

Conexiones humanísticas

por Frank Alberto Betances Reinoso

Inteligencia artificial en Otorrinolaringología: Desafíos éticos para una implementación responsable.

Introducción

La inteligencia artificial (IA) representa actualmente una de las transformaciones más disruptivas en el ámbito médico, con aplicaciones especialmente prometedoras en Otorrinolaringología. Desde la interpretación automatizada de tomografías computarizadas hasta el apoyo en diagnósticos diferenciales de patologías complejas, la IA está redefiniendo los paradigmas clínicos tradicionales.

No obstante, esta revolución tecnológica plantea desafíos éticos de gran calado que trascienden las meras consideraciones técnicas, interpelando directamente los fundamentos deontológicos y humanísticos de la práctica médica contemporánea.

El primer requisito ético de cualquier innovación científica es no engañar. Ofrecer promesas sin evidencia sólida o no comunicar adecuadamente la inmadurez de la tecnología erosiona la confianza social y profesional en la ciencia médica. En el contexto otorrinolaringológico, donde las decisiones clínicas afectan funciones vitales como la audición, el habla y la respiración, estos dilemas éticos adquieren una dimensión particularmente crítica.

Inmadurez tecnológica y ausencia de evidencia clínica sólida

La realidad detrás de las promesas

La implementación de soluciones de IA médica supone desafíos críticos relacionados con la validación clínica insuficiente y la robustez algorítmica limitada. La evidencia científica revela patrones sistemáticos de fallos en sistemas desplegados prematuramente, incluyendo errores diagnósticos derivados de algoritmos entrenados con "datasets" sesgados, incompletos o metodológicamente deficientes.

Los casos documentados ilustran estas vulnerabilidades: Watson for Oncology de IBM generó recomendaciones clínicamente "inseguras e incorrectas" según opiniones de médicos usuarios, mientras que múltiples algoritmos de diagnóstico por imagen para Covid-19 demostraron fallos metodológicos fundamentales al aprender a identificar características técnicas del equipamiento radiológico (marca del escáner, parámetros de

adquisición) en lugar de patrones patológicos de la enfermedad, comprometiendo su generalización clínica.

El desafío de la explicabilidad algorítmica

La opacidad inherente de los algoritmos de aprendizaje profundo genera una barrera crítica para la adopción clínica responsable: la imposibilidad de rastrear el proceso de toma de decisiones algorítmicas limita tanto la identificación de errores sistemáticos como la corrección de sesgos operacionales. Esta falta de interpretabilidad erosiona la confianza médica y compromete la seguridad del paciente.

El paradigma "explicabilidad-confianza-adopción" establece una secuencia causal fundamental: sin transparencia algorítmica no se puede generar confianza clínica informada, y sin esta confianza, la implementación responsable resulta inalcanzable. Los marcos regulatorios reconocen esta realidad: la FDA exige documentación detallada de los procesos de toma de decisiones para dispositivos médicos con IA, mientras que el AI Act europeo establece requisitos específicos de explicabilidad para sistemas de alto riesgo en salud.

La consecuencia práctica es limitante: ante un sistema opaco, el clínico se ve reducido a una respuesta binaria —aceptar o rechazar la recomendación— sin capacidad de evaluación crítica o integración contextual del conocimiento clínico.

Manipulación y sesgos en la presentación de resultados

Conflictos de interés y sesgo de publicación en IA médica

La presión comercial genera distorsiones sistemáticas en la literatura científica sobre IA médica, manifestándose en el sesgo de publicación hacia resultados favorables y la omisión deliberada de limitaciones metodológicas, fallos operacionales y análisis de riesgo. Esta práctica constituye una forma de “cherry-picking científico” que compromete la integridad de la evidencia disponible para la toma de decisiones clínicas.

El fenómeno trasciende la simple selección de datos e involucra la construcción de narrativas optimistas que ocultan la complejidad inherente a la implementación de sistemas de IA en entornos clínicos reales. La consecuencia directa es la promoción de la adopción prematura de tecnologías insuficientemente validadas, incrementando los riesgos para la seguridad del paciente y la efectividad terapéutica.

Esta instrumentalización de la investigación científica para fines comerciales refleja lo que Habermas conceptualizó como la colonización del mundo de la

vida por la racionalidad técnico-instrumental: la subordinación del conocimiento científico a imperativos del mercado, desplazando el compromiso primario con el bienestar del paciente y la validez científica.

Amplificación de desigualdades

La representatividad sesgada en los “datasets” de entrenamiento constituye un factor crítico que socava la equidad diagnóstica en sistemas de IA otorrinolaringológicos. La predominancia de datos provenientes de cohortes de países desarrollados introduce sesgos de selección que limitan la generalización de estos modelos a poblaciones subrepresentadas, generando disparidades sistemáticas en la precisión diagnóstica.

Esta problemática se ejemplifica claramente en algoritmos de triaje que han demostrado subestimar consistentemente la gravedad clínica en pacientes de minorías étnicas, vulnerando así el principio fundamental de justicia en la distribución equitativa de recursos sanitarios.

Riesgos de Alucinaciones y Falta de Contexto Clínico

Las "alucinaciones" algorítmicas

Los errores de interpretación de la IA, conocidos técnicamente como "alucinaciones", representan un riesgo sistémico para la calidad asistencial. Estos fallos se originan tanto por la complejidad opaca de los algoritmos como por limitaciones en los datos utilizados para su desarrollo.

Cuando estos sistemas generan información médica incorrecta pero aparentemente creíble —especialmente en herramientas de diagnóstico por imagen o asistencia clínica— pueden inducir errores diagnósticos con graves repercusiones para la seguridad del paciente y la responsabilidad institucional.

Limitaciones del contexto clínico

Los algoritmos de IA procesan datos con extraordinaria eficiencia, pero su funcionamiento se limita al reconocimiento de patrones matemáticos sin alcanzar la comprensión integral que demanda la práctica médica. Esta limitación cobra especial relevancia en especialidades como la Otorrinolaringología, donde cada patología trasciende su manifestación puramente orgánica.

Considérese, por ejemplo, cómo una pérdida auditiva no es meramente un dato audiológico: implica transformaciones en la comunicación intrafamiliar, aislamiento social potencial y adaptaciones laborales específicas. Del mismo modo, un trastorno vocal no se reduce a parámetros acústicos, sino que

puede comprometer la identidad profesional de un docente o cantante. Una intervención nasal conlleva no solo corrección funcional, sino expectativas estéticas que requieren comprensión empática de la autoimagen del paciente.

Esta complejidad contextual permanece fuera del alcance algorítmico, consolidando el rol insustituible del juicio clínico experto.

Deshumanización de la relación médico-paciente

La deshumanización del acto médico

La progresiva algoritmización de la práctica médica conlleva el riesgo de desarticular uno de los elementos más poderosos de la medicina: la relación médico y paciente. Esta relación, construida sobre pilares de confianza mutua, comprensión empática y compromiso compartido hacia la sanación, trasciende cualquier capacidad computacional.

En Otorrinolaringología, donde las patologías frecuentemente repercuten en funciones tan personales como la voz, la audición y el equilibrio, la dimensión relacional cobra especial relevancia. El paciente que padece una pérdida auditiva no solo requiere la adaptación de una prótesis auditiva, sino la comprensión de sus temores, la validación de su experiencia y el acompañamiento en dicho proceso.

Esta complejidad relacional permanece irremediablemente fuera del alcance algorítmico, constituyendo el territorio exclusivo de la humanidad médica.

La irreemplazabilidad del vínculo humano

En la práctica otorrinolaringológica diaria emerge una realidad contradictoria: mientras los algoritmos de IA pueden procesar síntomas y generar respuestas aparentemente comprensivas con mayor consistencia que algunos profesionales en los momentos difíciles, los pacientes mantienen una búsqueda instintiva de conexión humana auténtica durante sus procesos de enfermedad.

Esta búsqueda no traduce cierta nostalgia por los métodos tradicionales, sino el reconocimiento de una necesidad existencial fundamental. Cuando un paciente afronta una pérdida progresiva de audición o una cirugía vocal que puede alterar su identidad profesional, requiere de algo más que protocolos de tratamiento optimizados: necesita la presencia de alguien que pueda comprender genuinamente su vulnerabilidad, validar sus temores y acompañar su proceso de adaptación.

Los sistemas robóticos, pese a su eficiencia técnica, generan en los pacientes una sensación de vacío relacional que ninguna sofisticación algorítmica logra llenar. Esta limitación subraya que la Medicina, en su esencia más profunda, permanece como un encuentro entre seres humanos unidos por la experiencia compartida de la fragilidad.

La imperante necesidad de supervisión clínica y transparencia algorítmica

Supervisión humana como salvaguarda ética

La supervisión humana constante constituye un requisito ineludible para el uso ético y seguro de la IA. Toda decisión clínica asistida por IA debe ser validada por un profesional sanitario cualificado, quien nunca puede delegar por completo el juicio clínico en la tecnología. Los casos judiciales han dictaminado que los médicos siguen siendo responsables de validar los resultados de la IA antes de tomar decisiones críticas.

Transparencia y explicabilidad

Es obligatorio explicar al paciente cómo funciona el sistema de IA, cuál es su papel en el diagnóstico o tratamiento y detallar sus limitaciones. La transparencia algorítmica se está convirtiendo en un requisito regulatorio, no en una opción. Tanto la FDA como el AI Act europeo exigen que los sistemas de IA de alto riesgo sean suficientemente transparentes para permitir la interpretación humana de sus resultados.

Consentimiento informado: Obligación ética y legal

Marco normativo europeo

La normativa europea, especialmente el AI Act y el RGPD, exige que los pacientes sean informados de manera clara sobre el empleo de sistemas de IA en su atención sanitaria, asegurando su derecho a decidir y solicitar supervisión humana en las decisiones automatizadas. El consentimiento informado debe integrar la explicación del papel de la IA en el diagnóstico o tratamiento, los posibles riesgos asociados, sus limitaciones tecnológicas y la supervisión médica que se realiza.

Fundamentos bioéticos del consentimiento

La transparencia y el respeto a la autonomía del paciente requieren que toda intervención médica, incluida la mediada por IA, se realice tras un proceso informado donde se detallen los beneficios, riesgos y alternativas disponibles. El paciente tiene derecho a saber si sus datos serán utilizados

para entrenar algoritmos o si el diagnóstico ha sido asistido con tecnología automatizada, pudiendo decidir si acepta este modelo de atención.

Responsabilidad legal y ambigüedad jurídica

El dilema de la responsabilidad

La ambigüedad sobre quién es responsable ante un error de la IA —el médico, el hospital, el desarrollador— constituye uno de los desafíos éticos y legales más graves. Actualmente, los marcos regulatorios y judiciales aún están en evolución para determinar la responsabilidad ante daños a pacientes por decisiones algorítmicas. Esta incertidumbre legal genera ansiedad profesional y puede afectar la calidad de la atención.

Hacia una responsabilidad compartida

Se requiere un marco de responsabilidad compartida donde el desarrollador de software, la institución sanitaria y el profesional médico asuman sus respectivas obligaciones dentro de un sistema de supervisión y control riguroso.

Discusión. Hacia una IA ética en otorrinolaringología

Los cuatro principios bioéticos adaptados

La implementación responsable de la IA debe basarse en los cuatro principios bioéticos clásicos adaptados al contexto tecnológico:

- Autonomía: Los pacientes deben mantener su capacidad de decisión informada sobre el uso de IA en su atención.
- Beneficencia: La IA debe procurar un beneficio real y demostrable, no meramente teórico.
- No maleficencia: Debe evitarse todo daño por errores algorítmicos, sesgos o falta de supervisión.
- Justicia: Los avances de IA deben distribuirse equitativamente, no amplificar desigualdades existentes.

Medicina colaborativa como modelo ideal

El modelo óptimo para integrar la inteligencia artificial en Otorrinolaringología no contempla la sustitución del especialista, sino la evolución hacia un paradigma de medicina colaborativa. En este enfoque, la tecnología actúa como amplificador cognitivo del “expertise” clínico, potenciando las capacidades diagnósticas sin desplazar el juicio médico.

Esta sinergia tecnológica-humana permitiría automatizar procesos analíticos complejos y tareas repetitivas, redistribuyendo el tiempo clínico hacia competencias esencialmente humanas: la comunicación terapéutica, la evaluación integral del paciente y la toma de decisiones contextualmente informadas.

La medicina colaborativa preservaría así la dimensión relacional de la práctica médica, elemento fundamental que distingue el acto médico de un simple procesamiento de datos. Esta reconfiguración del trabajo clínico maximizaría la capacidad diagnóstica y terapéutica del otorrinolaringólogo, manteniendo intacto el núcleo humanístico de la excelencia asistencial.

Conclusiones

La inteligencia artificial en Otorrinolaringología presenta un potencial transformador significativo, sin embargo, su implementación clínica actual evidencia déficits éticos estructurales que requieren un abordaje inmediato y sistemático. La convergencia de factores críticos —validación científica insuficiente, evaluación acrítica de beneficios, distorsiones metodológicas en investigación comercial, y la omisión sistemática de riesgos operacionales como alucinaciones algorítmicas y descontextualización clínica— configura un escenario de incertidumbre que compromete tanto la seguridad del paciente como los principios fundamentales de la práctica médica.

El desafío trasciende la dimensión tecnológica para situarse en el núcleo de la identidad profesional médica: garantizar que el propósito central de la medicina —la mitigación del sufrimiento y el cuidado integral de la persona— mantenga primacía sobre la eficiencia algorítmica. La supervisión médica continua, la transparencia en los procesos de toma de decisiones automatizados, y el respeto riguroso del consentimiento informado constituyen no solamente buenas prácticas, sino imperativos deontológicos y marcos regulatorios ineludibles.

Nuestra especialidad médica debe asumir el liderazgo en la construcción de una cultura de evaluación crítica, transparencia metodológica y diálogo interdisciplinario que subordine la implementación de IA a los principios de beneficencia, no maleficencia y respeto por la autonomía del paciente. La actualización del principio hipocrático de “primum non nocere” en el contexto algorítmico exige esta aproximación ética como condición sine qua non para materializar el potencial transformador de estas tecnologías sin comprometer la dimensión humanística que define la excelencia en la práctica otolaringológica.

Bibliografía

1. Aranguren, J. L. (1994). *Ética y política*. Biblioteca Nueva.
2. Beauchamp, T. L., & Childress, J. F. (2019). *Principles of biomedical ethics* (8th ed.). Oxford University Press.
3. Comisión Europea. (2024). Reglamento (UE) 2024/1689 del Parlamento Europeo y del Consejo por el que se establecen normas armonizadas en materia de inteligencia artificial (Ley de IA). Diario Oficial de la Unión Europea.
4. Cortina, A. (2013). *¿Para qué sirve realmente la ética?* Paidós.
5. Food and Drug Administration. (2021). *Artificial Intelligence/Machine Learning (AI/ML)-Based Software as a Medical Device (SaMD) Action Plan*. FDA.
6. Habermas, J. (1968). *Ciencia y técnica como ideología*. Tecnos.
7. McCarthy, J. (1956). A proposal for the Dartmouth Summer Research Project on Artificial Intelligence. *AI Magazine*, 27(4), 12-14.
8. Obermeyer, Z., Powers, B., Vogeli, C., & Mullainathan, S. (2019). Dissecting racial bias in an algorithm used to manage the health of populations. *Science*, 366(6464), 447-453.
9. Parlamento Europeo y Consejo. (2016). Reglamento (UE) 2016/679 relativo a la protección de las personas físicas en lo que respecta al tratamiento de datos personales (RGPD). Diario Oficial de la Unión Europea.
10. Sociedad Española de Medicina Interna (SEMI). (2023). Médicos internistas piden una aplicación responsable y ética de la IA basada en un marco regulatorio, la transparencia algorítmica y la formación de profesionales. Comunicado de prensa.
11. Vayena, E., Blasimme, A., & Cohen, I. G. (2018). Machine learning in medicine: Addressing ethical challenges. *PLOS Medicine*, 15(11), e1002689.
12. World Health Organization. (2021). *Ethics and governance of artificial intelligence for health*. WHO Press.
13. World Medical Association. (2013). Declaration of Helsinki: Ethical principles for medical research involving human subjects. *JAMA*, 310(20), 2191-2194.