

ARTÍCULO ORIGINAL

Respuestas y precisión de inteligencias artificiales a *prompts* dirigidos al manejo de la lesión por presión

HIGHLIGHTS

1. Se comparó la precisión de las IA mediante *prompts* sobre lesión por presión.
2. La precisión de las IA fue del 31,9% con 23 aciertos.
3. El *hybrid prompting* presentó mayor precisión que el *zero-shot*.
4. Gemini® presentó mayor precisión para el manejo de la lesión por presión.

Letícia Pereira Pires¹ 
Maithê de Carvalho e Lemos Goulart¹ 
Fernanda Garcia Bezerra Góes¹ 
Yonara Cristiane Ribeiro¹ 
Luíza Lissonger Costa Guimarães¹ 
Lediane Farias de Sousa¹ 

RESUMEN

Objetivo: Comparar las respuestas y la precisión de diferentes Inteligencias Artificiales a partir de un conjunto de *prompts* para el manejo de la lesión por presión. **Método:** Estudio descriptivo, comparativo, de las respuestas de Inteligencias Artificiales, realizado entre octubre y noviembre de 2025, en dos momentos: 1) Estudio de *prompts* y 2) Prueba para la obtención de las respuestas. Se evaluaron Perplexity AI, Grok AI, Blackbox AI, OpenAI ChatGPT-4.0®, Gemini® y Claude AI, divididas aleatoriamente en los grupos *hybrid prompting* y *zero-shot*. Las respuestas se organizaron en una planilla de Excel y se analizaron de forma descriptiva. **Resultados:** La precisión general de las inteligencias artificiales fue del 31,9%. El grupo *hybrid prompting* alcanzó el 36,1% y el *zero-shot* obtuvo el 27,7%. La IA Gemini® demostró la mayor precisión y el ChatGPT-4.0®, la peor. **Conclusión:** La precisión presentada indica que la recomendación más segura es el modelo *hybrid prompting*, integrando la inteligencia artificial en el juicio clínico humano.

DESCRIPTORES: Inteligencia Artificial Generativa; Grandes Modelos de Lenguaje; Grupo de Enfermería; Razonamiento Clínico; Úlcera por Presión.

CÓMO REFERIRSE A ESTE ARTÍCULO:

Pires LP, Goulart MCL, Góes FGB, Ribeiro YC, Guimarães LLC, de Sousa LF. Respuestas y precisión de inteligencias artificiales a *prompts* dirigidos al manejo de la lesión por presión. Cogitare Enferm [Internet]. 2026 [cited "insert year, month and day"];31:e102398es. Available from: <https://doi.org/10.1590/ce.v31i0.102398es>

INTRODUCCIÓN

Las lesiones por presión (LP) constituyen una complicación frecuente y global, con una prevalencia que puede alcanzar el 72,5%, variando según el contexto¹. Desde 2013, el programa *Patients for Patient Safety* incluye la reducción de las LP como su sexta meta, reforzando la importancia de medidas que cualifiquen la asistencia y promuevan la seguridad del paciente².

La *National Pressure Injury Advisory Panel (NPIAP)* define la LP como daño localizado en la piel o el tejido subyacente, generalmente sobre prominencias óseas o áreas sometidas a presión, cizallamiento o factores químicos, físicos y biológicos³. Se clasifican en estadios que varían de uno (eritema no blanqueable) a cuatro (pérdida total del espesor con exposición de estructuras profundas), además de la lesión no clasificable, tisular profunda, relacionada con dispositivos médicos y en mucosas⁴.

La aparición de LP en pacientes institucionalizados se ha asociado a fallas asistenciales, como reposicionamiento inadecuado, ausencia de examen físico, déficit de recursos humanos y materiales, además de un dominio insuficiente sobre apósitos, curaciones y nuevas tecnologías². La actualización científica, asociada a la habilidad técnica del enfermero, es fundamental para la elección adecuada del apósito, contribuyendo a la calidad del cuidado y a la prevención de complicaciones³.

La tecnología en salud se ha expandido en las últimas décadas, destacándose la Inteligencia Artificial (IA), capaz de procesar grandes volúmenes de datos rápidamente y apoyar decisiones personalizadas⁵. La creciente integración de la IA a los servicios de salud indica una tendencia de uso ampliado, tornando esencial reconocer su potencial para perfeccionar conocimientos y prácticas. En ese escenario de transformación acelerada, discutir ventajas y limitaciones de esos modelos fortalece una postura crítica frente al uso de la IA en el cuidado en salud⁶.

En la salud, la IA se ha introducido mayoritariamente como una herramienta de apoyo a la toma de decisión clínica en situaciones complejas, tales como la identificación de factores de riesgo, la definición de pronóstico y la precisión diagnóstica de LP⁷. Un estudio alemán señaló la utilización de la IA por enfermeros para el reconocimiento, la clasificación, la reducción, la predicción y la estratificación de riesgo de LP, destacando el alto grado de especificidad⁸, evidenciando que la IA no solo complementa prácticas clínicas rutinarias, sino que también demuestra su aplicabilidad en escenarios complejos de la salud.

Se describen dos modelos principales de IA: la predictiva, orientada a la previsión de desenlaces a partir de algoritmos alimentados por un gran volumen de datos; y la generativa, que articula los algoritmos capaces de crear contenidos, incluyendo textos, imágenes, audio, entre otros⁹. Los *Large Language Models (LLM)*, o Modelos de Lenguaje de Gran Escala, son sistemas de IA entrenados con grandes volúmenes de datos para reconocer patrones y generar respuestas en lenguaje natural. Los LLM, cuando se optimizan mediante enfoques estructurados, como la integración con bases de conocimiento y directrices clínicas, pueden presentar un desempeño superior y alineamiento con las directrices y recomendaciones, incluso en el contexto de la enfermería¹⁰⁻¹¹. Los *chatbots* (espacio de interacción) alimentados por la IA generativa se consolidan como un recurso de apoyo, ofreciendo respuestas personalizadas a los usuarios⁶.

La IA puede elaborar recomendaciones y conductas personalizadas considerando el escenario terapéutico del paciente; no obstante, son respuestas dadas a partir de comandos llamados *prompts* y algoritmos sofisticados. Posibilita, además, identificar patrones clínicos al analizar extensos conjuntos de datos⁹, proporcionando información en tiempo real. Sin embargo, se destaca que incluso los modelos más perfeccionados

pueden no captar toda la complejidad de la pericia clínica de los profesionales en el manejo de las lesiones por presión o heridas crónicas¹⁰⁻¹¹.

La integración de esa tecnología en la práctica de enfermería exige un abordaje cauteloso y consciente por parte del profesional, siendo esencial que su uso refleje sus potenciales beneficios, limitaciones y riesgos para la seguridad del paciente⁶. Sin embargo, el uso de la IA debe comprenderse como una herramienta de apoyo, y no como sustituto del juicio clínico y de la toma de decisión profesional, debiendo utilizarse para potenciar el cuidado^{5,10-11}.

En ese sentido, la literatura plantea la necesidad de entrenar la IA con un gran volumen de datos, como los LLM; no obstante, en la práctica asistencial el profesional utiliza la herramienta más accesible, sin la inserción de datos estructurados. De esa forma, existe una laguna de conocimiento, que necesita ser cubierta, para entender cómo esas IA se comportan sin que esto se realice previamente. El uso de la IA en el cotidiano de la práctica, por lo tanto, debe estar aliado al rigor metodológico, al uso de datos fidedignos y al *prompt* solicitado, garantizando su confiabilidad y utilización como recurso facilitador en el manejo de las LP, auxiliando en la elección del apósito más adecuado y contribuyendo a la integridad de la piel del paciente.

Por consiguiente, el objetivo de este estudio es comparar las respuestas y la precisión de diferentes Inteligencias Artificiales a partir de un conjunto de *prompts* para el manejo de la lesión por presión.

MÉTODO

Estudio descriptivo y comparativo, de las respuestas de Inteligencias Artificiales a partir de un conjunto de *prompts*, realizado entre octubre y noviembre de 2025, en dos momentos: 1) Estudio de *prompts* y 2) Prueba para la obtención de las respuestas de las IA.

Se utilizaron las siguientes IA: Perplexity AI, Grok AI, Blackbox AI, OpenAI ChatGPT-4.0®, Gemini® y Claude AI. Las IA se seleccionaron por presentar plataformas cuyos términos de uso permiten el envío y el análisis de imágenes de LP, además de su indexación gratuita. Las IA Microsoft Copilot, DeepSeek y Meta AI se excluyeron por no realizar la lectura de imágenes conforme a los términos de uso.

Se optó por la división de las IA de forma aleatoria en dos grupos: el modelo *hybrid prompting*, articulando diferentes técnicas (*few-shot*, *role-prompting*, *chain-of-thought* y *chain-of-verification*), y el modelo *zero-shot* de la ingeniería de *prompts*¹². De acuerdo con la Pan American Health Organization, el término *prompt* se refiere a las instrucciones proporcionadas a un sistema de IA, con una pregunta breve o una solicitud detallada, que especifica el tono, el contexto, el formato y el objetivo deseado, cuya claridad y fundamentación contextual influyen directamente en la calidad de la respuesta generada¹³.

Se optó por las técnicas *hybrid prompting* y *zero-shot*, buscando atender de forma detallada y contextualizada la respuesta de la IA frente a las necesidades de decisión del enfermero para el manejo de la LP y, con base en sus potencialidades, para alcanzar la precisión y la confiabilidad en las respuestas¹⁴.

La técnica *hybrid prompting* integra diferentes técnicas de *prompt*, combinando fortalezas para perfeccionar el desempeño con respuestas más precisas y alineadas con los objetivos, lo que exige instrucciones más complejas¹⁵. La combinación ocurre por medio del *few-shot*, que proporciona ejemplos para orientar patrones de respuesta¹⁴; del *role prompting*, que atribuye un papel específico al modelo; del *chain-of-thought*, que lo lleva a explicitar el razonamiento paso a paso¹⁴; y del *chain-of-verification*,

estructurado en cuatro etapas: respuesta base, plan de verificación, comprobaciones y respuesta final, aumentando el rigor y la confiabilidad¹⁶. Por su parte, la estrategia *zero-shot* consiste en obtener una respuesta a partir de un único comando inicial, sin ofrecer ejemplos ni información adicional previa¹⁴.

Para la recolección de datos, que ocurrió en dos días del mes de noviembre de 2025, cada modelo fue alimentado con una imagen previamente seleccionada por el equipo, correspondiente a una LP no clasificable según el consenso NPIAP (Figura 1), insertada directamente en las IA utilizando cuentas (correo electrónico e inicio de sesión) creadas exclusivamente para el estudio, para minimizar sesgos derivados de *prompts* anteriores en el perfil individual de usuario. Las respuestas generadas se registraron en el perfil del proyecto y se guardaron como archivo digital mediante capturas de pantalla.



Figura 1. Imagen de lesión por presión no clasificable seleccionada por el equipo del estudio. Rio das Ostras, RJ, Brasil, 2025

Fuente: Feridas Crônicas - Recurso Educacional sobre Prevenção e Manejo da Lesão por Pressão (2020)¹⁷.

La LP no clasificable se caracteriza por la pérdida de la piel en todo su espesor, cuya profundidad y extensión reales no pueden medirse debido al recubrimiento por tejido necrótico o desvitalizado. Solo después de su remoción es posible la evaluación completa, evidenciando generalmente una LP estadio 3 o 4¹⁸⁻¹⁹.

La aleatorización por medio de sorteo ocurrió antes de la alimentación de la IA con la imagen, en un proceso en el cual un miembro del equipo de investigación escribió en papeles los nombres de las IA y los colocó en una caja. Se inició el sorteo por el grupo de *hybrid prompting* y, después del primero, se alternaron los grupos hasta finalizar el sorteo.

El grupo que recibió el *prompt zero-shot* estuvo compuesto por Perplexity AI, Grok AI y Blackbox AI, que recibieron la siguiente instrucción: "Responda el instrumento a continuación con base en la imagen". El referido instrumento fue desarrollado para este estudio y está constituido por 12 ítems redactados en formato de afirmaciones (Figura 2). Los tipos de apósito listados en la lista de verificación se definieron con base en las opciones más frecuentes utilizadas por enfermeros, según la literatura¹⁸⁻²⁰.

1. Clasifico esta lesión como:

() Estadio 1 () LPP por dispositivo médico () Agua destilada () Ringer () No realizaría la higienización con ninguno de estos materiales

() Estadio 2 () LPP de tejidos profundos () Clorhexidina () PHMB

() Estadio 3 () LPP en membrana mucosa () Jabón () Yodo

() Estadio 4 () No clasificable () Suero fisiológico 0,9% () Agua oxigenada

2. Identifico que la imagen presenta las siguientes estructuras afectadas:

() Epidermis () Cartilago () Ácidos grasos esenciales () Apósitos de espuma () Carbón activado

() Dermis () Ligamento () Azúcar () Crema/Spray barrera () Alginato de calcio

() Hipodermis () Hueso () Hidrofibra () Impregnados en plata () Presión negativa

() Fascia () Tendones () Hidropolímero () Colagenasa () Papaína

() Músculo () Ninguna estructura afectada () Hidrogel () Hidrocoloide

3. ¿Identifico los siguientes tejidos en la lesión?

() Tejido de granulación () Vitrina () Mantener la humedad de la lesión () Eliminar el tejido desvitalizado

() Esfacelo () Tejido necrótico () Controlar la carga bacteriana () Promover la granulación

() Escara () No identifico ninguno () Controlar el dolor

() Epíhole

4. Identifico exudado en la lesión:

() Si () Tal vez

() No

5. El borde de la lesión presenta signos de:

() Maceración () Hiperemia

() Epíhole () Ninguna alteración identificada

6. La piel perilesional presenta:

() Integridad conservada () Eritema

() Edema () Descamación

() Ninguna alteración identificada

7. Realizaría la higienización de esta LPP con:

8. Elijo como cobertura primaria para la realización de este apósito de LPP el/los siguiente(s) elemento(s):

9. El principal objetivo del apósito que elegí es:

10. Sobre el desbridamiento:

11. Elijo como cobertura secundaria para la realización de este apósito de LPP el/los siguiente(s) elemento(s):

12. Realicé el cambio del apósito con la siguiente frecuencia:

Figura 2. Instrumento elaborado para el grupo de IA con prompt zero-shot. Rio das Ostras, RJ, Brasil, 2025

Fuente: Los autores (2025).

El grupo compuesto por las IA OpenAI ChatGPT-4.0®, Gemini® y Claude AI recibió la técnica de *hybrid prompting*, con el *prompt* creado para este estudio: *“Usted es un enfermero especialista en curación de heridas con experiencia en evaluación, clasificación y manejo de lesiones por presión. Su público objetivo es un paciente que presenta una lesión tegumentaria. Su objetivo, a partir de la imagen adjunta, es clasificar la lesión por presión, identificar las estructuras afectadas y los tejidos presentes, identificar si hay presencia de exudado, indicar cómo realizar la limpieza de esa lesión, indicar el apósito primario y secundario y recomendar la frecuencia del cambio de apósitos. El formato de la respuesta debe ser en texto estructurado, con un tono informativo, profesional y conciso. Si la imagen es insuficiente para una decisión segura, informe explícitamente las limitaciones y describa qué datos clínicos adicionales son necesarios (ej.: signos vitales, comorbilidades, signos de infección y exudado). Siempre use un lenguaje de tipo ‘sugerencia’ cuando haya incertidumbre y evite afirmar un diagnóstico médico definitivo. Para llegar a la conclusión, exponga paso a paso el razonamiento y liste las comprobaciones realizadas (mínimo 3), por ejemplo: correspondencia con criterios de clasificación, verificación de signos de infección, compatibilidad del apósito con el exudado, y verificación de riesgos de maceración o trauma.”*

Las respuestas esperadas, considerando ambos los *prompts*, se definieron con base en la literatura actualizada^{18-19,21-22} sobre el tema y se disponen en el Cuadro 1.

Cuadro 1. Respuestas esperadas de las IA sobre el manejo de la LP. Rio das Ostras, RJ, Brasil, 2025

(continuar)

Ítem	Respuesta	Justificación
1. Clasifico esta lesión como:	No clasificable	Corresponde a una LP no clasificable, con pérdida total de la piel y daño tisular cuya profundidad no puede determinarse por estar cubierta por esfacelo o escara. Esa cobertura impide definir el estadio. Tras la remoción del tejido desvitalizado, la lesión podrá reclasificarse como LP 3 o 4, según la profundidad y la extensión del daño ¹⁶ .

Cuadro 1. Respuestas esperadas de las IA sobre el manejo de la LP. Rio das Ostras, RJ, Brasil, 2025

(conclusión)

Ítem	Respuesta	Justificación
2. Identifico que la imagen presenta las siguientes estructuras afectadas:	Dermis y epidermis	Epidermis y dermis como estructuras afectadas por la pérdida de espesor total, involucrando tanto la capa superficial como la capa más interna de la piel ¹⁷ .
3. Identifico los siguientes tejidos en la lesión	Tejido necrótico, esfacelo y granulación	Presencia de tejido necrótico (escara), esfacelo y tejido de granulación, evidenciando la coexistencia de áreas desvitalizadas y en proceso de reparación ¹⁷ .
4. Identifico exudado en la lesión:	No	No evidencia presencia de exudado, dado que no hay signos visibles sobre el lecho de la herida ¹⁷ .
5. El borde de la lesión presenta signos de:	Epíbole	Presenta epíbole en el borde de la lesión, caracterizada por el borde enrollado de la herida ¹⁷ .
6. La piel alrededor de la lesión presenta:	Eritema	Área de eritema, caracterizada por enrojecimiento de la piel derivado del aumento del flujo sanguíneo local en las áreas alrededor de la lesión ¹⁷ .
7. Realizaría la higienización de esta LP con:	Agua corriente o solución salina (suero fisiológico 0,9%) y Polihexametileno Biguanida (PHMB)	Solución salina 0,9% o PHMB, que son soluciones no citotóxicas, preservan fibroblastos y tejido de granulación y favorecen la cicatrización ¹⁷ . El agua corriente también es una alternativa segura y de bajo costo, especialmente en el cuidado domiciliario ¹⁹ .
8. Elijo como apósito primario para la realización de esta curación de LP el/los siguiente(s) ítem(s):	Papaína, hidrogel, ningún apósito o gasa humedecida con suero fisiológico 0,9%.	Papaína 10%, favorece el desbridamiento químico. Papaína al 2-4%, auxilia en la reparación tisular ²⁰ . Hidrogel, para mantener el ambiente húmedo. Gasa humedecida con suero fisiológico 0,9%, contexto de pocos recursos. Film adherente, enzimas o, en casos específicos, ningún apósito, para necrosis. El film ablanda la escara, mientras las enzimas la descomponen. Escara seca e íntegra, si el desbridamiento no forma parte del plan terapéutico, no aplicar apósito es aceptable ¹⁷ .
9. El principal objetivo del apósito que elegí es:	Remover el tejido desvitalizado, promover la granulación y mantener la humedad	Remoción del tejido no viable, para mantener la humedad adecuada, prevenir y controlar infecciones y olores, promover la granulación, la limpieza y el alivio del dolor ¹⁷ .
10. Sobre el desbridamiento:	Desbridamiento quirúrgico, autolítico, o mecánico o enzimático	Desbridamiento para remover el tejido no viable, reducir el riesgo de infección y favorecer la cicatrización. Realizado de forma quirúrgica, autolítica, mecánica o enzimática ¹⁷ .
11. Elijo como apósito secundario para la realización de esta curación de LP el/los siguiente(s) ítem(s):	Gasa seca y cinta adhesiva microporosa	La gasa seca proporciona protección y revestimiento adecuados al apósito primario, mientras que la cinta adhesiva microporosa desempeña la función de fijación del apósito a la piel ¹⁷ .
12. Realicé el cambio del apósito con la siguiente frecuencia:	Diariamente	Monitorear la evolución de la cicatrización, identificar alteraciones en el lecho de la herida y ajustar la conducta terapéutica ¹⁷ .

Fuente: Los autores (2025).

Las respuestas recolectadas de cada IA se insertaron en una planilla de *Microsoft Excel®* y se analizaron de forma descriptiva, dicotomizándose entre “correctas” e “incorrectas”. Se calcularon: 1) el porcentaje de aciertos en relación con cada ítem del instrumento; 2) el porcentaje de aciertos total y 3) el porcentaje de aciertos por grupo de IA. Se destaca que los ítems 1, 4 y 12 poseen solo una alternativa correcta. En cada pregunta, se consideraron “Correctas” solo las alternativas que corresponden integralmente a las afirmaciones esperadas; todas las demás se clasificaron como “Incorrectas”. Las respuestas presentadas con la opción “Tal vez” también se categorizaron como “Incorrectas”.

Para el universo de aciertos se consideraron los 12 ítems del instrumento y las seis IA analizadas, totalizando 72 aciertos posibles. Para la precisión general de las IA, se sumó el total de aciertos y se dividió por el universo de aciertos posibles. Para el cálculo de la precisión general de los ítems, se sumaron los aciertos del ítem y se dividió por el universo de aciertos posibles en el ítem, considerando la cantidad de IA, ese número fue de seis. En el cálculo de la precisión por grupos, se sumó el total de aciertos en cada grupo y se dividió por el universo de aciertos posibles, siendo 36 en cada uno. Para calcular cuál IA fue más asertiva, se utilizó el total de aciertos presentados dividido por el total de ítems del instrumento.

Se dispensó la presentación del estudio al Comité de Ética en Investigación, dado que la investigación no involucró seres humanos, conforme a la Ley n.º 14.874/2024.

RESULTADOS

La precisión general de las IA analizadas fue del 31,9% con 23 aciertos. No obstante, el porcentaje del grupo *hybrid prompting* fue del 36,1% con 13 aciertos y 23 errores en total, demostrando una precisión ligeramente superior al grupo *zero-shot*, que fue del 27,7% con 10 aciertos y 26 errores en total (Tabla 1).

Tabla 1. Porcentaje de aciertos según el instrumento de los seis modelos de IA. Rio das Ostras, RJ, Brasil, 2025

Ítem	N.º de Aciertos	N.º de Errores	Precisión de los ítems	Zero-Shot Aciertos	Hybrid Prompting Aciertos
Ítem 1	2	4	33,3%	0	2
Ítem 2	0	6	0,0%	0	0
Ítem 3	1	5	16,6%	0	1
Ítem 4	1	5	16,6%	1	0
Ítem 5	0	6	0,0%	0	0
Ítem 6	3	3	50,0%	2	1
Ítem 7	6	0	100,0%	3	3
Ítem 8	1	5	16,6%	0	1
Ítem 9	2	4	33,3%	0	2
Ítem 10	6	0	100,0%	3	3
Ítem 11	0	6	0,0%	0	0
Ítem 12	1	5	16,6%	1	0
Precisión general (%): 31,9%				27,7%	36,1%

Fuente: Los autores (2025).

Entre los ítems analizados, solo dos presentaron un acierto del 100%: el ítem 7, referente a la higienización de la lesión, y el ítem 10, relativo al desbridamiento, indicando consenso entre todas las IA en esos aspectos. En contraste, ítems fundamentales para la evaluación de la lesión no registraron acierto, como el ítem 2 (estructuras afectadas), el ítem 5 (borde de la lesión) y el ítem 11 (apósito secundario), evidenciando la incapacidad de los modelos de identificar elementos esenciales para el manejo de la LP. En los demás ítems, se observó variabilidad entre los dos tipos de *prompts*, con diferencias puntuales de desempeño.

Entre las IA analizadas, la Gemini® presentó la mayor precisión con seis (50,0%) aciertos, seguida de la Claude AI con cinco (41,6%) aciertos, Perplexity AI con cuatro (33,3%) aciertos y las Grok AI y Blackbox AI con tres (25,0%) cada una. La peor precisión fue presentada por el ChatGPT-4.0® con dos (16,6%) aciertos. Los resultados demostraron heterogeneidad en las respuestas.

El grupo de *hybrid prompting* presentó mayor precisión en la identificación y el manejo de la LP (Figura 3). En cuanto a la identificación de la LP, el grupo *zero-shot* no presentó aciertos referentes a la clasificación, a las estructuras afectadas ni al borde de la lesión. La Blackbox AI presentó una respuesta correcta en cuanto a la ausencia de exudado, y las Perplexity AI y Grok AI en el reconocimiento del eritema en la piel alrededor. Por su parte, el grupo *hybrid prompting* no fue asertivo en las respuestas referentes a estructuras afectadas, exudado y borde de la lesión. Tanto la Gemini® como la Claude AI acertaron la clasificación de la LP, siendo que la Gemini® también localizó el eritema en la piel alrededor y la Claude AI describió correctamente los tejidos presentes en la imagen.

En cuanto al manejo de la LP, la higienización con suero fisiológico 0,9% y la necesidad de desbridamiento fueron unánimemente indicadas entre las seis IA. El grupo *zero-shot* no fue asertivo en los ítems referentes al apósito primario, al objetivo de la curación y al apósito secundario. En el ítem cambio de apósito, solo la Perplexity AI respondió que el cambio debería ocurrir diariamente. En el grupo *hybrid prompting*, los ítems de apósito secundario y cambio de apósito no fueron respondidos correctamente por ninguna de las tres IA. La Gemini® se destacó por acertar además los ítems de apósito primario y objetivo, y la Claude AI en el objetivo.

IDENTIFICACIÓN DE LA LESIÓN							
Prompt	IA	Clasificación	Estructuras afectadas	Tejidos	Exudado	Borde de la lesión	Piel perilesional
Zero-Shot	Perplexity AI	Estadio 4	Epidermis, dermis, hipodermis, fascia y músculo	Tejido de granulación y tejido necrótico	SI	Hiperemia	Eritema
	Grok AI	Estadio 4	Dermis, hipodermis, fascia, músculo, hueso	Esófago, fibrina y tejido necrótico	SI	Hiperemia	Eritema
	BLACKBOX AI	Estadio 4	Epidermis, dermis, hipodermis y fascia	Esófago y tejido necrótico	No	Hiperemia	Integridad conservada
Hybrid Prompting	OpenAI Chat GPT-4.0®	Estadio 3 o 4	Epidermis, dermis y tejido subcutáneo	Tejido necrótico y tejido de granulación	Bajo o ausente	Eritema	Hiperemia
	Gemini®	No clasificable	Epidermis, dermis y tejido subcutáneo	Tejido necrótico y tejido de granulación	No se pudo evaluar de forma fiable	Hiperemia	Eritema
	Claude AI	No clasificable	Epidermis, dermis y (tejido subcutáneo probable)	Tejido necrótico, tejido de granulación y esófago	No se pudo evaluar adecuadamente	Hiperemia	Hiperemia
MANEJO DE LA LESIÓN POR PRESIÓN							
Prompt	IA	Limpieza	Apósito primario	Objetivo	Desbridamiento	Apósito secundario	Cambio de apósito
Zero-Shot	Perplexity AI	Suero fisiológico 0,9%	Hidrofibra, apósitos de espuma, apósito impregnado en plata y alginato de calcio	Eliminar el tejido desvitalizado, controlar la carga bacteriana y promover la granulación	Desbridamiento quirúrgico o enzimático	Apósito de película e hidrocoloide	Diariamente
	Grok AI	Suero fisiológico 0,9%	Hidrofibra, apósito impregnado en plata, colagenasa, alginato de calcio	Controlar la carga bacteriana	Desbridamiento quirúrgico	Venda y cinta de micropore	Cada dos días
	BLACKBOX AI	Suero fisiológico 0,9%	Hidrofibra, hidropolímero, apósitos de espuma, colagenasa, hidrogel e hidrocoloide	Eliminar el tejido desvitalizado y promover la granulación	Desbridamiento enzimático	Gasa seca, apósito de película, venda, cinta de micropore y esparadrapo	Cada tres días
Hybrid Prompting	OpenAI Chat GPT-4.0®	Suero fisiológico 0,9%	Hidrogel, apósito de plata, povidona yodada, alginato de plata; mantener la escara intacta si está seca y estable	No indicó el objetivo	Desbridamiento quirúrgico	Espuma absorbente	48-72 horas; si hay infección, cambiar diariamente
	Gemini®	Suero fisiológico 0,9%	Hidrogel	Eliminar el tejido necrótico	Desbridamiento autolítico	Apósito de espuma con borde de silicona	24-48 horas
	Claude AI	Suero fisiológico 0,9%	Hidrogel con o sin alginato de calcio	Mantener un ambiente húmedo y proteger la lesión	Desbridamiento quirúrgico o enzimático	Gasa estéril fijada con cinta de micropore, película transparente, espuma de poliuretano	24-48 horas

Figura 3. Respuestas de la IA sobre la identificación y el manejo de la LP. Rio das Ostras, RJ, Brasil, 2025

Leyenda: color rojo - error en la respuesta; color verde - acierto en la respuesta.

Fuente: Los autores (2025).

DISCUSIÓN

Los hallazgos de este estudio evidencian que, aunque las IA evaluadas reconocieron conductas fundamentales y consensuales en el manejo de la LP, como la higienización con suero fisiológico y la indicación de desbridamiento, la precisión general puede considerarse baja, dado que no alcanzó siquiera la mitad de los aciertos posibles. A pesar del desempeño satisfactorio de esas herramientas en tareas de menor complejidad, se verificaron limitaciones importantes, incluyendo inconsistencias en las interpretaciones y la posibilidad de ocurrencia de alucinaciones, lo que compromete la confiabilidad de las respuestas generadas⁷.

El *hybrid prompting* presentó una precisión discretamente superior al compararlo con el *zero-shot*, incluyendo mayor cautela en las respuestas. En el ítem referente al exudado, por ejemplo, dos modelos reconocieron las limitaciones de la imagen y presentaron sus respuestas en tono sugestivo cuando no había elementos suficientes para afirmaciones definitivas. La literatura refuerza que no se debe restringir el análisis a un único tipo de *prompt*, dado que *prompts* elaborados y estructurados tienden a favorecer respuestas refinadas y análisis más completos²³.

La ingeniería de *prompt* puede aumentar la precisión en sistemas de IA aplicados a la salud, como demostró un estudio europeo que observó un incremento del 80,1% al 99,6% en la comparación entre consultas realizadas con comandos simples y comandos más elaborados; no obstante, esa ganancia fue acompañada de un mayor tiempo de respuesta e impacto en la satisfacción del usuario²⁴. Se refuerza que, en el contexto clínico, es necesario validar también otros criterios, como la rapidez del sistema, la transparencia con la IA y la usabilidad de la técnica de *prompting*²⁵. Así, se tornan esenciales estudios de validación y evaluación que verifiquen si la IA presenta un desempeño adecuado para la síntesis de evidencias y contribuye de manera segura a la práctica profesional²⁶.

Existe una laguna en la literatura, especialmente en lo que se refiere a la escasez de referencias que aborden, de forma profunda y sistematizada, las técnicas de *prompting* aplicadas al manejo de la LP. Tal laguna restringe las comparaciones con la literatura y evidencia un campo aún en desarrollo. No obstante, otros campos clínicos han sido abordados internacionalmente, como un estudio desarrollado en Suecia para la evaluación del desempeño de LLM en la identificación de errores de medicación, demostrando que los mecanismos de verificación interna (*critic model*) aumentan la confiabilidad de la respuesta de la IA²⁷ y un estudio coreano que comparó diferentes estrategias de ingeniería de *prompts*, identificando el método *Chain-of-Thought* con respuestas más consistentes en el contexto terapéutico de la enfermería psiquiátrica²⁸.

En relación con las indicaciones de apósitos primarios, las IA presentaron distintas posibilidades. Aunque el alginato de calcio, las espumas y los apósitos con plata sean compatibles con heridas en las que hay pérdida de espesor total, tales opciones no correspondían a la lesión analizada²⁰. Aunque el instrumento de referencia no incluyera la hidrofibra como apósito primario, los modelos de IA la citaron con frecuencia. La literatura internacional reconoce esa opción como adecuada por favorecer el desbridamiento autolítico, indicado para remover tejido desvitalizado²⁹, compatible con el esfacelo en pequeña cantidad observado en la imagen.

Las IA tienen potencial para apoyar a los profesionales de salud, aumentando la eficiencia clínica y la seguridad del paciente¹. Sin embargo, competencias humanas como el razonamiento clínico, la toma de decisión y la interpretación permanecen insustituibles. Así, las recomendaciones de las IA deben ser evaluadas por el profesional antes de ser aplicadas⁷. Aunque sea pertinente problematizar las posibles causas del elevado número de errores de determinadas IA, como el ChatGPT-4.0®, no es posible, a partir de los datos de este estudio, inferir con precisión los factores determinantes

de esas fallas. Por lo tanto, el uso de IA en el manejo de la LP debe entenderse como una herramienta complementaria, preservando la autonomía del enfermero en las decisiones.

Cabe destacar que existe la necesidad apremiante de incorporar la alfabetización en inteligencia artificial en la formación y en la educación permanente de los profesionales de la salud, como destaca una revisión sistemática que investigó el nivel de alfabetización en IA entre profesionales de salud y estudiantes, en la que la mitad de los estudios incluidos encontró una alfabetización muy baja entre ambos³⁰. En ese sentido, no basta con solo utilizar la IA, sino desarrollar competencias para formular preguntas claras, contextualizadas y clínicamente relevantes, además de interpretar y validar los resultados generados. La ausencia de ese dominio puede llevar a respuestas imprecisas, equivocadas o potencialmente inseguras, incluso cuando se utilizan modelos aparentemente confiables. Así, la alfabetización en IA emerge como un componente esencial para el uso seguro, ético y efectivo de esas tecnologías en el manejo de lesiones por presión y en otros contextos asistenciales.

Este estudio presenta limitaciones en relación con la realización única de cada consulta a la IA, lo que limita la evaluación de la variabilidad de las respuestas, dado que los modelos generativos pueden producir resultados distintos en cada interacción. Además, el análisis se realizó sobre un número restringido de modelos de IA y versiones específicas disponibles en el período del estudio, lo que puede no reflejar el desempeño de herramientas pagas o de otras plataformas.

Este estudio no evaluó la concordancia interevaluadores, la consistencia de los modelos ni el tiempo de respuesta de cada IA. De esa forma, la tasa de precisión en cada IA puede verse influenciada por diferentes interpretaciones de los especialistas y profesionales de salud acerca de lo que se consideraría una respuesta correcta, lo que introduce variabilidad en la evaluación y se presenta como otra limitación del estudio. Además, no se midió la consistencia de las respuestas de los modelos, considerando que una IA puede parecer confiable y producir respuestas estables y, aun así, presentar errores sistemáticos desde la perspectiva clínica, lo que representa un riesgo en contextos asistenciales. El estudio tampoco consideró el tiempo de respuesta de los modelos, aspecto relevante para la aplicabilidad clínica, dado que los modelos más ágiles, aunque menos precisos, pueden ser más viables y presentar mejor usabilidad en el flujo de trabajo.

De esa forma, estos son puntos que orientan hacia estudios futuros que amplíen la noción de seguridad más allá de la precisión, incorporando también dimensiones como la consistencia, la usabilidad y la eficiencia en el contexto real de cuidado. Se recomienda, por lo tanto, que, en la evaluación de la aplicación de modelos de IA en la salud, se realicen análisis conjuntos de exactitud y consistencia, especialmente en el contexto multimodal del manejo de lesiones por presión, en el cual la complejidad interpretativa es elevada.

CONCLUSIÓN

La precisión de las IA relacionadas con el manejo de la LP, del 31,9%, fue considerada baja; no obstante, la IA Gemini® demostró la mayor precisión y el ChatGPT-4.0®, la peor. Aunque las inteligencias artificiales no hayan sido originalmente desarrolladas para el área de la salud, la precisión presentada indica que la recomendación más segura es el modelo *hybrid prompting*, integrando la inteligencia artificial en el juicio clínico humano.

Al evidenciar tanto potencialidades como limitaciones relevantes, se torna claro que la IA debe comprenderse como una herramienta complementaria en la práctica de

enfermería, exigiendo responsabilidad y uso consciente para asegurar que su aplicación fortalezca el cuidado, sin comprometer la seguridad del paciente y preservando el papel central del profesional de salud en la toma de decisiones respecto al manejo de las lesiones.

En un escenario marcado por la creciente utilización de esos modelos, es fundamental ampliar el debate sobre la integración responsable de la IA en el cuidado y señalar la necesidad de que futuras investigaciones avancen en el desarrollo de metodologías más robustas. En ese sentido, se torna esencial explorar estrategias de *prompting* en diferentes contextos asistenciales, contribuyendo al perfeccionamiento científico y ético del uso de esas tecnologías en la enfermería. Tal abordaje puede evitar la adopción indiscriminada de herramientas, cuya aplicabilidad en salud aún demanda más investigación y validación.

REFERENCIAS

1. Pott FS, Meier MJ, Stocco JGD, Petz FFC, Roehrs H, Ziegelmann PK. Pressure injury prevention measures: overview of systematic reviews. Rev Esc Enferm USP [Internet]. 2023 [cited 2025 Aug 15];57:e20230039. Available from: <https://doi.org/10.1590/1980-220X-REEUSP-2023-0039en>
2. Fontes TLA, de Oliveira BGRB, De Oliveira MF, da Silva MA, da Silva ARG, Pires BMFB, et al. Development of a checklist to prevent pressure injuries in patients with COVID-19. Rev enferm UFPE on line [Internet]. 2024 [cited 2025 Aug 15];18(1):e257602. Available from: <https://doi.org/10.5205/1981-8963.2024.257602>
3. Macêdo SM, Bastos LLAG, Oliveira RGC, Lima MCV, Gomes FCF. Selection criteria for primary dressings in the treatment of pressure ulcers in hospitalized patients. Cogitare Enferm [Internet]. 2021 [cited 2025 Aug 25];26:e74400. Available from: <http://dx.doi.org/10.5380/ce.v26i0.74400>
4. de Araújo CAF, Pereira SRM, de Paula VG, de Oliveira JA, de Andrade KBS, de Oliveira NVD. Evaluation of the knowledge of nursing professionals in the prevention of pressure ulcer in intensive care. Esc Anna Nery [Internet]. 2022 [cited 2025 Aug 25];26:e20210200. Available from: <https://doi.org/10.1590/2177-9465-EAN-2021-0200>
5. Kotp MH, Ismail HA, Basyouny HAA, Aly MA, Hendy A, Nashwan AJ, et al. Empowering nurse leaders: readiness for AI integration and the perceived benefits of predictive analytics. BMC Nursing [Internet]. 2025 [cited 2025 Aug 25];24(56):1-13. Available from: <https://doi.org/10.1186/s12912-024-02653-x>
6. Dal Sasso GTM, Lanzoni GMM, Alvarez AG, Barra DCC, Barbosa SFF. Potential contribution of Chatgpt® to learning about septic shock in intensive care. Texto Contexto Enferm [Internet]. 2024 [cited 2025 Aug 25];33:e20230184. Available from: <https://doi.org/10.1590/1980-265X-TCE-2023-0184en>
7. Vitorino LM, Yoshinari Junior GH, Lopes-Júnior LC. Artificial intelligence in nursing: advancing clinical judgment and decision-making. Rev Bras Enferm [Internet]. 2025 [cited 2025 Nov 19];78(4):e780401. Available from: <https://doi.org/10.1590/0034-7167.2025780401>
8. Seibert K, Domhoff D, Bruch D, Schulte-Althoff M, Fürstenau D, Biessmann F, et al. Application scenarios for artificial intelligence in nursing care: rapid review. J Med Internet Res [Internet]. 2021 [cited 2025 Aug 22];23(11):e26522. Available from: <https://www.jmir.org/2021/11/e26522>
9. Sampaio RC, Sabbatini M, Limongi R. Diretrizes para o uso ético e responsável da Inteligência Artificial Generativa: um guia prático para pesquisadores. São Paulo: Editora Intercom [Internet]. 2024 [cited 2025 Aug 20]. 62 p. Available from: <https://prpg.unicamp.br/noticias/lancamento-diretrizes-para-o-uso-etico-e-responsavel-da-inteligencia-artificial-generativa-um-guia-pratico-para-pesquisadores/>
10. Shi J, Chen H, Shi C, Su C, Yue S, Li W, et al. Development and comparative evaluation of knowledge graph-enhanced large language models for domain-specific question answering in nursing. BMC Nursing [Internet]. 2026 [cited 2026 May 5];25:530. Available from: <https://doi.org/10.1186/s12912-026-04647-3>

11. Gurler N, Sengul T, Bulbul SH, Akyaz DY, Guler O, Yantac AE, et al. Secure and accessible AI in nursing education: a modular agentic chatbot framework comparing ChatGPT-4o with an open-source LLM for chronic wound care. *Nurse Educ Pract* [Internet]. 2026 [cited 2026 May 5];93:104789. Available from: <https://doi.org/10.1016/j.nepr.2026.104789>
12. Bsharat SM, Myrzakhan A, Shen Z. Principled instructions are all you need for questioning LLaMA-1/2, GPT-3.5/4. *arXiv* [Internet]. 2024 [cited 2025 Sep 19];2:1-26. Available from: <https://arxiv.org/html/2312.16171v2>
13. Pan American Health Organization (PAHO). AI Prompt design for public health: using generative AI responsibly [Internet]. OPAS; 2025. [cited 2025 Sep 30]. 61 p. Available from: <https://iris.paho.org/items/b98d9c0f-6008-4b78-9066-01e1e81c29c5>
14. Schulhouff S, Llie M, Balepur N, Kahadze K, Liu A, Si C, et al. The prompt report: a systematic survey of prompt engineering techniques. *arXiv* [Internet]. 2024 [cited 2025 Sep 30];6:1-80. Available from: <https://arxiv.org/abs/2406.06608>
15. Vilakati, S. Prompt engineering for accurate statistical reasoning with large language models in medical research. *Front Artif Intell* [Internet]. 2025 [cited 2025 Nov 18];8:1658316. Available from: <https://doi.org/10.3389/frai.2025.1658316>
16. Dhuliawala S, Komeili M, Xu J, Raileanu R, Li X, Celikyilmaz, et al. Chain-of-verification reduces hallucination in large language models. *arXiv* [Internet]. 2023 [cited 2025 Oct 12];2:1-19. Available from: <https://arxiv.org/abs/2309.11495>
17. Bernardes, RM. Prevenção e manejo da lesão por pressão: segurança do paciente [Internet]. São Paulo: Universidade de São Paulo; 2020 [cited 2025 Sep 15]. Available from: http://eerp.usp.br/feridasronicas/recurso_educacional_lp_1_4.html
18. National Pressure Injury Advisory Panel (NPIAP). Prevention and treatment of pressure ulcers/injuries: the international guideline [Internet]. Washington; 2019 [cited 2025 Aug 15]. 405 p. Available from: <https://static1.squarespace.com/static/6479484083027f25a6246fcb/t/6553d3440e18d57a550c4e7e/1699992399539/CPG2019edition-digital-Nov2023version.pdf>
19. Potter PA, Perry AG, Stockert PS. Fundamentos de enfermagem. 11th. Rio de Janeiro: Guanabara Koogan; 2024. 1644 p.
20. Zhang C, Zhang S, Wu B, Zou K, Chen H. Efficacy of different types of dressings on pressure injuries: systematic review and network meta-analysis. *Nurs Open* [Internet]. 2023 [cited 2025 Aug 10];10(9):5857-67. Available from: <https://doi.org/10.1002/nop2.1867>
21. Holman, M. Using tap water compared with normal saline for cleansing wounds in adults: a literature review of the evidence. *J Wound Care* [Internet]. 2023 [cited 2025 Oct 10];32(8):507-12. Available from: <https://doi.org/10.12968/jowc.2023.32.8.507>
22. dos Santos TO, da Silva JVL, Rocha RG, Assad LG, da Costa CCP, Pires BMFB. Papain with urea cream in pressure injuries: a case series study. *Rev Enferm Atenção Saúde* [Internet]. 2024 [cited 2025 Oct 10];13(1):e202404. Available from: <https://seer.uftm.edu.br/revistaelectronica/index.php/enfer/article/view/6950>
23. O'Connor S, Peltonen LM, Topaz M, Chen LYA, Michalowski M, Ronquillo C, et al. Prompt engineering when using generative AI in nursing education. *Nurse Educ Pract* [Internet]. 2024 [cited 2025 Nov 19];74:103825. Available from: <https://doi.org/10.1016/j.nepr.2023.103825>
24. Wang MH, Jiang X, Zeng P, Li X, Chong KKL, Hou G, et al. Balancing accuracy and user satisfaction: the role of prompt engineering in AI-driven healthcare solutions. *Front Artif Intell* [Internet]. 2025 [cited 2025 Oct 19];8:1517918. Available from: <https://doi.org/10.3389/frai.2025.1517918>
25. Wang M, Cui J, Lee SMY, Lin Z, Zeng P, Li X, et al. Applied machine learning in intelligent systems: knowledge graph-enhanced ophthalmic contrastive learning with “clinical profile” prompts. *Front Artif Intell* [Internet]. 2025 [cited 2025 Oct 19];8:1527010. Available from: <https://doi.org/10.3389/frai.2025.1527010>

26. Flemyng E, Noel-Storr A, Macura B, Gartlehner G, Thomas J, Meerpohl JJ, et al. Position statement on artificial intelligence (AI) use in evidence synthesis across Cochrane, the Campbell Collaboration, JBI, and the collaboration for environmental evidence 2025. JBI Evid Synth [Internet]. 2025 [cited 2025 Oct 10];23(11):2162-6. Available from: <https://doi.org/10.11124/JBIES-25-00480>
27. Krifors A, Beskow T, Jonsson M, Lindner KJ, Calås J, Arwsbo V, et al. Improving medication error classification using a reasoning large language model. JAMIA Open [Internet]. 2026 [cited 2025 May 5];9(1):ooag004. Available from: <https://doi.org/10.1093/jamiaopen/ooag004>
28. Kim SK, Kim GM, Cha S, Jung M. Verification of the validity and reliability of therapeutic communication responses generated by large language models (LLMs): a comparative study of prompt-engineering strategies. J Korean Acad Psychiatr Ment Health Nurs [Internet]. 2025 [cited 2025 May 5];34(Spec):23-35. Available from: <https://doi.org/10.12934/jkpmhn.2025.34.S1.23>
29. Klein A, Ayoub N, Juhel C, Schuller R, Armstrong F, Pegalajar-Jurado A. Performance of a gelling fibre dressing in management of wounds in a community setting: a sub-analysis of the VIPES study. J Wound Care [Internet]. 2024 [cited 2025 Nov 19];33(7):464-73. Available from: <https://doi.org/10.12968/jowc.2024.0125>
30. Kimiafar K, Sarbaz M, Tabatabaei SM, Ghaddaripouri K, Mousavi AS, Raei M, et al. Artificial intelligence literacy among healthcare professionals and students: a systematic review. Front Health Inform [Internet]. 2023 [cited 2026 May 5];12:168. Available from: https://www.researchgate.net/publication/375600294_Artificial_Intelligence_Literacy_Among_Healthcare_Professionals_and_Students_A_Systematic_Review

Responses and accuracy of artificial intelligence systems to prompts on pressure injury management**ABSTRACT**

Objective: To compare the responses and accuracy of different Artificial Intelligences based on a set of prompts for the management of pressure injury. **Method:** A descriptive, comparative study of the responses of Artificial Intelligences, carried out between October and November 2025, in two stages: 1) Study of prompts and 2) Test to obtain the responses. The following were evaluated: Perplexity AI, Grok AI, Blackbox AI, OpenAI ChatGPT-4.0®, Gemini®, and Claude AI, randomly divided into the hybrid prompting and zero-shot groups. The responses were organized in an Excel spreadsheet and analyzed descriptively. **Results:** The overall accuracy of the artificial intelligences was 31.9%. The hybrid prompting group reached 36.1% and the zero-shot obtained 27.7%. The Gemini® AI demonstrated the highest accuracy and ChatGPT-4.0® the worst. **Conclusion:** The accuracy presented indicates that the safest recommendation is the hybrid prompt model, integrating artificial intelligence with human clinical judgment.

DESCRIPTORS: Generative Artificial Intelligence; Large Language Models; Nursing, Team; Clinical Reasoning; Pressure Ulcer.

Recibido en: 08/12/2025

Aprobado en: 07/05/2026

Editor asociado: Dra. Juliana Balbinot Reis Girondi

Autor correspondiente:

Letícia Pereira Pires

Universidade Federal Fluminense

Rua Recife, Lotes 1-7 - Jardim Bela Vista, Rio das Ostras - RJ, 28895-532

E-mail: leticiapires@id.uff.br

Contribución de los autores:

Contribuciones sustanciales a la concepción o diseño del estudio; o la adquisición, análisis o interpretación de los datos del estudio -

Pires LP, Goulart MCL, Góes FGB, Ribeiro YC, Guimarães LLC, de Sousa LF. Elaboración y revisión crítica del contenido intelectual del estudio - **Pires LP, Goulart MCL, Góes FGB, Ribeiro YC, Guimarães LLC, de Sousa LF.** Responsable de todos los aspectos del estudio, asegurando las cuestiones de precisión o integridad de cualquier parte del estudio - **Pires LP, Goulart MCL, Góes FGB, Ribeiro YC, Guimarães LLC, de Sousa LF.** Todos los autores aprobaron la versión final del texto.

Conflicto de intereses:

Los autores no tienen conflictos de intereses que declarar.

Disponibilidad de datos:

Los autores declaran que todos los datos están completamente disponibles en el cuerpo del artículo.

ISSN 2176-9133



Esta obra está bajo una [Licencia Creative Commons Atribución 4.0 Internacional](https://creativecommons.org/licenses/by/4.0/).