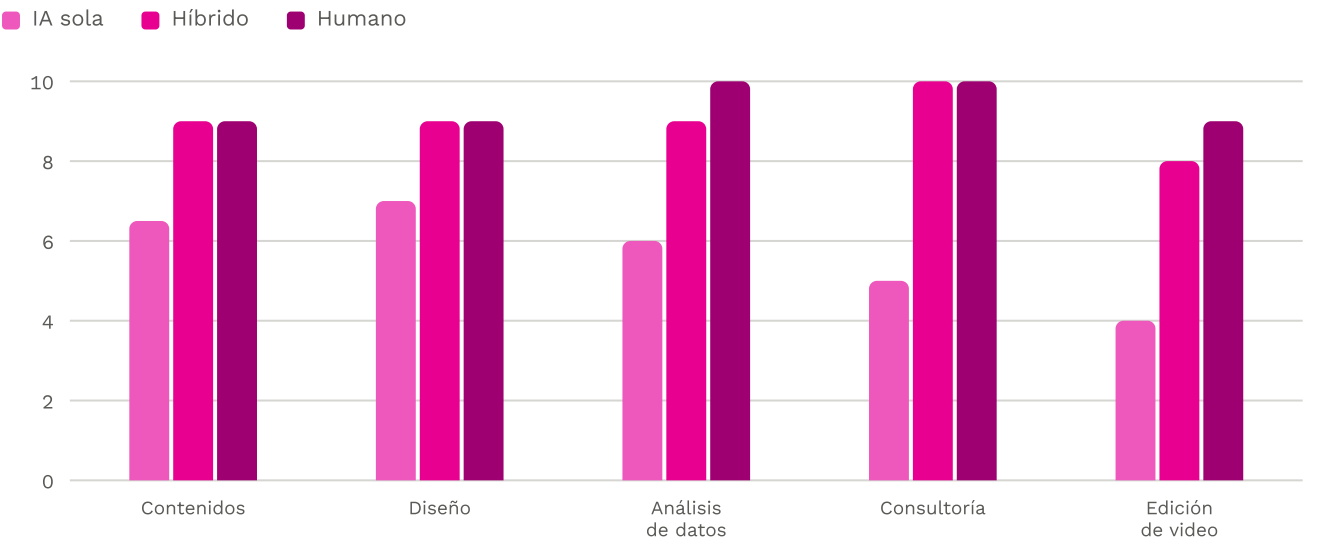
Los casos que siguen comparan los tres modelos operativos sobre escenarios reales de marketing, diseño, atención al cliente, análisis de datos, consultoría y producción audiovisual. Los números no son proyecciones: replican implementaciones documentadas por el grupo empresarial Wego Digital Hub, socio de la DMA y participante activo de la Comisión de Inteligencia Artificial.

Ahorro del modelo híbrido frente al modelo 100% humano



El ahorro se sostiene en los seis dominios, pero su magnitud cae a medida que el juicio humano pesa más en el resultado.

Calidad del resultado por modelo, sobre 10



La calidad del modelo híbrido iguala o casi iguala a la del humano puro en los cinco dominios puntuados sobre 10; la IA sin supervisión cae siempre. Atención al cliente se mide en tasa de resolución y se detalla en el Caso 3.

Seis casos comparativos

Caso 1 — Generación de contenidos

Escenario: Producción de cuatro artículos mensuales de 1.500 palabras.			
Dimensión	Modelo humano	Modelo IA	Modelo híbrido
Tiempo	16 hs / mes	2 hs / mes	6 hs / mes
Costo	USD 240	USD 5	USD 65
Calidad / resultado	9 / 10	6,5 / 10	9 / 10

En este flujo, Claude lidera la investigación y la síntesis de fuentes, ChatGPT estructura el borrador y propone variantes de titulares, Gemini valida datos y enlaces externos, y un editor humano refina tono, marca y narrativa. El resultado: reducción del 73%

del costo y del 62% del tiempo, manteniendo la calidad profesional. La calidad estable es lo que separa al híbrido del modelo “solo IA”, que cae a 6,5/10 sin supervisión y termina exigiendo reescritura.

Caso 2 — Diseño creativo para campañas digitales

Escenario: Campaña para Meta Ads y Google Ads: cinco piezas estáticas, tres variantes para testeo A/B, dos banners display y adaptaciones para formatos *story* y *feed*.

Dimensión	Modelo humano	Modelo IA	Modelo híbrido
Tiempo	8 hs	45 min	2 hs
Costo	USD 320	USD 2	USD 52
Calidad / resultado	9 / 10	7 / 10	9 / 10

Las herramientas generativas aceleran la exploración: en minutos se generan decenas de variantes que un diseñador hubiese tardado horas en bocetar. El refinamiento final, sin embargo, requiere ojo humano para la coherencia de marca, la

jerarquía visual y la adaptación a cada formato. El ahorro alcanza el 83,75%, pero, más relevante aún, el equipo creativo dedica su tiempo a lo que solo un humano puede hacer: decidir.

Caso 3 — Atención al cliente

Escenario: 1.000 consultas mensuales.

Dimensión	Modelo humano	Modelo IA	Modelo híbrido
Recursos	2 agentes	Agente conversacional	IA + humanos
Costo	USD 2.000	USD 300	USD 900
Calidad / resultado	85% de resolución	70% de resolución	92% de resolución

Este es el caso donde más claramente se vuelve evidente que la IA sola no basta. Un agente conversacional reduce el costo a la sexta parte, pero también baja la tasa de resolución y deteriora la experiencia. El modelo híbrido, en cambio, asigna a

la IA las consultas frecuentes y bien estructuradas, y deja a las personas los casos complejos o emocionalmente sensibles. El ahorro es del 55% y, paradójicamente, la calidad del servicio mejora frente al modelo puramente humano.

Caso 4 — Análisis de datos

Escenario: Reporte ejecutivo semanal.

Dimensión	Modelo humano	Modelo IA	Modelo híbrido
Tiempo	8 hs	45 min	2 hs
Costo	USD 560	USD 1	USD 88,50
Calidad / resultado	10 / 10	6 / 10	9 / 10

Los modelos conversacionales son extraordinarios para limpiar datos, generar tablas y proponer visualizaciones. Donde tropiezan es en la interpretación: distinguir un movimiento estacional de un cambio estructural, entender el contexto

comercial detrás de un número o anticipar la pregunta que hará el directorio. Cuando el analista valida y narra el análisis preliminar de la IA, el ahorro llega al 84,2% sin sacrificar el criterio que sostiene cada decisión.

Caso 5 — Consultoría estratégica

Escenario: Diseño de una estrategia de transformación digital.			
Dimensión	Modelo humano	Modelo IA	Modelo híbrido
Tiempo	20 hs	2 hs	12 hs
Costo	USD 2.000	USD 10	USD 1.010
Calidad / resultado	10 / 10	5 / 10	10 / 10
En este caso el modelo solo con IA produce análisis genéricos, sin contexto de cliente ni profundidad sectorial. El consultor humano, en cambio, alcanza la calidad máxima, pero a un costo difícil de escalar. El híbrido toma lo mejor de ambos: el consultor		delega a la IA la investigación, los <i>benchmarks</i> y la primera estructuración, y dedica su tiempo a la conversación con el cliente, al diagnóstico y a la decisión. El ahorro es del 49,5% manteniendo el rigor estratégico.	

Caso 6 — Edición avanzada de video

Escenario: producción de un video corporativo de 5 minutos a partir de 2 horas de material grabado, con montaje, <i>motion graphics</i> , corrección de color, musicalización y subtítulos profesionales.			
Dimensión	Modelo humano	Modelo IA	Modelo híbrido
Tiempo	16 hs	~12 hs de iteración	10 hs
Costo	USD 960	USD 1.400	USD 680
Calidad / resultado	9 / 10	4 / 10	8 / 10

Único caso de los seis en el que el modelo solo-IA resulta más caro que el humano puro: un 46% por encima, y con calidad inferior.

Este caso rompe el patrón de los anteriores y es, precisamente por eso, el más instructivo. La edición avanzada de video implica procesar grandes volúmenes de datos visuales en múltiples ciclos de iteración. El costo del modelo solo-IA se descompone en tres capas:

- 1. Análisis de material con modelos multimodales.**
Procesar dos horas de material bruto para que el modelo identifique los mejores cortes consume entre 4 y 8 millones de *tokens* solo en la fase de lectura (USD 60 a 120 en APIs de última generación). Cada ronda de refinamiento multiplica ese costo.
- 2. Creación de piezas audiovisuales mediante IA.**
Herramientas como Runway Gen-3 o Sora cobran por segundo de video generado. Producir 60 segundos de contenido utilizable puede requerir de 15 a 25 iteraciones —cada una fallida antes de encontrar la correcta—, con un costo acumulado de USD 400 a 600 solo en créditos de generación.

3. Suscripciones y herramientas auxiliares.
Corrección de color con IA, generación musical (Suno, Udio) y subtítulo automático de calidad profesional suman entre USD 80 y USD 150 mensuales adicionales.

El resultado no justifica la inversión: los modelos actuales pueden sugerir cortes y generar imágenes genéricas, pero no pueden leer la narrativa emocional de una pieza, respetar la identidad visual de una marca ni tomar las decisiones estéticas que diferencian un video corporativo de una producción correcta desde lo técnico, pero sin capacidad de diferenciación. La calidad del modelo solo-IA baja a 4/10 y el material producido exige tantas correcciones que resulta más barato y más rápido comenzar desde cero con el profesional.

El modelo híbrido recupera la lógica de los casos anteriores, pero con un diseño más selectivo: el editor senior usa IA únicamente en las tareas de bajo consumo de *tokens* —transcripción automática, montaje preliminar sobre el guion, sugerencias de música y subtítulos— y mantiene el control creativo completo sobre el montaje, los *motion graphics* y la

corrección de color. El resultado es un ahorro del 29% sobre el modelo humano puro sin sacrificar la calidad de la pieza final.

La lección es estructural: el consumo de *tokens* y de créditos de generación de video escala de forma no lineal con la complejidad del proyecto. Antes de

decidir usar IA para producción audiovisual avanzada, el análisis de ROI debe incluir una estimación realista del costo por iteración y del número de ciclos necesarios para alcanzar calidad profesional. En este dominio, la IA todavía no desplaza al editor senior: lo asiste en tareas puntuales.

En cinco de los seis casos, el modelo híbrido superó al humano puro en costo y al solo IA en calidad. En el sexto —edición avanzada de video— la IA resultó un 46% más cara que el profesional senior y con calidad inferior: una advertencia sobre los costos no lineales de la generación de video y los límites actuales de la automatización audiovisual.

La organización de 2030

Si proyectamos estos hallazgos hacia el futuro inmediato, lo que aparece no son empresas con menos personas, sino con personas distintas. Equipos más livianos, con mayor productividad individual, organizados alrededor de capacidades en lugar de jerarquías funcionales. Cada profesional pasará a coordinar un pequeño ecosistema de

agentes y modelos que amplifican su trabajo. Las funciones administrativas se concentrarán; las funciones de juicio y relación se multiplicarán.

Las organizaciones que liderarán esta transición invertirán hoy en tres frentes simultáneos: gobierno de datos, alfabetización en IA y rediseño de procesos. Sin gobierno no hay confianza, sin alfabetización no hay adopción y sin rediseño no hay retorno. Los tres son requisitos, no opciones.

Conclusión

Comparar licencias con salarios es la peor forma de medir el impacto de la IA. La verdadera ecuación incluye costos ocultos, riesgos de calidad, rediseño de procesos y, sobre todo, la decisión estratégica sobre dónde el juicio humano agrega valor irremplazable. La evidencia muestra que el modelo híbrido es hoy la frontera de eficiencia para casi todos los flujos del trabajo intelectual: mejor calidad que la IA sola, menor costo que el humano solo y, en muchos casos, mejor experiencia para el cliente final. El caso de la edición avanzada de video añade un matiz necesario: existen dominios donde el

consumo de *tokens* escala hasta superar el costo del profesional, y donde la calidad de la IA aún no justifica la inversión. Identificar esos límites con precisión es tan estratégico como identificar las oportunidades.

La pregunta para cada directorio ya no es si adoptar IA, sino con qué nivel de intencionalidad rediseñar el trabajo para que esa adopción produzca el retorno que promete. El ROI real de la IA se decide en el cruce entre proceso, herramienta y juicio humano. Es allí donde la Comisión de Inteligencia Artificial de la DMA propone seguir conversando.

Sobre la Comisión de Inteligencia Artificial

La IA avanza más rápido de lo que la industria puede procesar sola. Esa es la convicción de la que nace esta Comisión.

La Comisión de Inteligencia Artificial de la DMA surge de una convicción simple: la IA avanza más rápido de lo que la industria puede procesar sola. Recorrer lo que está pasando en el mundo, debatirlo con criterio, compartirlo con claridad y promover un uso responsable y legal de los datos

no es un ejercicio académico. Es una necesidad práctica para cualquier profesional, empresa o institución que quiera tomar decisiones informadas.

Este informe es una de las formas concretas en que esa convicción se traduce en acción.

A quiénes nos dirigimos

Empresas, agencias, profesionales, proveedores, estudiantes, medios y toda la comunidad regional de Iberoamérica.

Nuestro propósito

Informar, inspirar, promover, capacitar y conectar a profesionales de la industria en el uso con impacto de la inteligencia artificial.

Cómo trabajamos

Sobre tres ejes: *forecast* regional de IA, educación por niveles y verticales, y networking entre profesionales.

Integrantes de la Comisión

Hernán Litvac

Icomm

Ricardo Díez

Analytics Town

Fernando Cuscuela

Wego Digital Hub

María Laura Barreto

Doppler

Lionel Pieniazek

Shake Again

Natalia Borrazás

Growlat

Diego Amoros

iMetrics

Ariel Leiro

Data4Sales

Sebastián Antonini

Diecisiesonce



◆ MICROSITIO OFICIAL DE LA COMISIÓN

ia.dma.org.ar

Ediciones anteriores del informe, enlaces completos a todas las fuentes citadas, recursos de capacitación, agenda de actividades y novedades de la Comisión de Inteligencia Artificial de la DMA Argentina.

Este informe fue elaborado por los integrantes de la Comisión de Inteligencia Artificial de la DMA Argentina. Los datos citados corresponden a fuentes públicas verificadas a julio de 2026. Se permite su reproducción total o parcial citando la fuente.