

INTELIGÊNCIA ARTIFICIAL

na Prática Clínica

Um guia ético e prático para psicólogos

FABIO HERNANDEZ DE MEDEIROS

Davi de Barros Silva · Letícia Maria Cristino Fritsch
Janaína Souza de Lima · Suelen Vitoria Gomes de Oliveira

1ª EDIÇÃO · BRASÍLIA · 2026

Inteligência Artificial na Prática Clínica

Um guia ético e prático para psicólogos

AUTOR

Fabio Hernandez de Medeiros

COAUTORES

Davi de Barros Silva · Letícia Maria Cristino Fritsch
Janaína Souza de Lima · Suelen Vitoria Gomes de Oliveira

EDIÇÃO

1ª edição | ISBN 978-85-65721-60-8

EDITORA DO INSTITUTO WALDEN4

BRASÍLIA · 2026

Dados Internacionais de Catalogação na Publicação (CIP)
(Câmara Brasileira do Livro, SP, Brasil)

Inteligência artificial na prática clínica
[livro eletrônico] : um guia ético e prático para
psicólogos / Fabio Hernandez de Medeiros ...
[et al.]. -- Brasília, DF : Instituto Walden4,
2026.
PDF

Outros autores: Davi de Barros Silva, Letícia
Maria Cristino Fritsch, Janaína Souza de Lima, Suelen
Vitoria Gomes de Oliveira

Bibliografia
ISBN 978-85-65721-60-8

1. Inovações tecnológicas 2. Inteligência
artificial 3. Medicina e saúde 4. Psicologia
5. Psicoterapia I. Medeiros, Fabio Hernandez de.
II. Silva, Davi de Barros. III. Fritsch, Letícia
Maria Cristino. IV. Lima, Janaína Souza de. V.
Oliveira, Suelen Vitoria Gomes de.

26-378188.0

CDD-610.285

Índices para catálogo sistemático:

1. Inteligência artificial : Ciências médicas
610.285

Maria Alice Ferreira - Bibliotecária - CRB-8/7964



CRÉDITOS

Inteligência Artificial na Prática Clínica — Um guia ético e prático para psicólogos

1ª edição · Brasília: Editora do Instituto Walden4, 2026.

© 2026 Fabio Hernandez de Medeiros e coautores. Todos os direitos reservados. Reprodução somente mediante autorização prévia dos autores.

Autor: Fabio Hernandez de Medeiros.

Coautores: Davi de Barros Silva · Letícia Maria Cristino Fritsch · Janaína Souza de Lima · Suelen Vitoria Gomes de Oliveira.

Editora: Editora do Instituto Walden4 — www.walden4.com.br

COMO CITAR ESTA OBRA

APA · 7ª ed.

Medeiros, F. H., Silva, D. B., Fritsch, L. M. C., Lima, J. S., & Oliveira, S. V. G. (2026). *Inteligência artificial na prática clínica: Um guia ético e prático para psicólogos*. Editora do Instituto Walden4.

ABNT · NBR 6023

MEDEIROS, F. H.; SILVA, D. B.; FRITSCH, L. M. C.; LIMA, J. S.; OLIVEIRA, S. V. G. **Inteligência artificial na prática clínica: um guia ético e prático para psicólogos**. Brasília: Editora do Instituto Walden4, 2026.



Produção com auxílio de IA

Produzida com auxílio de IA (Claude Opus 4.8, Anthropic), sob revisão humana integral do autor.



Uso educacional

Conteúdo formativo. Não substitui a prática clínica, a supervisão profissional nem a assessoria jurídica/ética.

Editora do Instituto Walden4

A Editora do Instituto Walden4 tem o compromisso de disseminar o conhecimento produzido na área da Análise do Comportamento, tanto em termos científicos quanto profissionais. Buscando democratizar o acesso a esse conhecimento, muitos de nossos livros são disponibilizados gratuitamente. Todos os nossos livros estão disponíveis no formato digital online, o que significa que em questão de segundos você poderá estar desfrutando da leitura dos livros publicados por nós que despertarem seu interesse.

CONSELHO EDITORIAL

Dr. Gleidson Gabriel da Cruz

Dr. Márcio Borges Moreira

CONTATO

secretaria@walden4.com.br

www.walden4.com.br

[@instituto.walden4](https://www.instagram.com/instituto.walden4)

[facebook.com/iwalden4](https://www.facebook.com/iwalden4)

[youtube.com/user/instwalden4](https://www.youtube.com/user/instwalden4)

Sumário

1	O que é IA e o que não é	7
2	Ética e limites profissionais	12
3	O que o CFP e a legislação dizem	17
4	Uso seguro da IA	27
5	Prompts para psicólogos	33
6	Ferramentas de IA	39
7	Futuro da Psicologia com IA	44
§	Referências & encerramento	47

CAPÍTULO 1

O que é IA e o que não é



Objetivos de aprendizagem

Ao final deste capítulo, você será capaz de:

- Distinguir os conceitos de **IA**, **Machine Learning**, **Deep Learning** e **LLM**.
- Explicar, em linguagem simples, **como um LLM gera texto** — e por que isso não é “compreender”.
- Reconhecer os principais **mitos e limites** (alucinação, ausência de consciência, viés).
- Identificar **aplicações reais** de IA já documentadas em saúde mental.

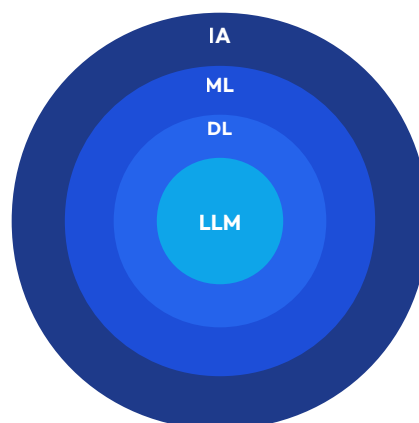
1.0 • Por que um psicólogo precisa entender isso

Você não precisa programar para usar IA com responsabilidade — mas precisa **entender o que ela é**. O Código de Ética exige competência técnica para qualquer ferramenta que entre na prática (Art. 1º, b; ver Cap. 3). Entender a IA é o que permite confiar na medida certa: nem rejeitar por medo, nem idolatrar por encantamento.

A boa notícia: os conceitos essenciais cabem em poucas páginas.

1.1 • As camadas: IA › ML › DL › LLM

Pense em **círculos concêntricos**, do maior para o menor — cada camada está contida na anterior:



- **Inteligência Artificial (IA)** — o campo mais amplo: qualquer sistema que executa tarefas que associamos à inteligência humana (decidir, reconhecer, recomendar). Inclui desde regras simples até os modelos atuais.

- **Machine Learning (ML)** — um ramo da IA em que o sistema **aprende padrões a partir de dados**, em vez de seguir regras escritas à mão. É a base das aplicações de IA em saúde mental revisadas na literatura (Singhal et al., 2024).
- **Deep Learning (DL)** — um tipo de ML que usa **redes neurais artificiais** com muitas camadas, capazes de capturar padrões complexos (imagens, linguagem, sinais). Redes neurais já são aplicadas a problemas psiquiátricos (Komatsu et al., 2021).
- **Large Language Model (LLM)** — um modelo de Deep Learning **especializado em linguagem**, treinado em enormes volumes de texto. ChatGPT, Claude e Gemini são LLMs.

**Fixe isto**

Todo LLM é Deep Learning; todo Deep Learning é Machine Learning; todo Machine Learning é IA. O contrário **não** vale.

1.2 • Como um LLM “*pensa*” (spoiler: ele não pensa)

Esta é a ideia mais importante do capítulo.

Um LLM funciona, em essência, como um **completador de texto extraordinariamente sofisticado**. Dado um trecho, ele calcula **qual a próxima palavra mais provável**, com base nos padrões estatísticos que extraiu de bilhões de frases durante o treino. Depois repete isso, palavra após palavra.

**Analogia**

É como o corretor preditivo do celular — só que com uma capacidade incomparavelmente maior de levar em conta o contexto. Quando o texto resultante soa inteligente, é porque os **padrões da linguagem humana** estão embutidos nas probabilidades, não porque há compreensão por trás.

O que isso implica para a clínica:

- O LLM **não entende** o sofrimento de quem escreveu o relato. Ele reconhece padrões textuais de sofrimento.
- Ele **não tem intenção, memória de você nem cuidado**. A “empatia” que aparece é forma linguística, não vínculo (ver Cap. 2).
- Ele pode produzir uma resposta **fluente e errada** com a mesma confiança de uma resposta correta.

Revisões recentes em saúde mental insistem nesse ponto: LLMs são promissores como apoio, mas seu funcionamento estatístico impõe cautela sobre o que se pode delegar a eles (Guo et al., 2024; Mesquiti & Nook, 2026).

1.3 • Mitos e limites

Mito	Fato
“A IA entende o paciente.”	Ela reconhece padrões de linguagem; não há compreensão nem consciência.
“Se respondeu com segurança, está certo.”	LLMs alucinam : geram informação falsa com tom convincente, inclusive referências e dados inexistentes.
“A IA é neutra e objetiva.”	Ela reflete vieses dos dados de treino — culturais, de gênero, de raça (Holm, 2024; ver Cap. 2).
“A IA aprende sozinha com a minha conversa em tempo real.”	O modelo não “aprende” durante o uso; e suas conversas só viram treino se a política permitir (Cap. 4).
“Quanto mais detalhe do caso eu der, melhor.”	Mais detalhe identificável = mais risco de quebra de sigilo, sem ganho clínico (Cap. 4).

Os três limites que mais importam ao psicólogo

1. **Alucinação** — o modelo pode inventar fatos, estudos e citações. Tudo precisa de validação humana.
2. **Ausência de consciência e responsabilidade** — a IA não responde eticamente; **você** responde (Código de Ética, Art. 2º, g).
3. **Viés de treino** — saídas podem reproduzir preconceitos; exigem leitura crítica.

1.4 • O que a IA já faz em saúde mental (aplicações reais)

Para não cair nem no ceticismo nem no exagero, vale conhecer o que **já está documentado** na literatura revisada por pares — sempre como **apoio**, não substituição:

- **Triagem e organização de informação** e apoio a pesquisa em saúde mental (Irshad et al., 2022; Mesquiti & Nook, 2026).
- **Apoio a tarefas clínicas acessórias** — psicoeducação, organização de materiais, rascunho de documentos (detalhado no Cap. 5).
- **Ferramentas conversacionais** que demonstram potencial e, ao mesmo tempo, limites importantes de segurança e eficácia (Guo et al., 2024; Singhal et al., 2024).
- **Pesquisa e análise** — apoio à análise de dados e à produção científica em psicologia (Mesquiti & Nook, 2026).



Card de equilíbrio

A evidência atual sustenta a IA como **assistente de tarefas**, não como agente clínico. Onde os estudos apontam bons resultados, há sempre um profissional humano no comando da decisão.

EM UMA FRASE

IA é um campo amplo cujo expoente atual, o LLM, é um **preditor de texto** muito capaz — útil como apoio, mas sem compreensão, consciência ou responsabilidade: por isso o julgamento clínico continua, inteiramente, com o psicólogo.

Referências deste capítulo

- Singhal et al. (2024), *Current Psychiatry Reports* — ML em saúde mental e o papel do clínico.
- Komatsu et al. (2021), *Biomedicines* — redes neurais aplicadas à psiquiatria.
- Irshad et al. (2022), *Journal of Health and Medical Sciences* — usos de IA na psicologia.
- Guo et al. (2024), *JMIR Mental Health* — revisão sistemática de LLMs em saúde mental.
- Mesquiti & Nook (2026), *Annual Review of Biomedical Data Science* — LLMs em pesquisa e tratamento.
- Holm (2024), *Frontiers in Psychiatry* — vieses e limites de objetividade.

Norma introduzida: dever de competência técnica (Código de Ética, Art. 1º, b) — aprofundado no Cap. 3.



INSTITUTO WALDEN4 · AVALIAÇÃO

Avaliação do Comportamento

Avaliação funcional detalhada, baseada na Análise do Comportamento Aplicada (ABA), para identificar necessidades e desenhar objetivos terapêuticos individualizados.

- Avaliação funcional do comportamento
- Definição de objetivos terapêuticos personalizados
- Plano de intervenção orientado por evidências

Quanto antes a avaliação, melhores os resultados.

Agende uma avaliação →

walden4.com.br/autismo/avaliacao · [@instituto.walden4](https://www.instagram.com/instituto.walden4)

INSTITUTO WALDEN4 · CIÊNCIA QUE CUIDA

CAPÍTULO 2

Ética e limites profissionais



Objetivos de aprendizagem

Ao final deste capítulo, você será capaz de:

- Reconhecer os **riscos éticos centrais** do uso de IA na prática psicológica.
- Aplicar **sigilo, transparência e responsabilidade** ao trabalho com IA.
- Distinguir **empatia simulada** de **cuidado real**.
- Identificar os **usos proibidos** da IA na clínica.

2.0 · Da norma à conduta

O Capítulo 3 apresenta as normas; este capítulo trata da **postura ética** ao usar IA — o que fazer quando a regra não basta e é preciso julgamento. A literatura especializada já mapeou esse terreno de forma sistemática: há catálogos de considerações éticas (Saeidnia et al., 2024), de violações concretas (Iftikhar et al., 2025) e propostas de frameworks de princípios (Pillay, 2025). Organizamos tudo em **cinco pilares**.

2.1 · Pilar 1 — Sigilo e confidencialidade

É o pilar mais concreto e o mais fácil de violar sem perceber. Inserir dados identificáveis de um paciente em uma IA em nuvem pode significar **expor informação sigilosa a terceiros** (Código de Ética, Art. 9º; Res. CFP 9/2024, Art. 3º; LGPD).

- **O risco:** a confidencialidade não se rompe só por má-fé — rompe-se por descuido, por colar um relato “só para resumir”.
- **A conduta:** anonimizar **sempre** antes de usar a IA (o Cap. 4 ensina como). O sigilo é dever indelegável; nenhuma ferramenta o assume por você.

2.2 · Pilar 2 — Viés algorítmico e justiça

LLMs aprendem com dados humanos e, com eles, herdam **preconceitos** — de gênero, raça, classe, cultura. Uma saída pode parecer técnica e, ainda assim, carregar um viés que prejudica determinados grupos (Holm, 2024; Zhang et al., 2025).

- **O risco:** aceitar acriticamente uma sugestão enviesada como se fosse objetiva.
- **A conduta:** ler toda saída com olhar crítico, atento a generalizações e estereótipos. A IA **não é neutra**; cabe ao psicólogo filtrar.

2.3 • Pilar 3 — Responsabilidade clínica indelegável

Este é o pilar que sustenta todos os outros. A **responsabilidade por qualquer decisão ou documento é, sempre, do psicólogo humano** — nunca da ferramenta (Código de Ética, Art. 1º, b e Art. 2º, g).

- **O risco:** o “viés de automação” — confiar na máquina por ser rápida e fluente, terceirizando o julgamento.
- **A conduta:** a IA **rascunha**; o psicólogo **valida, corrige e assina**. Se a saída estiver errada, a falha é profissional, não “da IA”. A discussão ética sobre apoio à decisão clínica reforça que a ferramenta informa, mas não decide (Heyen & Salloch, 2021).



Regra de bolso

Se você não consegue defender tecnicamente uma frase gerada pela IA, ela não pode entrar no seu documento.

2.4 • Pilar 4 — Transparência com o paciente

Quando a IA participa de algo que envolve o paciente (por exemplo, apoio à elaboração de materiais ou registros), há um dever de **informar** (Código de Ética, Art. 14; Res. CFP 9/2024, Art. 7º, parágrafo único).

- **O risco:** uso silencioso de IA que, se revelado, quebraria a confiança do vínculo.
- **A conduta:** ser transparente sobre o uso de tecnologia de apoio, na medida adequada a cada contexto. Transparência protege o vínculo e cumpre a norma.

2.5 • Pilar 5 — Empatia simulada × cuidado real

Um LLM pode produzir frases acolhedoras, mas isso é **forma linguística, não cuidado** (ver Cap. 1). A relação terapêutica — vínculo, presença, responsabilidade — é exatamente o que a IA **não** tem.

- **O risco:** confundir fluência empática com competência clínica; sugerir, implícita ou explicitamente, que a IA “atende”.
- **A conduta:** preservar o vínculo como núcleo insubstituível da clínica. A discussão atual contrasta a *eficiência* da IA com a *empatia* humana e conclui que uma não substitui a outra (Bharti & Singh, 2025; Gupta & Gandhi, 2025).

2.6 • Os usos proibidos (linha vermelha)

- ⊗ **É eticamente vedado usar IA para:**
- ✗ **Fazer diagnóstico psicológico** (Código de Ética, Art. 2º, q)
- ✗ **Substituir a avaliação clínica ou a psicoterapia**
- ✗ **Emitir parecer ou laudo automático** (Art. 2º, g — documento exige fundamentação técnico-científica)
- ✗ **Manejar crise** (ideação suicida, violência, risco) sem supervisão e rede presencial (Res. CFP 9/2024, Art. 5º)
- ✗ **Inserir dados identificáveis de pacientes** (Art. 9º; LGPD)

Estudos que catalogam o comportamento de “conselheiros de IA” mostram, empiricamente, por que essas linhas existem: sem essas barreiras, surgem violações éticas concretas e potencialmente danosas (Iftikhar et al., 2025; Saeidnia et al., 2024).

2.7 • Síntese: os cinco pilares

Pilar	Risco central	Conduta
Sigilo	Exposição de dado a terceiros	Anonimizar sempre (Cap. 4)
Viés	Aceitar saída enviesada como neutra	Leitura crítica
Responsabilidade	Terceirizar o julgamento à máquina	Validar e assinar; a IA só rascunha
Transparência	Uso silencioso de IA	Informar o paciente
Cuidado real	Confundir empatia simulada com vínculo	Preservar a relação humana

EM UMA FRASE

A ética do uso de IA na Psicologia se resolve em uma frase: **a ferramenta pode apoiar tarefas, mas o sigilo, o julgamento, a responsabilidade e o vínculo permanecem, integralmente, humanos.**

Referências deste capítulo

- Saeidnia et al. (2024), *Social Sciences* — considerações éticas centrais.
- Holm (2024), *Frontiers in Psychiatry* — vieses e limites de objetividade.
- Heyen & Salloch (2021), *BMC Medical Ethics* — ética do apoio à decisão clínica.
- Bharti & Singh (2025), *Science Insights* — eficiência × empatia na clínica.
- Zhang et al. (2025), *Nature Computational Science* — papéis de LLM e questões éticas.
- Iftikhar et al. (2025), *Proc. AAAI/ACM AIES* — violações éticas de “LLM counselors”.
- Gupta & Gandhi (2025), *International Review of Psychiatry* — letramento em IA e futuro.
- Pillay (2025), *Healthcare* — framework ético de cinco pilares.

Normativas aplicadas: Código de Ética (Art. 1º b, 2º g, 2º q, 9º, 14); Res. CFP 9/2024 (Art. 5º, 7º).
Detalhamento no Cap. 3.



INSTITUTO WALDEN4 · PSICOTERAPIA

Psicoterapia

Um caminho de autoconhecimento e transformação pessoal, em um espaço acolhedor e seguro — para adultos, crianças e adolescentes.

- Atendimento presencial e online
- Particular ou plano de saúde
- +15 anos de experiência · +5.000 pacientes atendidos

Dê o primeiro passo no seu autoconhecimento.

Marque sua sessão →

walden4.com.br/psicoterapia · [@instituto.walden4](https://www.instagram.com/instituto.walden4)

INSTITUTO WALDEN4 · CIÊNCIA QUE CUIDA

CAPÍTULO 3

O que o CFP e a legislação dizem



Objetivos de aprendizagem

Ao final deste capítulo, você será capaz de:

- Localizar o uso de IA dentro do arcabouço normativo brasileiro que rege a profissão.
- Reconhecer quais deveres do **Código de Ética** se aplicam diretamente ao uso de ferramentas de IA.
- Entender por que a **Resolução CFP nº 9/2024** (TDICs) alcança as ferramentas de IA em nuvem.
- Identificar o que a **LGPD** exige ao tratar dados de saúde mental em sistemas de IA.
- Diferenciar o que é **norma vinculante** do que é **orientação institucional** do CFP sobre IA.

3.0 · Por que este capítulo vem antes dos prompts

É tentador pular direto para “como usar a IA”. Mas, na Psicologia, **a pergunta não é apenas se a ferramenta funciona — é se o seu uso respeita o sigilo, a responsabilidade técnica e os direitos do paciente**. Nenhuma dessas exigências foi criada para a IA; elas já existiam e simplesmente **passam a valer também quando há uma IA no fluxo de trabalho**.

Em outras palavras: não existe (ainda) uma “lei da IA na Psicologia”. O que existe é um conjunto de normas — algumas de 2005, outras de 2018, 2024 e 2025 — que, lidas em conjunto, já dão respostas claras sobre o que pode e o que não pode. Este capítulo organiza esse conjunto em quatro camadas.



Mapa das quatro camadas normativas

1. **Código de Ética do Psicólogo** (Res. CFP 010/2005) — os deveres de fundo: sigilo, competência, responsabilidade.
2. **Resolução CFP nº 9/2024** — regula o exercício mediado por tecnologia (TDICs); alcança servidores remotos e, por extensão, IA em nuvem.
3. **Documentos do CFP sobre IA (2025)** — Nota de Posicionamento + cartilhas: orientação institucional, não resolução.
4. **LGPD** (Lei 13.709/2018) — proteção de dados pessoais sensíveis de saúde.

3.1 • O Código de Ética do Psicólogo (Resolução CFP nº 010/2005)

O Código de Ética é a camada mais antiga e, ainda assim, a mais decisiva para o uso de IA. Ele não menciona “inteligência artificial” — não precisava. Seus deveres são redigidos de forma ampla o bastante para alcançar qualquer ferramenta que o psicólogo coloque entre si e o paciente.

Quatro deveres concentram quase tudo o que importa aqui.

a) Competência técnica (Art. 1º, b e c)

“Assumir responsabilidades profissionais somente por atividades para as quais esteja capacitado pessoal, teórica e tecnicamente” (Art. 1º, b) e “prestar serviços psicológicos de qualidade [...] utilizando princípios, conhecimentos e técnicas reconhecidamente fundamentados na ciência psicológica” (Art. 1º, c).

Na prática com IA: usar uma ferramenta que você não compreende é, em si, um problema ético. Antes de colocar a IA no fluxo de trabalho, é preciso entender o que ela faz, como erra (alucinação, viés) e onde guarda os dados. O dever de competência fundamenta a própria existência deste livro: **letramento em IA é pré-requisito, não acessório.**

b) Responsabilidade técnico-científica dos documentos (Art. 2º, g)

Ao psicólogo é vedado “emitir documentos sem fundamentação e qualidade técnico-científica”.

Na prática com IA: um texto gerado por IA **não tem** fundamentação técnico-científica própria — ele é predição estatística de palavras. Se você usa a IA para apoiar a redação de um relatório, **a responsabilidade técnica continua integralmente sua.** A IA rascunha; o psicólogo valida, corrige e assina. Nunca o contrário.

c) Sigilo profissional (Art. 9º e 10)

“É dever do psicólogo respeitar o sigilo profissional a fim de proteger, por meio da confidencialidade, a intimidade das pessoas [...] a que tenha acesso no exercício profissional” (Art. 9º).

Na prática com IA: este é o ponto mais sensível de todo o livro. Colar o relato de um paciente — com nome, história, detalhes identificáveis — em um chatbot de IA pode significar **expor dado sigiloso a um servidor de terceiros.** O sigilo do Art. 9º é o motivo pelo qual o Capítulo 4 inteiro é dedicado a anonimização. A quebra de sigilo só é admitida em hipóteses estritas (Art. 10) e jamais por conveniência tecnológica.

d) Diagnóstico e exposição (Art. 2º, q)

É vedado “realizar diagnósticos, divulgar procedimentos ou apresentar resultados de serviços psicológicos em meios de comunicação, de forma a expor pessoas, grupos ou organizações”.

Na prática com IA: sustenta a **trava de segurança** do livro — nenhum prompt deste material faz diagnóstico, parecer automático ou manejo de crise. A IA não diagnostica; ela apoia tarefas acessórias.

**Dever de informar o paciente (Art. 14)**

“A utilização de quaisquer meios de registro e observação da prática psicológica [...] devendo o usuário ou beneficiário, desde o início, ser informado.” Se a IA entra no registro ou na observação da prática, há um dever de **transparência** com o paciente. Ver Capítulo 2.

3.2 • Resolução CFP nº 9/2024 — exercício mediado por TDICs**Por que esta resolução, e não a 11/2018****Correção normativa importante**

A antiga Resolução CFP nº 11/2018 (telepsicologia) está **REVOGADA**. A norma vigente é a Resolução CFP nº 9, de 18 de julho de 2024, que revogou expressamente a 11/2018 e a 04/2020 (Art. 9º). Qualquer material que ainda cite a 11/2018 como vigente está desatualizado.

A Resolução 9/2024 regula “o exercício da psicologia mediado por Tecnologia Digital da Informação e da Comunicação (TDICs) em território nacional” (Art. 1º). Embora tenha sido pensada principalmente para o atendimento online, sua definição de TDIC é ampla o suficiente para alcançar ferramentas de IA.

O elo com a IA: “servidores remotos”

O Art. 2º define exercício mediado por TDICs como aquele que envolve “emprego de métodos e técnicas psicológicas dependentes de **servidores remotos**”. O Art. 3º, III repete: “emprego de métodos e técnicas psicológicas mediante servidores remotos”.

Toda ferramenta de IA em nuvem (ChatGPT, Claude, Gemini etc.) roda em servidores remotos. A leitura mais prudente, portanto, é considerar que, ao processar informação ligada à prática clínica nesses serviços, o psicólogo se coloca **sob o alcance** desta resolução — e dos deveres que ela impõe. Não se trata de interpretação já pacificada pelo CFP, mas da postura mais conservadora e segura diante da norma.

Os deveres que mais importam para a IA

Artigo	O que exige	O que isso significa ao usar IA
Art. 3º, II	Responsabilidade ética no manuseio de dados sensíveis e implicações para o sigilo, a privacidade e a autonomia	Na prática: antes de digitar qualquer coisa numa IA, identifique se há dado sensível (ver o quadro abaixo). Havendo, anonimize ou não use — o que você cola passa a estar nas mãos de um terceiro.
Art. 3º, IV	Zelo ético: ser responsável por dados e informações sensíveis	Na prática: a responsabilidade pelo dado não se transfere para a empresa de IA quando você usa a ferramenta — ela continua inteiramente sua.
Art. 4º, I	Avaliar as condições de confidencialidade e privacidade das informações	Na prática: antes de adotar uma ferramenta, descubra onde ela guarda seus dados, por quanto tempo e se os usa para treinar o modelo — e só use a que você conhece e aprova.
Art. 4º, II	Ter as competências para o manejo das TDICs empregadas	Na prática: você precisa entender como a ferramenta funciona, como ela erra (alucinação, viés) e onde processa os dados antes de colocá-la no seu fluxo de trabalho.
Art. 4º, V	Embasar o uso (ou o não-uso) da TDIC em produção científica e ética	Na prática: decida usar — ou deixar de usar — a IA com base em evidência e fundamentos éticos , não por modismo ou comodidade.
Art. 5º	Encaminhar à rede presencial situações de risco de morte, violência, urgência e emergência	Na prática: diante de risco, violência ou emergência, encaminhe à rede presencial de imediato . O manejo de crise não é tarefa de IA — exige presença humana e rede de cuidado.
Art. 7º, parág. único	Especificar os recursos tecnológicos que garantem o sigilo e informar o cliente	Na prática: avise o paciente quando uma tecnologia (inclusive IA) participa do atendimento e quais recursos protegem o sigilo dele.

**O que conta como “informação sensível”?**

A LGPD (Art. 5º, II) define **dado pessoal sensível** como aquele que revela origem racial ou étnica, convicção religiosa, opinião política, filiação sindical ou a organização religiosa/filosófica/política, **dado referente à saúde ou à vida sexual**, e dado genético ou biométrico. Na clínica, isso abrange quase tudo o que um paciente traz: diagnóstico, sofrimento, história de vida, orientação sexual, religião. Por isso, no contexto clínico, **as informações que um paciente traz em regra configuram dado pessoal sensível** — e exigem proteção reforçada antes de chegar a qualquer IA (detalhamento jurídico em 3.4).

**Leitura-chave**

A Resolução 9/2024 não “proíbe” nem “autoriza” a IA — ela **transfere para o uso de IA os mesmos deveres de sigilo, competência e responsabilidade** que já regem o atendimento mediado por tecnologia. A ferramenta muda; o dever permanece.

3.3 • O posicionamento institucional do CFP sobre IA (2025)

Em 2025, o Conselho Federal de Psicologia passou a se manifestar diretamente sobre IA. É importante entender a **natureza** desses documentos: são **orientações institucionais**, não resoluções com força normativa autônoma. Eles interpretam e aplicam as normas existentes ao contexto da IA — e sinalizam para onde a regulação caminha.

São três documentos confirmados em fonte oficial do CFP.

1. Nota de Posicionamento sobre IA e Psicologia (03/07/2025) — posição institucional sobre o uso de IA na profissão. O CFP afirma:

“[...] reforçamos que, embora a IA possa gerar conteúdos com agilidade e oferecer apoio a diversas tarefas, sua aplicação exige supervisão e discernimento humanos. Cabe exclusivamente às psicólogas e aos psicólogos, em seus distintos contextos de trabalho — clínicos, institucionais, educacionais, organizacionais, jurídicos, comunitários, entre outros —, decidir se, quando e como integrar essas ferramentas à sua prática tendo como referência que o compromisso da Psicologia é com a garantia dos direitos fundamentais, com a promoção da saúde mental, com o desenvolvimento humano, com a proteção social e cidadania digna e com o exercício responsável da profissão, em todos os seus campos de atuação.”

2. Cartilha “Inteligência Artificial na Psicologia: guia para uma prática ética e responsável” (CFP, 2025) — guia de boas práticas para profissionais. Abre lembrando que a IA não é neutra:

“A Inteligência Artificial não é uma força neutra da natureza, trata-se de uma criação humana, moldada por valores e preconceitos humanos e pelos dados da sociedade em que emerge. Se não formos vigilantes, críticos e atualizados, a IA corre o

risco de se tornar mais uma ferramenta a serviço da reprodução e da amplificação de desigualdades e injustiças que já afligem nossa sociedade. O Conselho Federal de Psicologia e cada psicóloga e psicólogo brasileiro têm um compromisso ético e social com a promoção da equidade, da justiça social e da não discriminação.”

A cartilha traz recomendações práticas, entre elas:

Frente	Recomendação	Página
Supervisão de assistentes e chatbots	É vedado o uso de assistentes virtuais, chatbots, avatares ou robôs para a condução inteiramente autônoma de psicoterapias ou intervenções clínicas; exige-se supervisão humana contínua e suficiente.	p. 13 e 24
Gravação e transcrição de atendimentos	Obter consentimento informado prévio do paciente e usar apenas plataformas cujas condições de privacidade, sigilo e armazenamento sejam conhecidas e seguras.	p. 24
Avaliação psicológica	Uso de testes, plataformas de correção ou IA para laudos/diagnósticos exige supervisão e mediação humana direta ; o psicólogo é o único responsável pelo resultado final e pela validação.	p. 24
Triagem e agendamentos	Uso recomendado para tarefas administrativas rotineiras que não tocam o núcleo técnico (agenda, lembretes, fluxo de trabalho).	p. 22
Sistematização de prontuários	Governança rigorosa da informação, assegurando que o armazenamento cumpra todas as regras de sigilo profissional.	p. 25

3. Cartilha “Chatbots, Inteligência Artificial e sua Saúde Mental” (CFP, 2025) — material voltado ao **público geral**, que alerta para os riscos de tratar um chatbot como terapeuta. O guia reconhece que chatbots podem oferecer informação e apoio inicial, mas adverte que **não realizam avaliação clínica, não assumem responsabilidade ética e podem errar — ou até agravar — um quadro de saúde mental**, e que inserir dados identificáveis de um paciente em uma IA genérica representa risco ético e legal.

A direção dessas manifestações é clara e alinhada ao restante do arcabouço: a IA é admitida como **apoio**, sob responsabilidade humana integral, com proteção rigorosa de dados e sem substituir o vínculo terapêutico. É o mesmo princípio condutor deste livro.

3.4 • LGPD — Lei nº 13.709/2018

A Lei Geral de Proteção de Dados é a camada que falta para fechar o quadro. Ela não fala de Psicologia, mas trata do insumo central da clínica: **dados pessoais**.

Dado de saúde mental é dado sensível

A LGPD cria uma categoria especial — **dado pessoal sensível**. O Art. 5º, II define como tal o:

“dado pessoal sobre origem racial ou étnica, convicção religiosa, opinião política, filiação a sindicato ou a organização de caráter religioso, filosófico ou político, dado referente à saúde ou à vida sexual, dado genético ou biométrico, quando vinculado a uma pessoa natural”.

Informação sobre o estado psíquico, diagnóstico, sofrimento ou tratamento de uma pessoa é, portanto, **dado referente à saúde** — dado sensível — e recebe proteção reforçada.

Isso muda tudo na prática com IA: o que você cola num chatbot sobre um paciente não é “um texto qualquer” — é, juridicamente, **dado sensível de saúde de um terceiro**.

Três exigências que se aplicam diretamente

Conceito da LGPD	Exigência	Aplicação ao uso de IA
Base legal e finalidade	Todo tratamento de dado precisa de base legal e finalidade legítima e específica	Você precisa de fundamento para inserir o dado do paciente numa IA — e a “comodidade” não é base legal.
Consentimento / transparência	O titular deve saber como seus dados são tratados	Conecta com o dever de informar do Código (Art. 14) e da Res. 9/2024 (Art. 7°).
Segurança e dados sensíveis	Dados sensíveis exigem medidas de segurança reforçadas	Anonimizar antes de usar a IA; conhecer onde a ferramenta armazena e se usa o dado para treino (Cap. 4 e 6).



A saída mais segura

A LGPD define **anonimização** (Art. 5°, XI) como a “utilização de meios técnicos razoáveis e disponíveis no momento do tratamento, por meio dos quais um dado perde a possibilidade de associação, direta ou indireta, a um indivíduo”, e o resultado é o **dado anonimizado** (Art. 5°, III): “dado relativo a titular que não possa ser identificado, considerando a utilização de meios técnicos razoáveis e disponíveis na ocasião de seu tratamento”. O **Art. 12** estabelece que “os dados anonimizados não serão considerados dados pessoais para os fins desta Lei, salvo quando o processo de anonimização ao qual foram submetidos for revertido [...] ou quando, com esforços razoáveis, puder ser revertido”. Ou seja: o dado verdadeiramente anonimizado sai do regime mais rígido da LGPD — mas só enquanto **não puder ser reidentificado**. Por isso o Capítulo 4 trata anonimização não como “boa prática opcional”, mas como **a condição** para usar IA com responsabilidade.

3.5 • Síntese: a norma → a prática

Fonte	O que exige	O que isso significa ao usar IA
Código de Ética, Art. 1° b/c	Competência técnica	Entenda a ferramenta antes de usá-la.
Código de Ética, Art. 2° g	Documentos com base técnico-científica	A IA rascunha; você valida e assina.
Código de Ética, Art. 9°/10	Sigilo profissional	Não exponha dado identificável a servidores de terceiros.

Fonte	O que exige	O que isso significa ao usar IA
Código de Ética, Art. 2º q	Vedação a diagnóstico exposto	IA não diagnostica; é apoio.
Código de Ética, Art. 14	Informar o usuário	Transparência sobre o uso de tecnologia.
Res. CFP 9/2024, Art. 3º–4º	Sigilo, privacidade e competência em servidores remotos	Os deveres do atendimento mediado valem para a IA em nuvem.
Res. CFP 9/2024, Art. 5º	Encaminhar risco à rede presencial	Crise não se maneja por IA.
CFP — documentos de IA (2025)	Orientação: IA como apoio, sob responsabilidade humana	Direção institucional alinhada ao princípio condutor.
LGPD (13.709/2018)	Proteção reforçada a dado sensível de saúde	Anonimizar e conhecer a política de dados.

EM UMA FRASE

O Brasil ainda não tem uma “lei da IA na Psicologia” — mas o **Código de Ética, a Resolução CFP 9/2024 e a LGPD**, lidos em conjunto, já respondem com clareza: a IA é admitida como **apoio a tarefas acessórias, sob sigilo rigoroso, responsabilidade técnica indelegável e transparência** com o paciente.

Referências deste capítulo**Normativas (fontes primárias):**

- Conselho Federal de Psicologia. (2005). *Código de Ética Profissional do Psicólogo* (Resolução CFP nº 010/2005).
- Conselho Federal de Psicologia. (2024). *Resolução CFP nº 9, de 18 de julho de 2024 — exercício profissional mediado por TDICs*.
- Conselho Federal de Psicologia. (2025, 3 de julho). *CFP divulga posicionamento sobre Inteligência Artificial no contexto da prática psicológica*. site.cfp.org.br
- Conselho Federal de Psicologia. (2025). *Inteligência Artificial na Psicologia: guia para uma prática ética e responsável* (cartilha). [PDF](#)
- Conselho Federal de Psicologia. (2025). *Chatbots, Inteligência Artificial e sua Saúde Mental: um guia para navegar com mais segurança na nova fronteira digital* (cartilha). [PDF](#)
- Brasil. (2018). *Lei nº 13.709, de 14 de agosto de 2018 — Lei Geral de Proteção de Dados Pessoais (LGPD)*.

Científicas (da matriz do projeto):

- Pillay (2025), *Healthcare* — framework ético de cinco pilares.
- Heyen & Salloch (2021), *BMC Medical Ethics* — ética do apoio à decisão clínica (humano-IA).

**Nota de validação**

As citações dos documentos oficiais (CFP 2025 e LGPD) foram **verificadas literalmente na fonte** (Planalto e site do CFP), conforme a regra anti-alucinação do projeto.

PUBLICIDADE



idio
PSICOLOGIA INTELIGENTE

GESTÃO INTELIGENTE PARA PSICÓLOGOS

A teoria você aprendeu neste livro. Para a sua prática clínica, conte com a Idio.

Prontuário com IA, agenda, videochamada segura e financeiro — tudo em um só lugar.



Prontuário com IA

Resumos automáticos com a PsiquiA



Agenda & WhatsApp

Lembretes automáticos, menos faltas



Atendimento online

Vídeo integrado, sem instalar nada



Gestão financeira

Recibos e cobranças via Pix



Portal do paciente

Agendamentos e documentos



Acesso mobile

Onde você estiver



Criptografia TLS/repouso · conformidade com a LGPD · backups automáticos diários

CONVITE ESPECIAL PARA LEITORES

Comece a usar hoje mesmo.

Plano Start **grátis para sempre** (até 3 pacientes ativos). Sem cartão de crédito.

idio.com.br

@idio.psi



CAPÍTULO 4

Uso seguro da IA



Objetivos de aprendizagem

Ao final deste capítulo, você será capaz de:

- Aplicar técnicas de **anonimização** e **pseudonimização** a casos clínicos antes de usar qualquer IA.
- Identificar com precisão **o que nunca deve ser inserido** em uma ferramenta de IA.
- Configurar suas contas para **reduzir a retenção de dados** e sair do uso para treino.
- Reconhecer e evitar os **erros mais comuns** que levam à quebra de sigilo.

4.0 • A regra de ouro



Trate todo chatbot como uma sala de espera lotada

Tudo o que você “fala” ali pode ser ouvido, guardado e processado por terceiros. Você nunca diria o nome e a história de um paciente em voz alta numa sala cheia — então não digite isso num chat.

O Capítulo 3 mostrou *por que* o sigilo se aplica à IA (Código de Ética, Art. 9º; Res. CFP 9/2024, Art. 3º II/IV; LGPD). Este capítulo mostra **como** cumprir esse dever na prática, passo a passo. A literatura sobre uso responsável de IA em saúde mental converge para um ponto: a proteção de dados não é uma etapa final, é a **condição de entrada** (Saeidnia et al., 2024; Stade et al., 2024).

4.1 • Anonimização e pseudonimização de casos

Há duas formas de descaracterizar um dado antes de usá-lo numa IA. Entender a diferença evita uma falsa sensação de segurança.

Técnica	O que é	Reversível?	Quando usar
Anonimização	Remover todos os elementos que permitem identificar a pessoa, direta ou indiretamente	Não (idealmente)	Sempre que possível — é o padrão-ouro
Pseudonimização	Substituir identificadores por códigos/apelidos, mantendo uma chave separada	Sim, com a chave	Quando você precisa recontactar ao caso depois — a chave nunca vai para a IA

**Cuidado com a reidentificação**

Anonimizar não é só trocar o nome. Um conjunto de detalhes aparentemente inofensivos — profissão rara + cidade pequena + idade + evento específico — pode identificar alguém tão bem quanto um CPF. **Anonimização real remove também o contexto identificável.**

Exemplo lado a lado**Identificável (NÃO inserir)**

“Meu paciente João Carlos, 34 anos, dentista em [cidade pequena], casado com a Fernanda, relatou que desde a morte do pai em março piora a insônia. Trabalha na Clínica OdontoVida.”

**Anonimizado (uso seguro)**

“Um paciente adulto, profissional da saúde, relata piora de insônia associada a um luto recente (falecimento de um familiar próximo há alguns meses). Como organizar a psicoeducação sobre higiene do sono para esse contexto?”

Repare que a versão segura **preserva o que importa clinicamente** (luto recente, insônia, necessidade de psicoeducação) e **descarta o que só serve para identificar** (nome, cônjuge, cidade, local de trabalho).

4.2 • O que NUNCA inserir em uma IA

Uma lista mental simples, para usar antes de cada uso:

**Os “nunca” do sigilo digital**

Nome completo, apelido único ou iniciais reconhecíveis do paciente



CPF, RG, número de prontuário, telefone, e-mail, endereço



Nome de familiares, empregador, escola, clínica



Cidade pequena + combinação de traços que estreite a identificação



Imagens, áudios ou prints da sessão com rosto/voz



Datas exatas que, somadas a outros dados, identifiquem (ex.: “deu à luz em 12/03 no Hospital X”)



Qualquer trecho que, lido por um conhecido do paciente, revelaria de quem se trata

O teste prático: “Se esse texto vazasse amanhã, alguém conseguiria descobrir quem é o paciente?” Se a resposta não for um “não” confiante, **não cole**.

Esse cuidado é a tradução direta do dever de sigilo (Código de Ética, Art. 9º) e do tratamento reforçado de dado sensível de saúde pela LGPD (Cap. 3). A literatura que cataloga violações éticas de “conselheiros de IA” aponta a exposição inadvertida de dados como uma das falhas mais frequentes (Iftikhar et al., 2025).

4.3 • Boas práticas de configuração

Anonimizar é a primeira camada. A segunda é **configurar a ferramenta** para reduzir o rastro que ela guarda.

a) Saia do treino do modelo

A maioria das ferramentas grandes permite **desativar o uso das suas conversas para treinar o modelo**. Procure em *Configurações* → *Dados / Privacidade* e desative. Os nomes exatos da opção, verificados em junho de 2026 (ver tabela do Cap. 6):

- **ChatGPT** (OpenAI): *Configurações* → *Controles de dados* → “Improve the model for everyone”.
- **Claude** (Anthropic): *Configurações* → *Privacidade* → “Help Improve Claude”.
- **Gemini** (Google): *Conta Google* → *Dados e privacidade* → “Atividade dos apps Gemini”.
- **Perplexity**: *Configurações* → *Preferências* → “AI data retention” (vem ligado por padrão).
- **Copilot** (Microsoft): *Configurações* → *Privacidade* → “Training on conversation activity”.

Como esses nomes mudam com frequência, reconfira na própria ferramenta.

b) Prefira contas e planos com garantias de dados

Versões corporativas/educacionais costumam oferecer **compromisso de não usar dados para treino** e maior controle de retenção. Avalie a política antes de adotar (dever do Art. 4º, I da Res. 9/2024).

c) Gerencie histórico e retenção

- Apague conversas sensíveis após o uso.
- Use chats temporários/anônimos quando disponíveis.
- Não acumule no histórico aquilo que você não inseriria de novo hoje.

d) Separe o ambiente profissional do pessoal

Use uma conta dedicada ao trabalho de apoio clínico, sem misturar com buscas pessoais — facilita controlar o que entra e o que fica.



Importante

Desativar o treino **não** torna seguro inserir dado identificável. É uma camada extra, não um substituto da anonimização. A ordem é sempre: **(1) anonimizar** → **(2) configurar** → **(3) usar**.

4.4 • Erros comuns (e como evitá-los)

Erro comum	Por que é perigoso	Como evitar
“Só troquei o nome”	Contexto ainda identifica (4.1)	Remova traços que estreitam a identificação
Colar a transcrição inteira da sessão	Mesmo sem nome, o relato pode ser único	Resuma e abstraia o caso antes
Subir áudio/print da sessão	Voz e imagem são identificadores diretos	Nunca subir mídia da sessão
Usar conta pessoal com treino ativo	O dado pode alimentar o modelo	Desative o treino; conta dedicada
Confiar na saída sem revisar	IA alucina e enviesa	Você valida e assina (Cap. 3, Art. 2º g)
Achar que “apagar” some na hora	Pode haver retenção temporária no servidor	Não inserir é mais seguro que apagar depois
Manejar fala de crise no chat	IA não substitui rede de cuidado	Encaminhar à rede presencial (Res. 9/2024, Art. 5º)



Checklist — antes de colar no chat

- ☐ Removi nome, contatos e documentos?
- ☐ Removi familiares, empregador, instituições?
- ☐ O conjunto de detalhes ainda permite identificar? (se sim, abstraia mais)
- ☐ Não há áudio, imagem ou print da sessão?
- ☐ O treino do modelo está desativado nesta conta?
- ☐ Eu validarei e assumirei a responsabilidade pela saída?
- ☐ Isto é uma tarefa de **apoio** (não diagnóstico, parecer ou crise)?

Se todos estiverem marcados, o uso está dentro das boas práticas. Se algum falhar, **pare e ajuste**.

EM UMA FRASE

Uso seguro de IA na clínica = **anonimizar sempre, configurar a ferramenta** para reter e treinar o mínimo, e **nunca delegar à máquina** a responsabilidade técnica e o sigilo que são, por lei e por ética, do psicólogo.

Referências deste capítulo

- Saeidnia et al. (2024), *Social Sciences* — considerações éticas centrais no uso de IA em saúde mental.
- Stade et al. (2024), *npj Mental Health Research* — roadmap de desenvolvimento e uso responsável.
- Iftikhar et al. (2025), *Proc. AAAI/ACM AIES* — catálogo de violações éticas de “LLM counselors”.

• Holm (2024), *Frontiers in Psychiatry* — trade-offs éticos e limites de objetividade.

Normativas aplicadas: Código de Ética (Art. 9º — sigilo); Res. CFP 9/2024 (Art. 3º II/IV, Art. 4º I, Art. 5º); LGPD (dado sensível de saúde, segurança). *Detalhamento no Cap. 3.*



Nota de validação

Os nomes exatos das opções de privacidade de cada ferramenta (4.3) foram **verificados nas páginas oficiais em junho de 2026** (mesma base da tabela do Cap. 6). Como mudam com frequência, reconfira na própria ferramenta antes de usar.



INSTITUTO WALDEN4 · RÁDIO IW4

Rádio iW4 — ciência e conhecimento ao seu alcance

Rádio 24 horas com leituras de artigos científicos, palestras, aulas e muito mais, dedicada à Análise do Comportamento e ciências afins.

- Sempre disponível · 24/7
- Conteúdo para desenvolvimento profissional
- No navegador e no app Android (Google Play)

Conhecimento de qualidade, sempre que você quiser.

Ouçá agora, é gratuito →

walden4.com.br/radio · [@instituto.walden4](https://www.instagram.com/instituto.walden4)

INSTITUTO WALDEN4 · CIÊNCIA QUE CUIDA

CAPÍTULO 5

Prompts para psicólogos



Objetivos de aprendizagem

Ao final deste capítulo, você será capaz de:

- Utilizar **prompts seguros** para tarefas de apoio, reconhecendo seus limites.
- Estruturar um bom prompt (contexto → instrução → formato).
- Aplicar a **trava de segurança** e a **anonimização** a cada uso.

5.0 • Como ler este capítulo



Trava de segurança (vale para TODOS os prompts)

Nenhum prompt aqui faz — ou deve ser adaptado para fazer — **diagnóstico, avaliação psicológica, psicoterapia, parecer/laudo automático ou manejo de crise**. Todos executam **tarefas acessórias** que continuam sob sua responsabilidade técnica (Cap. 2 e 3). E todos pressupõem **dados anonimizados** (Cap. 4).

Cada prompt segue a mesma estrutura: **Contexto** (para que serve) · **Limite** (o que ele NÃO faz) · **Prompt pronto** (com botão “copiar”) · **Selo** ✓ *Uso de apoio*.

Anatomia de um bom prompt

Um prompt eficaz tem três partes: (1) **papel/contexto** (“Você é um assistente que apoia um psicólogo em...”), (2) **instrução clara** (o que fazer) e (3) **formato de saída** (lista, tabela, tom, extensão). A literatura sobre IA aplicada ao trabalho clínico de apoio mostra que instruções bem-estruturadas melhoram a utilidade e reduzem desvios (Stade et al., 2024; McCrudden et al., 2026).

5.1 • Documentos clínicos (apoio à redação, não substituição do registro)

Este conjunto **NÃO** serve para: **gerar laudo, parecer ou diagnóstico, nem substituir o prontuário**. Serve para **organizar e revisar a forma de um texto que você redige e valida** (Código de Ética, Art. 2º, g).

Prompt 1 — Revisar a clareza de um texto já escrito por você

Contexto: melhorar a redação de uma anotação que você mesmo escreveu, sem alterar o conteúdo clínico.

Limite: não cria conteúdo clínico; só revisa forma.

Você é um assistente de redação que apoia um psicólogo. Revise o texto abaixo apenas quanto à clareza, coesão e gramática, SEM acrescentar informações, interpretações ou conteúdo clínico novo. Mantenha o sentido original. Aponte ao final, em lista, o que alterou.

Texto: "[cole aqui o seu texto, já anonimizado]"

✓ Uso de apoio

Prompt 2 — Estruturar tópicos de uma anotação de evolução (modelo genérico)

Contexto: obter um esqueleto neutro para organizar suas próprias anotações.

Limite: gera **estrutura**, não conteúdo do caso.

Sugira uma estrutura de tópicos (apenas títulos de seções) para organizar uma anotação de evolução clínica genérica em psicologia. Não invente conteúdo de caso; apresente somente o esqueleto reutilizável.

✓ Uso de apoio

5.2 · Planejamento terapêutico (organização das ideias do psicólogo)

Este conjunto **NÃO** serve para: definir o tratamento no lugar do psicólogo nem prescrever conduta. Serve para **organizar e ampliar ideias** que você avalia criticamente.

Prompt 3 — Levantar possibilidades para você avaliar

Limite: a IA propõe; a decisão clínica é sua.

Você apoia um psicólogo no planejamento. Para o tema clínico abstrato abaixo (sem dados de paciente), liste possíveis eixos de trabalho psicoeducativo e organizacional que um profissional PODERIA considerar, apresentando-os como opções a avaliar criticamente — não como prescrição.

Tema (genérico e anonimizado): "[ex.: manejo de ansiedade de desempenho em adultos]"

✓ Uso de apoio

5.3 • Supervisão (estruturação de estudo de caso)

Este conjunto **NÃO** serve para: substituir o supervisor humano. Serve para **organizar a apresentação** de um caso já anonimizado para levar à supervisão. A literatura explora a IA como apoio à supervisão e ao treino — sempre complementar ao supervisor (Sheperis & Sadeh-Sharvit, 2023; Shafran et al., 2026).

Prompt 4 — Roteiro para apresentar um caso na supervisão

Ajude a organizar a APRESENTAÇÃO de um estudo de caso para supervisão em psicologia. Forneça um roteiro de seções (ex.: motivo de busca, hipóteses de trabalho, intervenções, dúvidas para o supervisor), em formato de checklist. Não preencha com conteúdo; só o roteiro.

✓ Uso de apoio

Prompt 5 — Gerar perguntas para levar ao supervisor

Com base no tema clínico genérico abaixo (sem dados de paciente), sugira perguntas reflexivas que um psicólogo poderia levar à supervisão para aprofundar sua própria análise. Apresente como perguntas, não respostas.

Tema: "[genérico e anonimizado]"

✓ Uso de apoio

5.4 • Psicoeducação (materiais para pacientes)

Este conjunto **NÃO** serve para: orientar um paciente específico sem a curadoria do psicólogo. Serve para **rascunhar materiais educativos genéricos** que você revisa, adapta e valida antes de entregar.

Prompt 6 — Rascunho de material psicoeducativo genérico

Você apoia um psicólogo na criação de material psicoeducativo. Escreva um texto curto, acessível e acolhedor (linguagem leiga, sem jargão) explicando, de forma geral, o tema abaixo para o público adulto. Inclua um aviso de que o material é informativo e não substitui acompanhamento profissional. Evite prometer resultados.

Tema: "[ex.: o que é higiene do sono]"

✓ **Uso de apoio**

Revise e adapte antes de entregar a qualquer paciente.

5.5 • Redes sociais / comunicação profissional

Este conjunto NÃO serve para: fazer divulgação sensacionalista, prometer resultados ou expor casos. Deve respeitar o Código de Ética na publicidade (Art. 20: identificação com CRP, sem previsão taxativa de resultados, sem sensacionalismo).

Prompt 7 — Rascunho de post educativo para o público geral

Escreva um rascunho de post curto para redes sociais, de um psicólogo, com teor EDUCATIVO sobre o tema abaixo. Regras:

- linguagem acessível e respeitosa;
- NÃO prometa resultados nem use tom sensacionalista;
- inclua um convite leve a procurar acompanhamento profissional;
- deixe um espaço "[Nome - CRP 00/00000]" para identificação.

Tema: "[ex.: a importância das pausas no trabalho]"

✓ **Uso de apoio**

Confira a conformidade com o Art. 20 antes de publicar.

5.6 • Antes de usar qualquer prompt (resumo)**Confira sempre**

- ☐ O conteúdo está **anonimizado** (Cap. 4)?
- ☐ A tarefa é de **apoio** (não diagnóstico/parecer/crise)?
- ☐ Vou **revisar e assumir a responsabilidade** pela saída?
- ☐ O resultado respeita **sigilo, viés e transparência** (Cap. 2)?

EM UMA FRASE

Um bom prompt para psicólogo é aquele que **economiza tempo em tarefas acessórias** — revisar, organizar, rascunhar — **sem nunca tocar** no que é insubstituível: o juízo clínico, o vínculo e a responsabilidade técnica.

Referências deste capítulo

- Sheperis & Sadeh-Sharvit (2023), *J. Technology in Counselor Education & Supervision* — supervisão apoiada por IA.
- McCrudden et al. (2026), *JMIR Formative Research* — documentação clínica com IA.
- Shafran et al. (2026), *Behaviour Research and Therapy* — IA em treino/supervisão de TCC.
- Stade et al. (2024), *npj Mental Health Research* — desenvolvimento e uso responsável.

Normativas aplicadas: Código de Ética (Art. 2º g, 2º q, 9º, 20); Res. CFP 9/2024 (Art. 5º); LGPD. Detalhamento no Cap. 3 e 4.



EDITORA WALDEN4 · LIVROS

Editora Walden4 — sobre o ombro de gigantes

Publicações especializadas em Análise do Comportamento e práticas baseadas em evidências científicas, conectando a ciência psicológica ao público.

- 13 anos de experiência · 62 livros publicados
- Da sua ideia à publicação, com curadoria editorial
- Disponível no Google Play e em PDF

Leve a ciência do comportamento para a sua estante.

Conheça o catálogo →

walden4.com.br/livros · [@instituto.walden4](https://www.instagram.com/instituto.walden4)

INSTITUTO WALDEN4 · CIÊNCIA QUE CUIDA

CAPÍTULO 6

Ferramentas de IA



Objetivos de aprendizagem

Ao final deste capítulo, você será capaz de:

- Comparar as principais ferramentas de IA segundo **critérios relevantes ao psicólogo**.
- Avaliar a **política de dados** de uma ferramenta antes de adotá-la (dever da LGPD e da Res. CFP 9/2024, Art. 4º, I).
- Escolher a ferramenta adequada à **tarefa de apoio** pretendida.

6.0 • Como comparar com responsabilidade

Antes da tabela, três avisos:

1. **Privacidade vem primeiro.** Para o psicólogo, o critério decisivo não é “qual é mais inteligente”, e sim **o que a ferramenta faz com os dados**. Avaliar isso é dever normativo (Res. CFP 9/2024, Art. 4º, I; LGPD).
2. **Tudo muda rápido.** Preços, limites e políticas de dados de IA mudam de um mês para o outro. Os valores abaixo foram **verificados nas páginas oficiais de cada produto em junho de 2026** e devem ser reconferidos na fonte antes de qualquer decisão.
3. **Comparar é factual, não promocional.** Descrevemos finalidades típicas observáveis e verificadas; **não inventamos capacidades** (regra anti-alucinação do projeto).

A literatura que revisa LLMs em saúde mental reforça que a escolha de ferramenta deve considerar adequação à tarefa, limites e segurança — não apenas desempenho (Guo et al., 2024; Tan et al., 2025).

6.1 • Tabela comparativa

⚠ Os campos de **custo** e **política de dados** mudam com frequência — valores e nomes de opções **verificados nas páginas oficiais em junho de 2026**; reconfira na fonte antes de decidir. Clique nos cabeçalhos para ordenar; use o campo abaixo para filtrar.

Ferramenta	Finalidade típica	PT-BR	Custo (modelo)	Política de dados / treino	Adequação a tarefas de apoio
ChatGPT (OpenAI)	Assistente geral de texto: rascunho, revisão, organização	Bom	Gratuito; pagos Go (US\$ 8/mês) e Plus (US\$ 20/mês)	Permite desativar o uso para treino em Configurações → Controles de dados → “Improve the model for everyone”	Redação e organização de materiais (Cap. 5)
Claude (Anthropic)	Assistente geral de texto, forte em textos longos e análise cuidadosa	Bom	Gratuito (Sonnet 4.5); Pro US\$ 20/mês	Permite desativar em Configurações → Privacidade → “Help Improve Claude”; conversas sinalizadas para revisão de segurança ainda podem ser usadas	Revisão e síntese de textos longos (Cap. 5)
Gemini (Google)	Assistente geral integrado ao ecossistema Google	Bom	Gratuito; plano pago via Google AI Pro	Desative em Conta Google → Dados e privacidade → “Atividade dos apps Gemini”	Apoio a redação e pesquisa
NotebookLM (Google)	Trabalhar sobre documentos que você fornece (resumir, perguntar)	Bom	Gratuito; recursos ampliados no NotebookLM Plus (assinatura Google AI)	Para uso individual, suas fontes não são usadas para treino, salvo se você enviar feedback ao Google	Estudar materiais/artigos que você já tem
Perplexity	Busca com IA e citação de fontes	Bom	Gratuito; Pro US\$ 20/mês	Usa dados para treino por padrão; desative em Configurações → Preferências → “AI data retention”	Pesquisa com rastreabilidade de fontes
Copilot (Microsoft)	Assistente integrado ao Windows/Office	Bom	Gratuito; recursos pagos integrados ao Microsoft 365	Desative o treino em Configurações → Privacidade → “Training on conversation activity”	Apoio à produtividade em documentos

Selo de privacidade

Cada linha traz um selo —  “permite desativar treino”,  “verificar política” — conforme a coluna de política de dados.

6.2 • Critérios para escolher (independente da moda)

Ao decidir, pergunte-se, nesta ordem:

1. **Privacidade** — a ferramenta permite desativar o uso para treino? Onde os dados ficam? (Cap. 4)
2. **Finalidade** — a tarefa é de redação, de pesquisa com fontes, ou de trabalhar sobre documentos seus? Ferramentas diferentes brilham em coisas diferentes.
3. **Idioma** — qualidade em PT-BR para o seu uso.
4. **Custo × necessidade** — o plano gratuito já resolve a sua tarefa de apoio?
5. **Rastreabilidade** — preciso de fontes citadas (Perplexity) ou de trabalhar meus próprios documentos (NotebookLM)?



Regra que não muda

Independentemente da ferramenta, **anonimize antes e valide depois**. A ferramenta é a parte fácil; a responsabilidade é sua (Cap. 2 e 3).

6.3 • Qual usar para quê (guia rápido)

Sua tarefa de apoio	Ponto de partida sugerido
Revisar/organizar um texto seu	ChatGPT ou Claude
Trabalhar em cima de artigos/PDFs que você tem	NotebookLM
Pesquisar algo com fontes citadas	Perplexity
Apoio dentro do Office/Windows	Copilot
Síntese de textos longos	Claude

Sugestões de ponto de partida, não recomendações exclusivas — todas exigem a mesma checagem de privacidade.

EM UMA FRASE

A melhor ferramenta de IA para o psicólogo é aquela cuja **política de dados você conhece e aprova** e que se ajusta à **tarefa de apoio** específica — a “inteligência” do modelo é secundária diante do sigilo.

Referências deste capítulo

- Guo et al. (2024), JMIR Mental Health — revisão sistemática de LLMs em saúde mental.
- Tan et al. (2025), Journal of Psychology and AI — competência de LLMs em contextos de psicoterapia.

Norma aplicada: dever de avaliar a política de dados antes de usar (LGPD; Res. CFP 9/2024, Art. 4º, I). *Detalhamento no Cap. 3.*

**Nota de validação**

Os campos de custo e política de dados foram verificados nas páginas oficiais de cada produto em junho de 2026. Dado o ritmo de mudança do setor, reconfira na fonte antes de qualquer decisão.



INSTITUTO WALDEN4 · AUTISMO (TEA)

Atendimento Especializado em Autismo (TEA)

Cuidado integral e multidisciplinar baseado em evidências, para crianças e adolescentes com Transtorno do Espectro Autista — e as famílias que caminham com eles.

- 8 especialidades: ABA, fonoaudiologia, terapia ocupacional, nutrição e mais
- Programas individualizados para cada fase do desenvolvimento
- +15 anos de experiência · foco total na família

Cada criança merece o melhor suporte.

Conheça os programas →

walden4.com.br/autismo · [@instituto.walden4](https://www.instagram.com/instituto.walden4)

INSTITUTO WALDEN4 · CIÊNCIA QUE CUIDA

CAPÍTULO 7

Futuro da Psicologia com IA



Objetivos de aprendizagem

Ao final deste capítulo, você será capaz de:

- Analisar criticamente **cenários emergentes** do uso de IA na Psicologia.
- Pesquisar **potencial** × **risco** de cada cenário.
- Reconhecer o que permanece **insubstituível** no trabalho do psicólogo.

7.0 • Como pensar o futuro sem ingenuidade

Falar do futuro da IA na Psicologia exige equilíbrio entre dois erros: o **deslumbramento** (achar que a máquina resolverá tudo) e o **negacionismo** (achar que nada mudará). A literatura mais recente adota um terceiro caminho — **cautela informada**: mapear o que é possível, o que é desejável e o que é arriscado (Stade et al., 2024; Zhang et al., 2025; Gupta & Gandhi, 2025).

Examinamos quatro cenários. Para cada um: o **potencial**, o **risco** e o **limite ético** que não se negocia.

7.1 • Supervisão assistida

O **cenário**: ferramentas de IA que ajudam o supervisor e o supervisionando a organizar casos, gerar perguntas reflexivas e estruturar feedback de treino (Sheperis & Sadeh-Sharvit, 2023; Shafran et al., 2026).

- **Potencial**: ampliar o acesso à supervisão, padronizar parte do treino, liberar tempo para a discussão que realmente importa.
- **Risco**: terceirizar o julgamento formativo; achatamento da singularidade do caso em respostas genéricas.
- **Limite ético**: a IA **apoia** a supervisão; o supervisor humano permanece responsável pela formação e pela avaliação.

7.2 • Documentação automática

O **cenário**: assistência à produção de registros e materiais a partir de informações fornecidas pelo profissional (McCrudden et al., 2026; Stade et al., 2024).

- **Potencial**: reduzir a carga burocrática, um fator real de sobrecarga e desgaste.
- **Risco**: registros genéricos ou imprecisos; **quebra de sigilo** se dados identificáveis forem inseridos; documentos sem fundamentação técnica própria.

- **Limite ético:** todo documento exige **fundamentação técnico-científica** e validação humana (Código de Ética, Art. 2º, g); dados sempre **anonimizados** (Cap. 4). A IA rascunha; o psicólogo assina.

7.3 • Simulação de pacientes para treino

O cenário: “pacientes virtuais” gerados por IA para estudantes praticarem entrevista, escuta e manejo, em ambiente seguro (Shafran et al., 2026; Tan et al., 2025).

- **Potencial:** treino repetível e de baixo risco, sem expor pessoas reais; espaço para errar e aprender.
- **Risco:** aprender padrões irreais (um paciente simulado não reage como um humano); confundir competência técnica com competência relacional.
- **Limite ético:** simulação é **complemento** do treino supervisionado, nunca substituto da experiência clínica real com supervisão.

7.4 • Riscos transversais e o lugar insubstituível do psicólogo

Além dos cenários, há riscos que atravessam todos eles e exigem vigilância contínua:

Risco transversal	Por quê	Como mitigar
Viés de automação	Confiar na máquina por ser fluente e rápida	Validação crítica sempre (Cap. 2)
Erosão do sigilo	Mais ferramentas = mais pontos de vazamento	Anonimização e política de dados (Cap. 4, 6)
Viés algorítmico	Modelos herdaram preconceitos dos dados	Leitura crítica das saídas (Cap. 2)
Desigualdade de acesso	Nem todos terão as mesmas ferramentas	Letramento e governança (Gupta & Gandhi, 2025)
Vácuo regulatório	A norma corre atrás da tecnologia	Governança ética e decisão humano-IA (Heyen & Salloch, 2021)

A taxonomia de papéis que os LLMs podem assumir em saúde mental converge para a mesma conclusão deste livro: **quanto mais a tarefa se aproxima do juízo clínico e do cuidado, menos cabe à IA e mais cabe ao humano** (Zhang et al., 2025).

EM UMA FRASE

O futuro da Psicologia com IA não é a substituição do psicólogo, e sim a **liberação do psicólogo das tarefas acessórias** para que ele se dedique ao que só um humano faz: cuidar, vincular-se e decidir com responsabilidade.

Referências deste capítulo

- Stade et al. (2024), *npj Mental Health Research* — roadmap de desenvolvimento responsável.
- Shafran et al. (2026), *Behaviour Research and Therapy* — IA em treino/supervisão.

- Zhang et al. (2025), *Nature Computational Science* — taxonomia de papéis de LLM e ética.
- Tan et al. (2025), *Journal of Psychology and AI* — competência de LLMs em psicoterapia.
- Heyen & Salloch (2021), *BMC Medical Ethics* — ética do apoio à decisão clínica baseada em aprendizado de máquina.
- Gupta & Gandhi (2025), *International Review of Psychiatry* — letramento em IA e futuro da área.
- Sheperis & Sadeh-Sharvit (2023), *J. Technology in Counselor Education & Supervision* — supervisão apoiada por IA.
- McCrudden et al. (2026), *JMIR Formative Research* — documentação clínica com IA.

Norma aplicada: reforço do posicionamento ético central + Código de Ética (Art. 2º, g). *Detalhamento no Cap. 3.*

REFERÊNCIAS

Referências & encerramento



- Bharti, J., & Singh, R. (2025). Between care and code: Clinicians' perspectives as an ethical lens for AI in mental health – Insights from an interview-based case study. *Science Insights*, 47(5), 2041–2048. <https://doi.org/10.15354/si.25.or1302>
- Guo, Z., Lai, A., Thygesen, J., Farrington, J., Keen, T., & Li, K. (2024). Large language models for mental health applications: Systematic review. *JMIR Mental Health*, 11, e57400. <https://doi.org/10.2196/57400>
- Gupta, D. R., & Gandhi, T. (2025). AI and ethics in mental health: Personal reflections on the future of psychiatry. *International Review of Psychiatry*, 38(1–3), 122–125. <https://doi.org/10.1080/09540261.2025.2600617>
- Heyen, N., & Salloch, S. (2021). The ethics of machine learning-based clinical decision support: An analysis through the lens of professionalisation theory. *BMC Medical Ethics*, 22(1), Article 112. <https://doi.org/10.1186/s12910-021-00679-3>
- Holm, S. (2024). Ethical trade-offs in AI for mental health. *Frontiers in Psychiatry*, 15, Article 1407562. <https://doi.org/10.3389/fpsyt.2024.1407562>
- Iftikhar, Z., Xiao, A., Ransom, S., Huang, J., & Suresh, H. (2025). How LLM counselors violate ethical standards in mental health practice: A practitioner-informed framework. *Proceedings of the AAAI/ACM Conference on AI, Ethics, and Society*, 8(2), 1311–1323. <https://doi.org/10.1609/aies.v8i2.36632>
- Irshad, S., Azmi, S., & Begum, N. (2022). Uses of artificial intelligence in psychology. *Journal of Health and Medical Sciences*, 5(4). <https://doi.org/10.31014/aior.1994.05.04.242>
- Komatsu, H., Watanabe, E., & Fukuchi, M. (2021). Psychiatric neural networks and precision therapeutics by machine learning. *Biomedicines*, 9(4), 403. <https://doi.org/10.3390/biomedicines9040403>
- McCrudden, K. E., Swirbul, M. S., Peake, E. E., Rodio, M. J., & Padmanabhan, A. (2026). AI-powered documentation for mental health providers: Retrospective observational mixed methods study. *JMIR Formative Research*, 10, e84628. <https://doi.org/10.2196/84628>
- Mesquiti, S., & Nook, E. C. (2026). Large language models in mental health research and treatment. *Annual Review of Biomedical Data Science*. Advance online publication. <https://doi.org/10.1146/annurev-biodatasci-092524-113527>
- Pillay, Y. (2025). Ethical decision-making guidelines for mental health clinicians in the artificial intelligence (AI) era. *Healthcare*, 13(23), 3057. <https://doi.org/10.3390/healthcare13233057>
- Saeidnia, H. R., Hashemi Fotami, S. G., Lund, B. D., & Ghiasi, N. (2024). Ethical considerations in artificial intelligence interventions for mental health and well-being: Ensuring responsible implementation and impact. *Social Sciences*, 13(7), 381. <https://doi.org/10.3390/socsci13070381>

- Shafraan, R., Bond, L., Carlbring, P., Cohen, Z., Creed, T. A., Davey, E., Egan, S. J., Freeman, D., Hollon, S. D., Jacobson, N. C., Johnson, C., Kaysen, D. L., McGuinness, D. L., Patel, V., Pozuelo, J. R., Santos, H., Singla, D. R., Stirman, S. W., Taylor, D., & Wade, T. D. (2026). From innovation to implementation: Artificial intelligence in cognitive behaviour therapy training and supervision. *Behaviour Research and Therapy*, 197, Article 104945. <https://doi.org/10.1016/j.brat.2025.104945>
- Sheperis, D., & Sadeh-Sharvit, S. (2023). Using AI-supported supervision in a university telemental health training clinic. *Journal of Technology in Counselor Education and Supervision*, 4(1). <https://doi.org/10.61888/2692-4129.1093>
- Singhal, S., Cooke, D. L., Villareal, R. I., Stoddard, J. J., Lin, C.-T., & Dempsey, A. G. (2024). Machine learning for mental health: Applications, challenges, and the clinician's role. *Current Psychiatry Reports*, 26(12), 694–702. <https://doi.org/10.1007/s11920-024-01561-w>
- Stade, E. C., Stirman, S. W., Ungar, L. H., Boland, C., Schwartz, H. A., Yaden, D. B., Sedoc, J., DeRubeis, R. J., Willer, R., & Eichstaedt, J. C. (2024). Large language models could change the future of behavioral healthcare: A proposal for responsible development and evaluation. *npj Mental Health Research*, 3(1), Article 12. <https://doi.org/10.1038/s44184-024-00056-z>
- Tan, K. S., Cervin, M., Leman, P., Nielsen, K., Kumar, P., & Medvedev, O. N. (2025). AI meets psychology: An exploratory study of large language models' competence in psychotherapy contexts. *Journal of Psychology and AI*, 1(1), Article 2545258. <https://doi.org/10.1080/29974100.2025.2545258>
- Zhang, R., Meng, H., Neubronner, M., & Lee, Y.-C. (2025). Computational and ethical considerations for using large language models in psychotherapy. *Nature Computational Science*, 5(10), 854–862. <https://doi.org/10.1038/s43588-025-00874-x>

Encerramento

Este livro nasceu de uma convicção simples: **a melhor resposta à chegada da Inteligência Artificial na Psicologia não é proibir nem idolatrar, mas compreender**. Ao longo dos sete capítulos, procuramos mostrar que a IA pode aliviar o psicólogo de tarefas acessórias — organizar, revisar, rascunhar, pesquisar — sem jamais tocar naquilo que é o coração da profissão: o vínculo, a escuta, o julgamento clínico e a responsabilidade ética.

As normas brasileiras — o Código de Ética, a Resolução CFP nº 9/2024 e a LGPD — não foram escritas para a IA, mas já oferecem um mapa seguro: sigilo rigoroso, competência técnica, transparência com o paciente e responsabilidade indelegável. A evidência científica reunida aqui aponta na mesma direção: quanto mais perto do cuidado, menos lugar para a máquina e mais lugar para o humano.

Que este material sirva de bússola — não de manual fechado — para um uso da tecnologia que seja, antes de tudo, **ético, seguro e a serviço das pessoas**.

Obra produzida com auxílio de Inteligência Artificial (Claude Opus 4.8, da Anthropic), sob revisão humana integral.
Uso educacional — não substitui a prática clínica, a supervisão profissional nem a assessoria jurídica/ética formal.

Inteligência Artificial na Prática Clínica — Um guia ético e prático para psicólogos

Fabio Hernandez de Medeiros e coautores · 1ª edição · Brasília: Editora do Instituto Walden4, 2026



INSTITUTO WALDEN4 · PSQUIATRIA INFANTIL

Psiquiatria para Crianças e Adolescentes

Atendimento psiquiátrico especializado em crianças e adolescentes, presencial e online — diagnóstico cuidadoso e tratamento baseado em evidências, integrado à equipe multidisciplinar.

- TDAH, ansiedade, humor e outras condições da infância
- Tratamento medicamentoso responsável, com acompanhamento contínuo
- Presencial (Asa Norte) e online, por telepsiquiatria

Cada fase do desenvolvimento merece o cuidado certo.

Agende uma consulta →

walden4.com.br/psiquiatra · [@instituto.walden4](https://www.instagram.com/instituto.walden4)

INSTITUTO WALDEN4 · CIÊNCIA QUE CUIDA