

## **Análise do potencial do ChatGPT-3 como ferramenta auxiliar na prescrição medicamentosa em odontopediatria**

**Michelle Alonso Coutinho**

Pós-graduação em Cirurgia e Traumatologia Bucomaxilofacial, UNIFESO, Teresópolis/RJ, Brasil.  
Mestre em Desordens Temporomandibular e Dor Orofacial, São Leopoldo Mandic - Campinas/SP, Brasil.  
Programa de Pós-Graduação, Faculdade de Odontologia, Universidade Veiga de Almeida, Rio de Janeiro/RJ, Brasil

✉ [coutinho.ma@yahoo.com](mailto:coutinho.ma@yahoo.com)

**Hislene Oliveira Ramos**

Faculdade de Odontologia, Universidade Veiga de Almeida, Rio de Janeiro/RJ, Brasil.

**Andréa Vaz Braga Pintor**

Programa de Pós-Graduação, Faculdade de Odontologia, Universidade Veiga de Almeida, Rio de Janeiro/RJ, Brasil.

✉ [andrea\\_pintor@hotmail.com](mailto:andrea_pintor@hotmail.com)

**Marcela Baraúna Magno**

Programa de Pós-Graduação, Faculdade de Odontologia, Universidade Veiga de Almeida, Rio de Janeiro/RJ, Brasil. Departamento de Odontopediatria e Ortodontia, Universidade Federal do Rio de Janeiro, Rio de Janeiro/RJ, Brasil.

✉ [marcela.magno@hotmail.com](mailto:marcela.magno@hotmail.com)

Recebido em 4 de julho de 2025

Aceito em 11 de novembro de 2025

### **Resumo:**

Avaliou-se o potencial do ChatGPT-3 (cGPT-3) como ferramenta auxiliar na terapêutica medicamentosa em odontopediatria, por meio da análise da qualidade das respostas fornecidas às 45 perguntas sobre a forma terapêutica mais indicada, analgésicos, anti-inflamatórios e anestésicos locais, em diferentes cenários clínicos. As respostas foram analisadas e pontuadas quanto à precisão (0 a 5 pontos), completude (0 a 2 pontos), objetividade (0 a 2 pontos) e pontuação total (0 a 9 pontos). Os grupos de medicamentos foram comparados pelo teste de Kruskal-Wallis ( $\alpha=5\%$ ). A média  $\pm$  desvio padrão da pontuação total foi de  $3,71 \pm 2,58$ , e a mediana  $\pm$  desvio interquartil dos escores de precisão, completude e objetividade foi de  $3,0 \pm 3,0$ ,  $1,0 \pm 1,0$  e  $0,0 \pm 0,0$ , respectivamente. De forma geral, a maioria das respostas foi considerada 'Completamente incorreta' ( $n=12$ , 26,6%) quanto a sua precisão, 'incompleta' ( $n=23$ , 51,1%) quanto a sua completude, e com 'muita quantidade (mais de 50%) de informações extras' ( $n=44$ , 95,6%) quanto a objetividade. Independentemente do grupo de medicamento, a maioria das respostas não foram objetivas e apresentaram 'muita quantidade (>50%) de informações extras', não influenciando nos escores de objetividade ( $p=0,502$ ). Entretanto, os escores de precisão ( $p=0,009$ ) e completude ( $p=0,020$ ), bem como a média da pontuação total ( $p=0,016$ ), foram menores para as respostas sobre anestésicos locais, em comparação as de analgésicos. Pode-se concluir que, clinicamente, o ChatGPT-3 deve ser utilizado com cautela e não deve ser a única fonte de informação utilizada sobre terapêutica medicamentosa em odontopediatria.

**Palavras-chave:** Inteligência artificial, modelos de linguagem na clínica, odontopediatria, prescrições de medicamentos, tomada de decisão clínica.

## Analysis of the potential of ChatGPT-3 as an auxiliary tool in drug prescription in pediatric dentistry

### Summary:

The potential of ChatGPT-3 (cGPT-3) as an auxiliary tool in drug therapy in pediatric dentistry was evaluated, through the quality of the answers provided to the 45 questions about the most indicated therapeutic form, analgesics, anti-inflammatories and local anesthetics in different clinical scenarios. The responses were analyzed and scored for accuracy (0 to 5 points), completeness (0 to 2 points), objectivity (0 to 2 points), and total score (0 to 9 points). The drug groups were compared using the Kruskal-Wallis test ( $\alpha=5\%$ ). The mean  $\pm$  standard deviation of the total score was  $3.71\pm2.58$ , and the median  $\pm$  interquartile deviation of the accuracy, completeness, and objectivity scores was  $3.0\pm3.0$ ,  $1\pm1.0$ , and  $0.0\pm0.0$ , respectively. In general, most of the answers were considered 'Completely incorrect' ( $n=12$ , 26.6%) in terms of accuracy, 'incomplete' ( $n=23$ , 51.1%) in terms of completeness, and with 'a lot (more than 50%) of extra information' ( $n=44$ , 95.6%) in terms of objectivity. Regardless of the drug group, most of the answers were not objective and presented 'a lot (>50%) of extra information', not influencing the objectivity scores ( $p=0.502$ ). However, the accuracy ( $p=0.009$ ) and completeness ( $p=0.020$ ) scores, as well as the mean total score ( $p=0.016$ ), were lower for the answers on local anesthetics, compared to those on analgesics. It can be concluded that, clinically, ChatGPT-3 should be used with caution and should not be the only source of information used on drug therapy in pediatric dentistry.

**Keywords:** Artificial intelligence, language models in the clinic, pediatric dentistry, drug prescriptions, clinical decision-making.

## Análisis del potencial de ChatGPT-3 como herramienta auxiliar en la prescripción de fármacos en odontopediatría

### Resumen:

Se evaluó el potencial de ChatGPT-3 (cGPT-3) como herramienta auxiliar en la terapia farmacológica en odontopediatría, a través de la calidad de las respuestas proporcionadas a las 45 preguntas sobre la forma terapéutica más indicada, analgésicos, antiinflamatorios y anestésicos locales en diferentes escenarios clínicos. Las respuestas se analizaron y calificaron por precisión (0 a 5 puntos), integridad (0 a 2 puntos), objetividad (0 a 2 puntos) y puntuación total (0 a 9 puntos). Los grupos de fármacos se compararon mediante la prueba de Kruskal-Wallis ( $\alpha=5\%$ ). La media  $\pm$  desviación estándar de la puntuación total fue de  $3,71\pm2,58$ , y la mediana  $\pm$  desviación intercuartílica de las puntuaciones de precisión, completitud y objetividad fue de  $3,0\pm3,0$ ,  $1\pm1,0$  y  $0,0\pm0,0$ , respectivamente. En general, la mayoría de las respuestas se consideraron "Completamente incorrectas" ( $n = 12$ , 26,6%) en términos de precisión, "incompletas" ( $n = 23$ , 51,1%) en términos de completitud y con "mucho (más del 50%) de información adicional" ( $n = 44$ , 95,6%) en términos de objetividad. Independientemente del grupo de drogas, la mayoría de las respuestas no fueron objetivas y presentaron "mucho (>50%) información adicional", no influyendo en los puntajes de objetividad ( $p = 0,502$ ). Sin embargo, las puntuaciones de precisión ( $p = 0,009$ ) y completitud ( $p = 0,020$ ), así como la puntuación total media ( $p = 0,016$ ), fueron menores para las respuestas sobre anestésicos locales, en comparación con las de analgésicos. Se puede concluir que, clínicamente, ChatGPT-3 debe usarse con precaución y no debe ser la única fuente de información utilizada sobre la terapia farmacológica en Odontopediatría.

**Palabras clave:** Inteligencia artificial, modelos de lenguaje en la clínica, odontopediatría, prescripción de medicamentos, toma de decisiones clínicas.

## INTRODUÇÃO

Em 1950, Alan Turing acreditava que no final do século XX seria possível falar de máquinas pensantes e sugeriu que os seres humanos poderiam resolver problemas utilizando informações disponíveis e raciocínio lógico, e que as máquinas poderiam fazer o mesmo (Turing, 1950). Deste pensamento, foi proposto o termo inteligência artificial (IA) (McCarthy *et al.*, 1955). Com o avanço da tecnologia, surgiram os *chatbot*, programas de computador capazes de manter uma conversa com um usuário humano por meio de aplicativos de mensagens utilizando uma interface conversacional através da parametrização de palavras-chave. Com isso, as pessoas perceberam o potencial dessa tecnologia para alcançar diferentes perspectivas (Çakmakoglu, 2023; Huang *et al.*, 2023). O *Chat Generated Pre-Trained Transformer-3* (ChatGPT-3, Chat transformador pré-treinado generativo-3) é um *chatbot* de IA desenvolvido pela empresa OpenAI, lançado em novembro de 2022, que apresenta modelos de linguagem grande (do inglês, *large language models* - LLMs) (ChatGPT, 2024). LLMs são aplicações de IA treinadas em centenas de *terabytes* de dados textuais. Este *chatbot* está sendo bastante estudado para auxiliar na educação de uma forma geral, por ser de uso gratuito e acessível online (Eggmann *et al.*, 2023).

Na capacitação de cirurgiões-dentistas para o atendimento pediátrico, entende-se que o sucesso da anestesia local é de suma importância, pois alivia a ansiedade, promove a confiança e contribui para uma experiência segura e positiva. No contexto odontopediátrico, os critérios de seleção para anestésicos locais e medicamentos fundamentam-se no conhecimento sobre o fármaco indicado para o caso clínico, bem como na sua dosagem calculada em função da superfície ou peso corporal da criança, com o objetivo de se obter o efeito máximo da medicação, com total ausência ou mínimos efeitos adversos (relação risco/benefício) (Andrade, 2014). Ademais, na odontopediatria, a prescrição medicamentosa com fins profiláticos ou terapêuticos, requer que o profissional considere o desenvolvimento dos órgãos e tecidos da criança, bem como as particularidades do paciente e a gravidade do processo inflamatório e/ou infeccioso. O cirurgião-dentista, portanto, deve estar apto a escolher e prescrever o medicamento com melhor ação e menor toxicidade para os pacientes infantis (Scarparo, 2021).

No entanto, estudos anteriores reportam que o conhecimento de dentistas sobre o uso e prescrição de medicamentos para crianças varia de baixo (Rubanenko, 2021) a moderado

(Capan, Duman e Kalaoglu, 2023), e que o cálculo da dose do medicamento é uma das possíveis barreiras encontradas. Diante desse cenário, a IA poderia ser uma ferramenta em potencial para auxiliar cirurgiões-dentistas neste processo. A precisão, abrangência e/ou utilidade das respostas geradas pelo ChatGPT em saúde foram avaliados previamente (Johnson *et al.*, 2023; Mago e Sharma, 2023), entretanto, a sua aplicação na terapêutica medicamentosa em odontopediatria não foi relatada. Assim, com vistas a sanar essa lacuna do conhecimento, o objetivo deste trabalho foi avaliar o potencial do ChatGPT-3 como uma ferramenta de uso clínico a terapêutica medicamentosa em odontopediatria, através da avaliação da qualidade das respostas fornecidas pelo chat.

## METODOLOGIA

### Procedimento de role-playing com IA

O presente estudo utilizou a técnica de *role-playing* com a inteligência artificial GPT-3 (OpenAI), acessada via a interface do ChatGPT. O objetivo foi simular uma situação clínica na qual um cirurgião-dentista busca orientações sobre prescrição medicamentosa para pacientes pediátricos. A interação entre o usuário e o chat incluiu três etapas:

1) Introdução do papel e cenário clínico - Informar para o chat GPT-3 quem está falando com ele e a situação – “Eu sou um dentista e estou com dúvidas sobre a prescrição de medicamentos para uma criança. Pode me ajudar?”,

2) Orientação de base teórica - Detalhe - “Incorpore orientações reportadas por Eduardo Dias de Andrade em seu livro *Terapêutica Medicamentosa em Odontologia*” (Andrade, 2014) (figura 1).

3) Execução da tarefa – as perguntas eram feitas de forma direta e suas respostas inseridas na tabela para posterior avaliação dos especialistas.

**Figura 1** – *Role-playing* realizado no CHATGPT-3

### **PERGUNTAS E RESPOSTAS DO CHATGPT – 3**

#### **1) Olá, bom dia**

Oi! Bom dia! Como posso ajudar você hoje?

#### **2) “Eu sou um dentista e estou com dúvidas sobre a prescrição de medicamentos para uma criança. Pode me ajudar?”**

Claro, ficarei feliz em ajudar com suas dúvidas sobre a prescrição de medicamentos para uma criança. No entanto, é importante lembrar que não sou um profissional de saúde e não posso fornecer aconselhamento médico específico. Posso oferecer informações gerais e diretrizes gerais para considerar ao prescrever medicamentos a crianças. Por favor, compartilhe suas dúvidas e estarei à disposição para fornecer orientações gerais. No entanto, é sempre recomendável consultar um médico pediatra ou especialista em medicina infantil para obter aconselhamento específico para o caso de seu paciente.

#### **3) Incorpore orientações reportadas por Eduardo Andrade em seu livro *Terapêutica Medicamentosa em Odontologia***

Peço desculpas, mas meu conhecimento se limita a informações disponíveis até setembro de 2021 e não tenho acesso direto a livros específicos, como "Terapêutica Medicamentosa em Odontologia" de Eduardo Andrade. No entanto, posso oferecer algumas diretrizes gerais sobre prescrição de medicamentos em odontologia pediátrica.

Ao prescrever medicamentos para crianças, considere o seguinte:

1. **Peso e Idade:** A dose de um medicamento muitas vezes é baseada no peso e na idade da criança. Certifique-se de calcular a dose apropriada com base nessas informações.
2. **Tipo de Medicamento:** Alguns medicamentos são mais seguros e apropriados para crianças do que outros. O uso de medicamentos específicos depende da condição a ser tratada.
3. **Forma Farmacêutica:** Medicamentos líquidos são geralmente mais fáceis de administrar a crianças pequenas do que comprimidos ou cápsulas.
4. **Interações Medicamentosas:** Esteja ciente de possíveis interações entre o medicamento prescrito e outros medicamentos que a criança possa estar tomando.
5. **Alergias:** Certifique-se de perguntar sobre alergias conhecidas a medicamentos antes de prescrever qualquer substância.
6. **Orientação dos Pais ou Responsáveis:** Forneça orientações claras sobre a administração do medicamento, incluindo dosagem, frequência e duração do tratamento. Certifique-se de que os pais ou responsáveis compreendam as instruções.
7. **Acompanhamento:** Acompanhe o progresso da criança durante o tratamento e esteja disponível para responder a quaisquer dúvidas ou preocupações dos pais ou responsáveis.

Lembrando que a prescrição de medicamentos deve ser feita por um profissional de saúde licenciado e com conhecimento em odontologia pediátrica. Se você deseja informações específicas do livro de Eduardo Andrade, sugiro consultar o livro diretamente ou entrar em contato com outros profissionais que possam ter experiência com as orientações apresentadas no livro.

**Fonte:** Autoria própria.

### Elaboração das perguntas

Quarenta e cinco perguntas sobre terapêutica medicamentosa em odontopediatria foram elaboradas pelos autores (equipe constituída de 1 aluno de graduação, 2 doutores em odontopediatria e 1 especialista e mestre em cirurgia buco-maxilo facial). As perguntas foram elaboradas e categorizadas conforme os grupos de medicamentos: anti-inflamatórios (n=17), analgésicos (n=17), e anestésicos locais (n=10). Uma pergunta (n=1) sobre a forma farmacêutica mais adequada para uso em odontopediatria foi elaborada.

Para assegurar que as questões estavam redigidas em linguagem clara e compatível com o nível de escolaridade de ensino superior completo (simulando perguntas de um dentista, conforme etapa 1 do role-playing), maximizando a compreensão pela IA, a legibilidade das questões foi avaliada através do *Flesch Kincaid readability test* ([Good Calculators, 2023](#)) (Kincaid *et al.*, 1975).

### Avaliação das respostas

A avaliação quanto a precisão, completude e objetividade das respostas fornecidas pela IA foi feita por dois especialistas, ambos doutores em odontopediatria, com 9 ou mais anos de docência. A avaliação dos especialistas utilizou Escalas Likert com diferentes intervalos foram consideradas para avaliar a precisão, completude e objetividade das respostas do ChatGPT-3. As escalas utilizadas no presente estudo foram baseadas em escalas previamente descritas na literatura (Johnson *et al.* 2023), onde a modificação realizada foi a inclusão de uma pontuação para os casos nos quais a IA não respondeu à pergunta inserida. Os avaliadores que designaram as pontuações foram previamente treinados e calibrados para a aplicação das escalas.

Em relação a sua precisão as respostas poderiam receber escore (Kappa = 0,86):

- 0 – Não respondeu diretamente à pergunta;
- 1 – Completamente incorreto;
- 2 – Mais incorreto que correto;
- 3 – Aproximadamente igual entre correto e incorreto;

4 – Mais correto que incorreto;

5 – Completamente correto.

Em relação a completude as respostas poderiam receber escore (Kappa = 1,0):

0 – Incompleto, aborda alguns aspectos da questão, mas partes significativas estão faltando ou incompleta;

1 – Adequado, aborda todos os aspectos da questão e fornece a quantidade mínima de informações necessárias para ser considerado completo;

2 – Abrangente, aborda todos os aspectos da questão.

Em relação a objetividade as respostas poderiam receber escore (Kappa = 0,80):

0 - Muita quantidade (mais de 50%) de informações extras;

1- Pouca quantidade (entre 50% e 20%) de informações extras;

2 – Resposta objetiva com menos de 20% de informações extras.

A soma dos escores de cada critério resultou em uma pontuação total de 0 a 9 para cada resposta, em que valores mais altos indicam maior qualidade.

### **Análise estatística**

Os resultados foram tabelados e reportados através de análise descritivas. Comparações entre os escores das respostas para precisão (0 a 5), completude (0 a 2) e objetividade (0 a 2) entre os grupos de medicamentos (analgésicos, anti-inflamatórios e anestésicos locais) foram avaliadas através do teste de Kruskal-Wallis ( $\alpha=5\%$ ). A distribuição paramétrica dos dados relacionados a pontuação total das respostas foi analisada através do teste de Shapiro-Wilk e diferenças entre essas pontuações nos grupos de medicamentos (analgésicos, anti-inflamatórios e anestésicos locais) foram analisadas através do teste Kruskal-Wallis ( $\alpha=5\%$ ).

## RESULTADOS

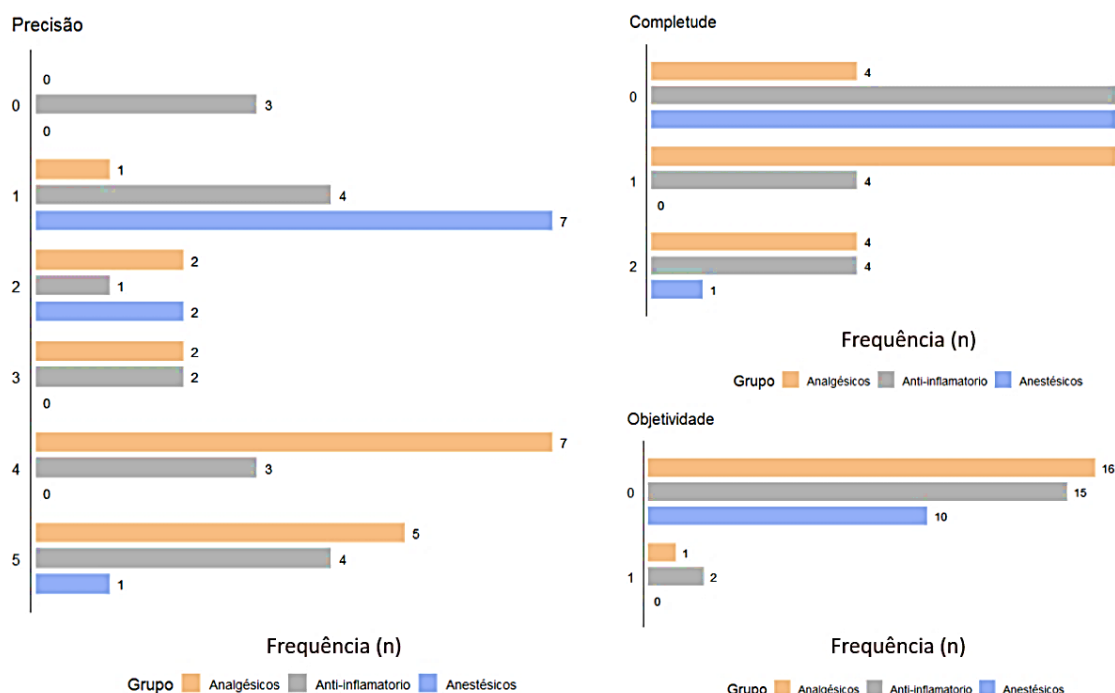
A legibilidade das questões recebeu pontuação entre 0 e 30 (média 22,4, desvio padrão 15,7), indicando que seria necessário Ensino Superior Completo (ESC) para uma facilidade de leitura. A lista de perguntas realizadas e suas respectivas legibilidades, assim como as respostas fornecidas pelo Chat-GPT3 e a avaliação destas quanto à precisão, completude e objetividade, encontram-se descritas na Tabela 1, em anexo a esse artigo (<https://www.e-publicacoes.uerj.br/sustinere/article/view/89795/58702>).

A média  $\pm$  desvio padrão da pontuação total das 45 respostas relacionadas aos medicamentos foi de  $3,71 \pm 2,58$ , e a mediana  $\pm$  desvio interquartílico dos escores de precisão, completude e objetividade foram de  $3,0 \pm 3,0$ ,  $1,0 \pm 1,0$  e  $0,0 \pm 0,0$ , respectivamente.

A resposta para a primeira pergunta, sobre forma farmacêutica mais recomendada na odontopediatria, foi considerada ‘completamente correta’ (escore 5), ‘adequada’ (escore 1) e com ‘muita quantidade (>50%) de informações extras’ (escore 0). Quanto a precisão e completude das 17 respostas relacionadas a analgésicos a maioria foi considerada ‘mais correta que incorreta’ (n=7, 41,18%) ou ‘completamente correta’ (n=5, 29,41%), e ‘adequada’ (n=9, 52,94%). Enquanto sobre anti-inflamatórios, o chat respondeu de forma ‘incompleta’ (n=9, 52,94%) e ‘mais incorreta que correta’ (n=4, 23,5%) ou ‘completamente correta’ (n=4, 23,5%). Vale destacar que o ChatGPT-3 não respondeu ao questionamento em três perguntas sobre anti-inflamatórios. Considerando as respostas sobre anestésicos locais a maioria foi classificada como ‘completamente incorreta’ (n=7, 70%) e ‘incompleta’ (n=9, 90%).

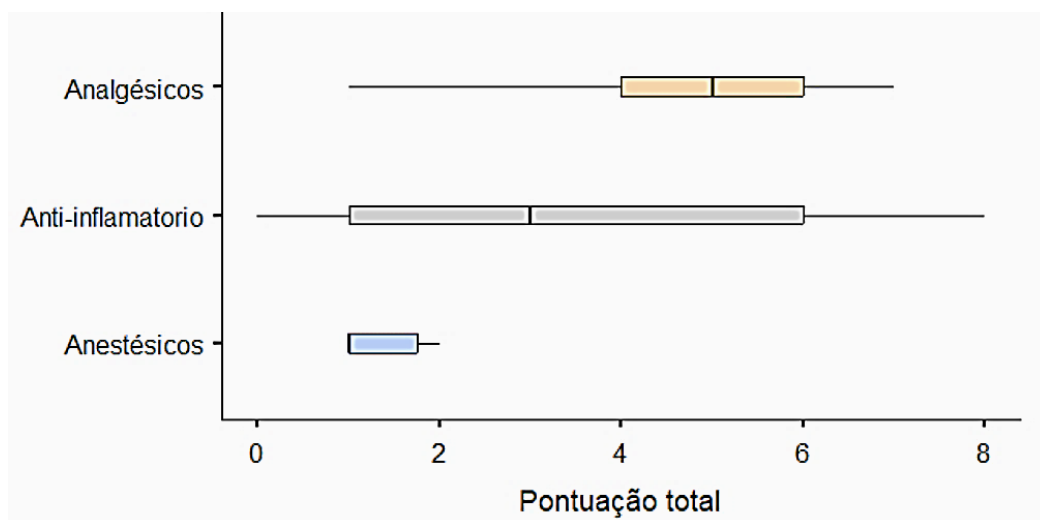
A maioria das respostas não foi objetiva e apresentou ‘muita quantidade (>50%) de informações extras’ tanto sobre analgésicos (n=16, 91,18%), anti-inflamatórios (n=15, 88,23%) e anestésicos locais (n=10, 100%). A figura 2 reporta as frequências absolutas da classificação de cada pergunta, em relação ao grupo de medicamentos.

**Figura 2** – Frequências absolutas da classificação de cada escore, em relação a precisão, completude e objetividade, de acordo com o grupo de medicamentos (analgésicos, anti-inflamatórios e anestésicos).



**Fonte:** Autoria própria.

O grupo do medicamento não influenciou nos escores de objetividade ( $p=0,502$ ), entretanto, os escores de precisão ( $p=0,009$ ) e completude ( $p=0,020$ ), bem como a média da pontuação total ( $p=0,016$ , figura 3), foram menores para as respostas sobre anestésicos locais em comparação as de analgésicos (Tabela 2).

**Figura 3** – Box plot referente a pontuação total das respostas de acordo com o grupo de medicamentos (analgésicos, anti-inflamatórios e anestésicos).

**Fonte:** Autoria própria.

**Tabela 2** – Média  $\pm$  desvio padrão da pontuação total das respostas, e mediana  $\pm$  desvio interquartílico, de acordo com o grupo de medicamento.

| Grupo de medicamento | Pontuação total                | Precisão                       | Compleitude                  | Objetividade  |
|----------------------|--------------------------------|--------------------------------|------------------------------|---------------|
| Analgésicos          | 4,82 $\pm$ 1,85 <sup>a</sup>   | 4,00 $\pm$ 2,00 <sup>a</sup>   | 1,0 $\pm$ 0,0 <sup>a</sup>   | 0,0 $\pm$ 0,0 |
| Anti-inflamatórios   | 3,56 $\pm$ 2,94 <sup>a,b</sup> | 3,00 $\pm$ 3,00 <sup>a,b</sup> | 0,0 $\pm$ 1,0 <sup>a,b</sup> | 0,0 $\pm$ 0,0 |
| Anestésicos          | 1,80 $\pm$ 1,87 <sup>b</sup>   | 1,00 $\pm$ 0,75 <sup>b</sup>   | 0,0 $\pm$ 0,0 <sup>b</sup>   | 0,0 $\pm$ 0,0 |
| p valor              | <b>0,016</b>                   | <b>0,009</b>                   | <b>0,020</b>                 | <b>0,502</b>  |

Letras sobrescritas (a, b) diferentes denotam diferenças estatísticas entre os grupos ( $p < 0,05$ )

**Fonte:** Autoria própria.

## DISCUSSÃO

A aplicação de modelos de linguagem de grande porte (LLMs) na saúde, incluindo a odontopediatria, tem crescido significativamente, oferecendo suporte teórico a cirurgiões-dentistas com dúvidas sobre prescrição medicamentosa em pacientes pediátricos (ANDRADE,

2014; CAPAN *et al.*, 2023; EGGMANN *et al.*, 2023; HUANG *et al.*, 2023; MACHADO *et al.*, 2024; RUBANENKO *et al.*, 2021). No entanto, persistem questionamentos sobre a validade e os limites do conhecimento gerado por essas ferramentas no contexto clínico.

O presente estudo elaborou 45 perguntas abertas sobre indicação, contraindicação, posologia, dose máxima e efeitos adversos de analgésicos, anti-inflamatórios e anestésicos locais, incluindo os mais utilizados na prevenção e controle da dor em procedimentos odontopediátricos (Ferreira e Lopes, 2016). As perguntas elaboradas neste estudo foram discursivas com o objetivo de simular dúvidas comuns no cotidiano clínico de cirurgiões-dentistas, visto que estudos prévios sugerem que o desempenho do ChatGPT-3 em perguntas discursivas é comparável ao de questões objetivas na área médica (Johnson *et al.*, 2023).

Identificou-se uma maior prevalência de respostas classificadas como ‘completamente incorretas’, ‘incompletas’ e com ‘muita quantidade de informações extras’, com desempenho particularmente inferior em temas relacionados a anestésicos locais. Dessa forma, esta IA não pode ser considerada confiável, além de não ser direta com relação as respostas, principalmente para questões relacionadas a anestésicos locais de uso odontológico. Refletindo sobre a complexidade do manejo em terapêutica medicamentosa em odontopediatria e o objetivo desse trabalho, um paralelo com Mira *et al.* (2024) pode ser traçado devido ao baixo grau de concordância entre superespecialistas e o ChatGPT em respostas sobre apneia obstrutiva do sono, considerando que, em ambos, a avaliação humana de aspectos clínicos do paciente é fundamental para a decisão terapêutica do profissional. Observou-se, entretanto, baixa qualidade nas respostas fornecidas pelo ChatGPT, contrastando com literatura que apontam desempenho superior a 83% nesse formato (Machado, Jassé e Arantes; 2024), principalmente, quando consideramos que as perguntas foram classificadas quanto à sua legibilidade. Esta pode ser influenciada por fatores como comprimento e complexidade da frase, vocabulário familiar, comprimento da palavra, voz ativa, entre outros parâmetros (Dubay, 2004). Havendo, o presente trabalho, obtido o nível médio de leitura de  $22,4 \pm 15,7$  através do *Flesch Kincaid readability test*, é possível inferir a necessidade de nível universitário (odontologia, farmácia, medicina) para a compreensão e, conseqüentemente, resolução às questões elaboradas, o que está alinhado com o objetivo deste trabalho.

O excesso de informação encontrada nas respostas pode ser explicado pela complexidade das perguntas realizadas, o que pode fazer o ChatGPT gerar respostas altamente detalhas e com informações externas numa frequência alta, devido a abundância de informações disponíveis na internet, o uso crescente de modelos de linguagem de IA e à sua capacidade para relacionar o prompt de comando aos dados dentro de seu arsenal (Çakmakoglu, 2023; Gilson *et al.*, 2023). Resultados anteriores reportam que o ChatGPT-3 sofre uma diminuição na sua performance quando o nível de dificuldade das perguntas aumenta (Gilson *et al.*, 2023; Johnson *et al.*, 2023). Vale ressaltar que as perguntas foram realizadas apenas em língua portuguesa, e o *chatbot* estudado pode apresentar performance diferente em outras línguas. Estudos futuros em diferentes idiomas são encorajados, para avaliar essa variável.

Os grupos de medicamentos analgésico e anti-inflamatórios apresentaram escores baixos de precisão, completude, objetividade e pontuação total nas respostas do ChatGPT-3, e significativamente semelhantes, indicando que esse modelo de IA não é recomendado como fonte primária de informação para esses medicamentos. Resultados semelhantes foram observados ao avaliar o desempenho do ChatGPT-3 em perguntas sobre radiologia bucomaxilofacial, onde pontuações de mediana e moda foram semelhantes para marcos anatômicos da região da cabeça e pescoço, patologias bucomaxilofaciais e características radiográficas de algumas dessas patologias (Mago e Sharma, 2023). Essa uniformidade de desempenho, mesmo em contextos distintos, demonstra que o ChatGPT-3 não é capaz de oferecer respostas com a precisão e especificidade exigidas para a prática clínica odontológica. Embora alguns estudos relatem resultados mais promissores do modelo em avaliações clínicas de diagnóstico e prognóstico (rastreamento do forame apical, verificação do comprimento de trabalho, projeção de patologias periapicais, morfologias radiculares e previsões de retratamento e descoberta de fraturas radiculares verticais), esses desempenhos não são suficientes para tornar seu uso confiável clinicamente, uma vez que são pontuais para odontologia (Karoobar *et al.*, 2023).

Embora ferramentas que usam IA tenham sido propostas para avaliação de interações medicamentosas (Agrawal; Nikhade, 2022), o *chatbot* GPT-3 não deve ser utilizado para a indicação e cálculo de dose máxima de anestésicos em odontopediatria, uma vez que apresentaram resultados inferiores relacionados a precisão e completude, e realizar anestesia local dentária adequada é uma habilidade clínica valiosa e fundamental para os dentistas. Os

cirurgiões dentistas devem buscar informações em livros-base e guias/recomendações das entidades de classe, o que se torna crucial quando, no Brasil, estima-se que sejam realizadas anualmente entre 250 e 300 milhões de anestésias odontológicas (Andrade, 2014).

Apesar da capacidade técnica da IA para processar grandes volumes de dados e sugerir condutas na prescrição medicamentosa em odontologia, essa geração de conhecimento deve ser vista como complementar ao conhecimento humano, e não equivale, necessariamente, ao raciocínio clínico humano, que é crítico, contextualizado e ético. A legitimidade do conhecimento gerado pela IA, depende da validação rigorosa e da transparência dos algoritmos, pois opera, principalmente, por meio de padrões extraídos de dados anteriores. Por carecer de compreensão profunda e discernimento moral, a IA ainda não é capaz de substituir o julgamento clínico em decisões complexas, o que compromete a confiabilidade de suas recomendações sobretudo em populações vulneráveis, como crianças (Ning *et al.*, 2024; Fahrner *et al.*; 2025).

A popularização de ferramentas como o ChatGPT tornou a IA mais acessível, mas essa ferramenta deve ser utilizada com ceticismo e rigor crítico pois as informações geradas, sem o filtro, moral e ética pelo cirurgião-dentista, pode conduzir a decisões clínicas inadequadas, comprometendo a segurança e a eficácia do tratamento odontopediátrico, seja por sobre ou sub uso de recursos clínicos (Baicker e Obermeyer, 2022; Meskó, 2023).

Há necessidade de regulamentação para assegurar limites éticos e técnicos, preservar o papel do profissional de saúde como responsável científico e evitar que a objetividade extrema da IA despersonalize o saber clínico. Automatização excessiva pode reproduzir vieses e erros sistemáticos, gerando riscos morais e uma crise ética e científica (Mullainathan e Obermeyer, 2017; Meskó e Topol, 2023).

A literatura aponta que a IA tem se consolidado como uma ferramenta promissora na odontologia auxiliando na gestão do paciente, diagnóstico, prevenção de falhas clínicas e apoio à tomada de decisões (Ding *et al.*, 2023). No entanto, seu uso em prescrição medicamentosa é pouco explorado, concentrando-se mais em radiologia (Ahmed *et al.*, 2021). A versão pública atual do ChatGPT-3 não deve ser utilizada como única fonte para terapêutica medicamentosa em odontopediatria, pois não foi projetado para fornecer orientação médica e, é inadequado para tomada de decisão clínica, devendo as respostas serem revisadas

rigorosa e cientificamente porque o *chat* gera informações falsas, ainda que convincentes, a fim de evitar complicações éticas e odonto-legais (Kim, 2023; Eggmann *et al.*, 2023; Johnson *et al.*, 2023). Deve-se lembrar que o dentista possui autoridade no ato da prescrição e erros relacionados a essa ação são de responsabilidade ética e legal do cirurgião dentista.

Espera-se que, o uso da IA na odontologia, se expanda com o avanço tecnológico, desenvolvendo modelos mais sofisticados (Thurzo *et al.*, 2022). Considerando que o ChatGPT aprende a partir de dados fornecidos através de treinamento de duas etapas e se aprimora através da modelagem baseada em IA, é recomendável que outras pesquisas em farmacologia odontopediátrica sejam realizadas após incorporação de literatura médica, bases de dados farmacológicas e evidências do mundo real para validação do desempenho clínico e das informações fornecidas. Isso é especialmente importante, porque a IA na odontologia, tem avançado lentamente e de forma restrita quando comparada com outras áreas, devido às preocupações com privacidade por parte dos pacientes e os dados usados em treinamento (Johnson *et al.*, 2023; Huang *et al.*, 2023).

## CONCLUSÃO

Respostas ‘completamente incorretas’ e ‘incompletas’ foram as mais prevalentes entre as 45 respostas sobre terapêutica medicamentosa em odontopediatria. As respostas sobre anestésicos locais apresentaram significativamente menores escores de precisão e completude, bem como pontuação total. A objetividade das respostas foi semelhante entre os grupos de medicamentos, sendo consideradas ‘muita quantidade de informações extras’. Na rotina clínica, esta IA deve ser utilizado com cautela e não deve ser a única fonte de consulta quando se trata da terapêutica medicamentosa em odontopediatria.

## AGRADECIMENTOS

O presente trabalho foi realizado com apoio da Fundação de Amparo à Pesquisa do Estado do Rio de Janeiro – FAPERJ (204.517/2024) e a Coordenação de Aperfeiçoamento de Pessoal de Nível Superior – CAPES (88887.894282/2023-00).

## CONTRIBUIÇÃO DOS AUTORES

**Michelle Alonso Coutinho:** Colaborou na redação técnica, na revisão especializada do manuscrito e realizou adequação do manuscrito às normas editoriais da revista. Todos os autores revisaram e aprovaram a versão final submetida, assumindo responsabilidade integral pelo conteúdo e pelos resultados apresentados.

**Hislena Oliveira Ramos:** Responsável pela redação inicial do manuscrito, participou do planejamento do estudo e contribuiu para a redação científica. Todos os autores revisaram e aprovaram a versão final submetida, assumindo responsabilidade integral pelo conteúdo e pelos resultados apresentados.

**Andréa Vaz Braga Pintor:** Participou da análise de dados, contribuiu para a redação científica e realizou revisão técnica do manuscrito. Todos os autores revisaram e aprovaram a versão final submetida, assumindo responsabilidade integral pelo conteúdo e pelos resultados apresentados.

**Marcela Baraúna Magno:** Responsável pelo delineamento metodológico do estudo, condução da coleta e análise dos dados, bem como pela redação inicial e revisão técnica do manuscrito. Todos os autores revisaram e aprovaram a versão final submetida, assumindo responsabilidade integral pelo conteúdo e pelos resultados apresentados.

## REFERÊNCIAS

AGRAWAL, P.; NIKHADE, P. Artificial intelligence in dentistry: past, present, and future. *Cureus*, v. 14, n. 7, e27405, 2022. DOI: <https://doi.org/10.7759/cureus.27405>.

AHMED, Naseer *et al.* Artificial Intelligence Techniques: Analysis, Application, and Outcome in Dentistry—A Systematic Review. *BioMed Research International*, v. 2021, p. 9751564, 2021. Disponível em: <https://doi.org/10.1155/2021/9751564> Acesso em: 14 nov. 2024.

ANDRADE, Eduardo D. de. *Terapêutica medicamentosa em odontologia*. 3 ed. São Paulo, Brasil: Editora Artes Médicas, 2014. 256 p.

BAICKER, Katherine; OBERMEYER, Ziad. Overuse and underuse of health care: new insights from economics and machine learning. *JAMA Health Forum, American Medical Association*, 2022. p. e220428-e220428. doi:10.1001/jamahealthforum.2022.0428

ÇAKMAKOĞLU, Ezgi Eroğlu. The place of ChatGPT in the future of dental education. *Journal of Clinical Trials and Experimental Investigations*, v. 2, n. 3, p. 121-129, 2023. DOI: <https://doi.org/10.5281/zenodo.8210063>

CALCULADORA DA DOSE DE MEDICAMENTOS EM PEDIATRIA. Bibliomed, Disponível em: <https://www.bibliomed.com.br/calculadoras/peddose/index.cfm> . Acesso em: 23 set. 2023.

CALCULADORA DE MEDICAMENTOS EM ODONTOPEDIATRIA. Departamento de Medicina Social – UFPEL. Disponível em: <https://dms.ufpel.edu.br/casca/modulos/odonto-calc#comp/odonto-calc> Acesso em: 23 set. 2023.

CÁLCULO DE DOSES DE MEDICAMENTOS. Biblioteca Virtual em Saúde. Disponível em: [aps.bvs.br/apps/calculadoras/?page=2](https://aps.bvs.br/apps/calculadoras/?page=2) . Acesso em: 23 set. 2023.

CAPAN, Belen S.; DUMAN, Canan; KALAOGLU, Elif E. Antibiotic prescribing practices for prophylaxis and therapy of oral/dental infections in pediatric patients—results of a cross-sectional study in Turkey. *GMS Hygiene and Infection Control*, v. 18, 2023. DOI: <https://doi.org/10.3205/dgkh000437>

CHATGPT. In: WIKIPÉDIA, a enciclopédia livre. Flórida: Wikimedia Foundation, 2024. Disponível em: <https://pt.wikipedia.org/w/index.php?title=ChatGPT&oldid=69499531> . Acesso em: 23 set. 2023.

DING, Hao *et al.* Artificial intelligence in dentistry—A review. *Frontiers in Dental Medicine*, v. 4, p. 1085251, 2023. DOI: <https://doi.org/10.3389/fdmed.2023.1085251>

DUBAY, William H. The principles of readability. Online Submission, 2004. 76 p. Disponível em: <http://files.eric.ed.gov/fulltext/ED490073.pdf> . Acesso em: 30 de nov. 2024.

EGGMANN, Florin *et al.* Implications of large language models such as ChatGPT for dental medicine. *Journal of Esthetic and Restorative Dentistry*, v. 35, n. 7, p. 1098-1102, 2023. DOI: <https://doi.org/10.1111/jerd.13046>

FAHRNER, L. John *et al.* The generative era of medical AI. *Cell*, v. 188, n. 14, p. 3648-3660, 2025. <https://doi.org/10.1016/j.cell.2025.05.018>

FERREIRA, Tânia R.; LOPES, Luciane C. Analysis of analgesic, antipyretic, and nonsteroidal anti-inflammatory drug use in pediatric prescriptions. *Jornal de pediatria*, v. 92, p. 81-87, 2016. DOI: <https://doi.org/10.1016/j.jped.2015.04.007>

FLESCH KINCAID CALCULATOR. Good Calculators – Free online calculators. Disponível em: <https://goodcalculators.com/flesch-kincaid-calculator/> . Acesso em: 25 out. 2023.

GILSON, Aidan *et al.* How does ChatGPT perform on the United States medical licensing examination? The implications of large language models for medical education and knowledge assessment. *JMIR Medical*

Education, v. 9, n. 1, p. e45312, 2023. DOI: 10.2196/45312. Disponível em: <https://mededu.jmir.org/2023/1/e45312>

HUANG, Hanyao *et al.* ChatGPT for shaping the future of dentistry: the potential of multi-modal large language model. International Journal of Oral Science, v. 15, n. 1, p. 29, 2023. DOI: <https://doi.org/10.1038/s41368-023-00239-y>

JOHNSON, Douglas *et al.* Assessing the accuracy and reliability of AI-generated medical: an evaluation of the Chat-GPT model. Research Square, 2023. Preprint, rs.3.rs-2566942. DOI: <https://doi.org/10.21203/rs.3.rs-2566942/v1>

KAROBARI, Mohmed I. *et al.* Evaluation of the diagnostic and prognostic accuracy of artificial intelligence in endodontic dentistry: A comprehensive review of literature. Computational and Mathematical Methods in Medicine, v. 2023, p. 7049360, 2023. DOI: <https://doi.org/10.1155/2023/7049360>

KIM, Seong-Gon. Using ChatGPT for language editing in scientific articles. Maxillofacial Plastic and Reconstructive Surgery, v. 45, n. 1, p. 13, 2023. DOI: <https://doi.org/10.1186/s40902-023-00381-x>

KINCAID, J. Peter *et al.* Derivation of new readability formulas (automated readability index, fog count, and Flesch reading ease formula) for Navy enlisted personnel. Research Branch Report, p. 8–75, 1975.

MACHADO, Hellen C. de P.; JASSÉ, Fernanda F. de A.; ARANTES, Diandra C. Uso de aplicativo de inteligência artificial como fonte de informação para resolução de provas de odontologia. Revista Brasileira de Odontologia Legal, v. 11, n. 1, p. 7-18, 2024. DOI: <https://doi.org/10.21117/rbol-v11n12024-535>

MAGO, Jyoti; SHARMA, Manoj. The Potential Usefulness of ChatGPT in Oral and Maxillofacial Radiology. Cureus, v. 15, n. 7, 2023. DOI: <https://doi.org/10.7759/cureus.42133>

MCCARTHY, John *et al.* A proposal for the dartmouth summer research project on artificial intelligence, august 31, 1955. AI magazine, v. 27, n. 4, p. 12-12, 2006. DOI: <https://doi.org/10.1609/aimag.v27i4.1904>

MESKÓ, Bertalan. The ChatGPT (generative artificial intelligence) revolution has made artificial intelligence approachable for medical professionals. Journal of Medical Internet Research, v. 25, p. e48392, 2023. DOI: 10.2196/52865. <https://doi.org/10.2196/52865>

MESKÓ, Bertalan; TOPOL, Eric J. The imperative for regulatory oversight of large language models (or generative AI) in healthcare. NPJ Digital Medicine, v. 6, n. 1, p. 120, 2023. Doi: <https://doi.org/10.1038/s41746-023-00873-0>

MIRA, Felipe. A. *et al.* Chat GPT for the management of obstructive sleep apnea: do we have a polar star? *European Archives of Oto-Rhino-Laryngology*, v. 281, p. 2087-2093, 2024. DOI: <https://doi.org/10.1007/s00405-023-08270-9>

MULLAINATHAN, Sendhil; OBERMEYER, Ziad. Does machine learning automate moral hazard and error?. *American Economic Review*, v. 107, n. 5, p. 476-480, 2017. doi:10.1257/aer.p20171084

NING, Yilin *et al.* Generative artificial intelligence and ethical considerations in health care: a scoping review and ethics checklist. *The Lancet Digital Health*, v. 6, n. 11, p. e848-e856, 2024. [https://doi.org/10.1016/s2589-7500\(24\)00143-2](https://doi.org/10.1016/s2589-7500(24)00143-2)

RUBANENKO, Moran *et al.* Assessment of the Knowledge and Approach of General Dentists Who Treat Children and Pediatric Dentists Regarding the Proper Use of Antibiotics for Children. *Antibiotics*, v. 10, n. 10, p. 1181, 2021. DOI: <https://doi.org/10.3390/antibiotics10101181>

SCARPARO, Angela. *Odontopediatria: Bases teóricas para uma prática clínica de excelência*. 1ª edição. Editora Manole. 1. ed. Barueri, SP: Manole, 2021. 544 p.

THURZO, Andrej. *et al.* Where is the Artificial Intelligence Applied in Dentistry? Systematic Review and Literature Analysis. *Healthcare*, v. 10, n. 7, p. 1269, 2022. Disponível em: <https://doi.org/10.3390/healthcare10071269> . Acesso em: 14 nov. 2024.

TURING, Alan. M. "Computing Machinery and Intelligence." *Mind*, v. 59, n. 236, p. 433-460, 1950. Disponível em: <http://www.jstor.org/stable/2251299> . Acesso em: 18 jan. 2025).



Este trabalho está licenciado com uma Licença [Creative Commons - Atribuição 4.0 Internacional](https://creativecommons.org/licenses/by/4.0/).