

Competènc+IA

Com la IA et pot ajudar en el teu dia a dia

Butlletí d'Intel·ligència Artificial

Número 5 | Maig 2026

Organitza

Societat Catalana | Xarxa d'Unitats |
de Lípids i Arterioesclerosi

sòcxula



L'Acadèmia
FUNDACIÓ ACADÈMIA DE CIÈNCIES MÈDIQUES
I DE LA SALUT DE CATALUNYA I DE BALEARIS



Aquesta entitat dona suport als Objectius de Desenvolupament Sostenible

Amb la col·laboració de Daiichi-Sankyo



Daiichi-Sankyo



Editorial

El salt cap a un ús més avançat de ChatGPT

En els primers números de *Competènc+IA* hem anat construint les bases per utilitzar ChatGPT amb criteri professional. Vam començar pel més essencial: entendre que la qualitat d'una resposta depèn molt de la qualitat del prompt. Per això vam introduir la fórmula **CRIDA** –context, rol, instrucció, dades i acotació– com una manera senzilla d'estructurar millor les peticions i obtenir respostes més útils, precises i aplicables.

Després vam veure que ChatGPT no s'ha d'utilitzar només “tal com ve de sèrie”. **Configurar** bé l'eina és clau per adaptar-la al nostre perfil clínic, docent o investigador. Personalitzar l'estil de resposta, definir instruccions personalitzades, gestionar la memòria, revisar el control de dades i aplicar criteris bàsics de seguretat permet passar d'una IA genèrica a una IA molt més alineada amb les nostres necessitats professionals.

En el número següent vam fer un pas més i vam abordar l'organització del treball amb **Projectes**. Aquesta funcionalitat permet agrupar xats, fonts, documents i instruccions sota un mateix objectiu, cosa especialment útil quan treballem en tasques que duren dies o setmanes: preparar una revisió científica, una sessió clínica, un protocol, una proposta de recerca o material docent. També vam veure com gestionar millor els xats per mantenir l'espai de treball més ordenat i eficient.

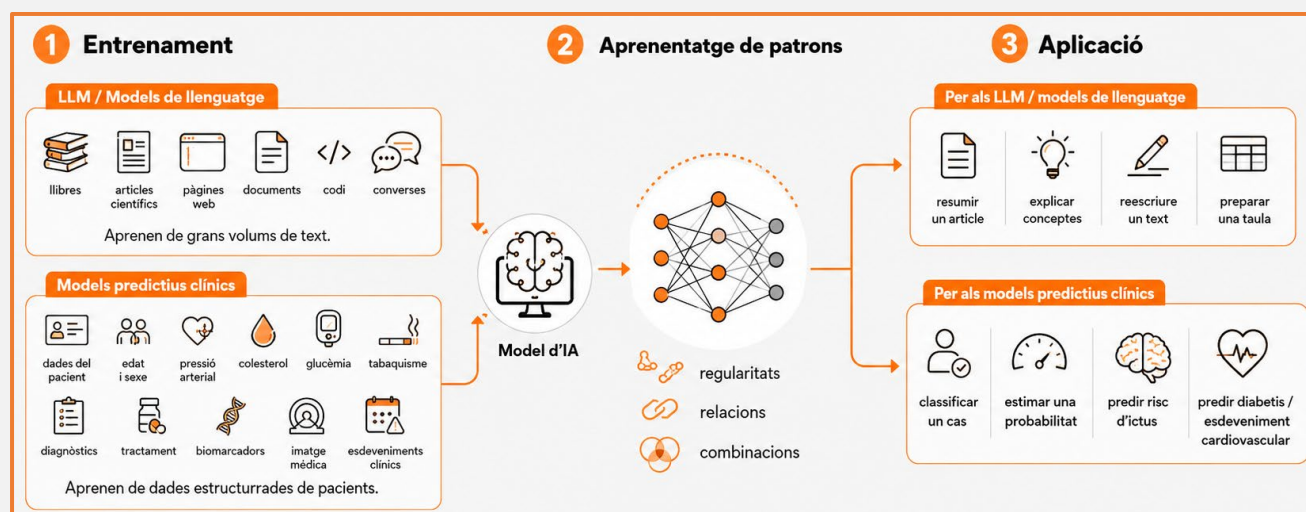
Finalment, en el número anterior vam introduir el menú **Apps**, que permet connectar ChatGPT amb eines externes. Aquest punt és important perquè ChatGPT deixa de ser només una finestra de conversa i comença a funcionar com un entorn de treball connectat, capaç d'interactuar amb serveis digitals, cercadors acadèmics, eines de disseny, plataformes formatives o aplicacions de productivitat.

El número de maig fa un nou pas en aquesta evolució. Un cop apreses les bases, és el moment d'entrar en eines que permeten treballar amb ChatGPT d'una manera més avançada, estructurada i professional. Ens centrarem en quatre funcionalitats especialment rellevants per a l'ús professional: **Recerca profunda, Llenç, Xats de grup i GPTs personalitzats**. Totes quatre representen un canvi important en la manera d'utilitzar la IA: permeten cercar informació amb més profunditat, redactar documents llargs amb més control, treballar de manera col·laborativa i crear assistents especialitzats per a tasques concretes. L'objectiu és que la IA deixi de ser només una eina que respon i es converteixi en un entorn de treball que ajuda a pensar, organitzar, crear i col·laborar millor.

Entenent la IA: la qualitat de les dades d'aprenentatge determina la qualitat de la IA.

La intel·ligència artificial pot semblar gairebé màgica: li donem dades i fa prediccions, li fem una pregunta i ens ofereix una resposta coherent en pocs segons. La IA aprèn a partir dels exemples amb què ha estat entrenada. Entrenar una IA vol dir exposar-la a grans volums d'informació perquè identifiqui patrons, relacions i combinacions de dades. Després utilitza aquests patrons per generar una resposta, classificar un cas o estimar una probabilitat.

Els models de llenguatge, com ChatGPT, Gemini, Claude o Copilot, s'entrenen sobretot amb grans volums de text. En canvi, els models predictius clínics necessiten dades estructurades de pacients: edat, sexe, pressió arterial, colesterol, glucosa, tabaquisme, diagnòstics, tractaments, biomarcadors, imatges o esdeveniments clínics. El seu objectiu és detectar combinacions de variables associades a un resultat, com predir un ictus, diabetis o un esdeveniment cardiovascular.



Idea clau



La qualitat de la resposta o de la predicció depèn de la qualitat de les dades amb què la IA ha après. Un LLM necessita bons textos per respondre millor. Un model clínic predictiu necessita bones dades de pacients per predir millor.

Dades reals representatives i ben recollides



Models útils



Prediccions precises

Dades errònies, incompletes, esbiaixades o artificials



Models erronis



Prediccions errònies, esbiaixades o artificials

Quan les dades semblaven reals, però no ho eren

Una notícia recent publicada a *Nature* il·lustra molt bé aquest problema. Diversos investigadors han detectat que desenes de models d'IA per predir ictus i diabetis havien estat entrenats amb bases de dades d'origen dubtós, disponibles en plataformes obertes com Kaggle. Algunes d'aquestes dades semblaven clíniques, però presentaven patrons molt poc compatibles amb cohorts reals: gairebé cap valor perdut, distribucions massa regulars i variables fisiològiques amb una variabilitat artificial.

El problema, segons *Nature*, és que alguns d'aquests models podrien haver arribat a utilitzar-se en entorns clínics, incloent hospitals a Espanya i Indonèsia, i fins i tot apareixen vinculats a eines web de càlcul de risc o a sol·licituds de patent.

Aquest cas posa en evidència una debilitat important del nostre ecosistema científic:

Avaluem el model utilitzat

Arquitectura del model
Corba ROC i AUC
Algoritme

Però no avaluem

D'on surten les dades	Validació externa
Qui les ha recollit	Traçabilitat clínica
Si són pacients reals	Són representatives



Idea clau:

La IA és una aliada potent però només si es construeix sobre dades reals, verificades, representatives i clínicament significatives.

Basu M. [Dozens of AI disease-prediction models were trained on dubious data](#). *Nature*, 15 abril 2026.

Entenent la IA: Cost ecològic d'utilitzar la IA. Estem fent un ús proporcional, necessari i eficient?

Quan fem servir ChatGPT, Gemini, Copilot o qualsevol altre model d'IA, la sensació és que tot passa "al núvol". Escrivim una pregunta, reps una resposta i sembla que el procés sigui gairebé màgic. Però el núvol no és immaterial sinó que són centres de dades, servidors, xips, electricitat, refrigeració, aigua i materials crítics.

La IA ens pot ser de molta ajuda, però també té una petjada ecològica. Per això, cada vegada serà més important valorar el seu ús.

Algunes xifres per posar-hi escala

Un estudi recent de Google sobre Gemini Apps¹ va estimar que un prompt de text mitjà consumia aproximadament 0,24 Wh, generava 0,03 g de CO₂ equivalent i utilitzava uns 0,26 mL d'aigua, l'equivalent aproximat a cinc gotes. En aquest sentit, si extrapolem, generar una imatge podria consumir de 3 a 12 mL d'aigua en funció dels models usats i la resolució de la imatge.

D'on surt el cost ambiental de la IA?

El cost ecològic de la IA no prové només d'una conversa individual, sinó de tota la cadena tecnològica que fa possible que el model respongui.

Entrenament dels models

Abans de respondre, els models s'han d'entrenar amb grans volums de dades, utilitzant milers de xips especialitzats durant setmanes o mesos.

Inferència: cada pregunta també compta

Cada vegada que fem una consulta, el model calcula una resposta. Una pregunta pot tenir un cost petit, però milions de consultes diàries converteixen aquest cost en una magnitud rellevant.

Refrigeració dels centres de dades

Els centres de dades i els servidors generen calor i necessiten sistemes de refrigeració. Això implica consum directe d'aigua o indirecte a través de la generació elèctrica. Això és especialment rellevant perquè els centres de dades no sempre se situen en zones amb abundància d'aigua.

Fabricació de xips i servidors

La IA necessita GPUs, memòria, xarxes i infraestructures. Produir aquest maquinari requereix minerals, metalls, energia i processos industrials complexos.

Renovació i residus electrònics

La cursa per models més grans i ràpids accelera la renovació d'equips i pot contribuir al creixement dels residus electrònics.

Idea clau



Una pregunta puntual a la IA té un impacte ecològic petit, però l'escala importa: milions de consultes diàries poden tenir un cost global rellevant. A més, no totes les tasques consumeixen igual: ordenar text té un cost baix, mentre que generar imatges, vídeos o respostes llargues i iteratives requereix molts més recursos.

La indústria també s'ha adonat del problema

Les grans empreses d'IA ja no competeixen només per crear models més grans, sinó també models més petits, ràpids i eficients. OpenAI ha impulsat versions mini i nano, Google ha presentat models Flash-Lite i els models oberts xinesos continuen apostant per arquitectures més compactes.

Aquest moviment busca reduir latència, costos i consum computacional. En altres paraules: fer una IA més pràctica, escalable i potencialment més sostenible.

Com podem fer un ús més sostenible de la IA?

No cal deixar d'utilitzar la IA. Cal utilitzar-la millor.

1. Utilitza la IA quan aporti valor real

No totes les tasques necessiten IA. Reserva-la per allò que realment millora la qualitat, estalvia temps o ajuda a pensar millor.

2. Escribe prompts més concrets

Un bon prompt, amb context i objectiu clar, evita respostes vagues i redueix iteracions innecessàries.

3. Demana respostes ajustades a la necessitat

No sempre cal una resposta llarga. Pots demanar una resposta breu, una taula, cinc punts clau o un resum de 150 paraules.

4. Genera imatges o vídeos només quan siguin necessaris

Són recursos molt útils per a docència, recerca i divulgació, però acostumen a tenir més cost computacional que el text.

5. Tria el model adequat a la tasca

Per a tasques senzilles, no sempre cal utilitzar el model més potent. La competència digital també consisteix a saber escollir bé l'eina.

6. Reutilitza el context

Per a treballs llargs, utilitza Projectes, documents de referència i instruccions estables. Això evita repetir informació constantment i fa la interacció més eficient.

Missatge final



La IA és una eina molt potent però amb un cost ecològic important. Cada resposta depèn d'una infraestructura física que consumeix electricitat, aigua, xips i materials. De manera que l'hem d'incorporar amb criteri, de forma proporcional i responsable, i per a tasques on realment aporta valor

Domina la IA: funcionalitats avançades de ChatGPT que et poden canviar la manera de treballar

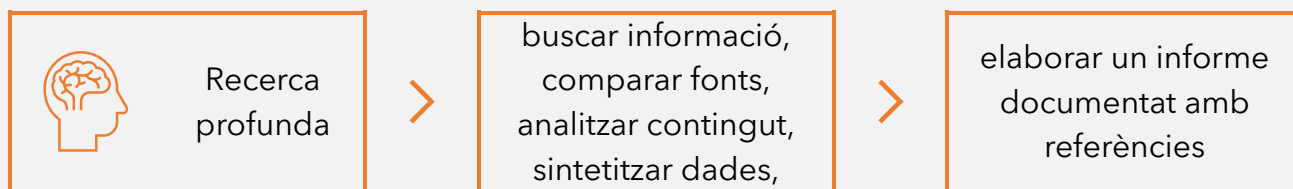
Fins ara hem vist com escriure bons prompts, configurar ChatGPT, personalitzar-ne les respostes, organitzar la feina amb Projectes i connectar-lo amb altres eines mitjançant Apps. Ara fem un pas més i entrem en algunes funcionalitats que permeten utilitzar ChatGPT d'una manera més avançada i professional.

Aquestes funcions permeten treballar amb més profunditat, més continuïtat i més control. Ens ajuden a buscar informació complexa, redactar documents llargs, compartir converses amb altres persones i crear assistents adaptats a tasques concretes.

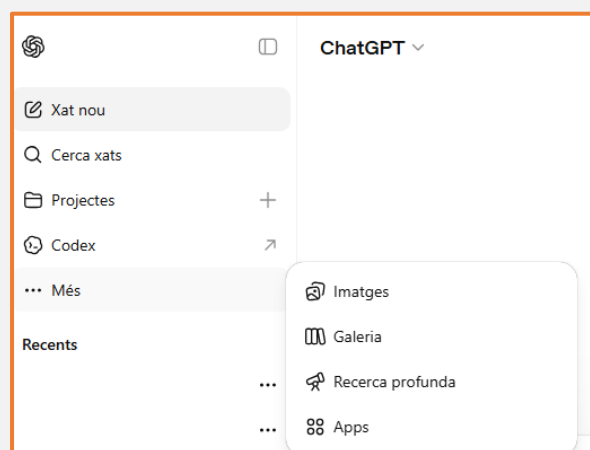
En aquest número veurem quatre funcionalitats especialment útils:

Recerca profunda: quan necessites anar més enllà d'una resposta ràpida

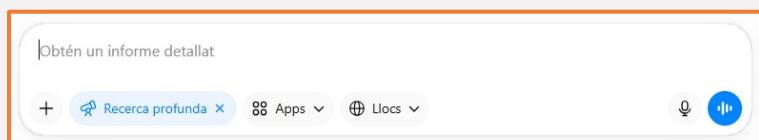
La funció **Recerca profunda** ("Deep Research") és una funcionalitat avançada de ChatGPT pensada per investigar preguntes complexes. Permet obtenir respostes molt més completes, contrastades i estructurades que una conversa habitual amb el xat. A diferència d'una resposta estàndard –que es genera gairebé immediatament–, la Recerca profunda dedica més temps a:



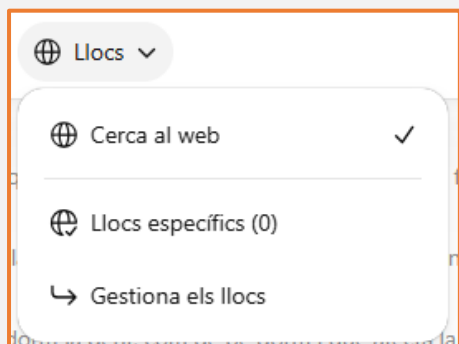
Com s'activa la recerca profunda?



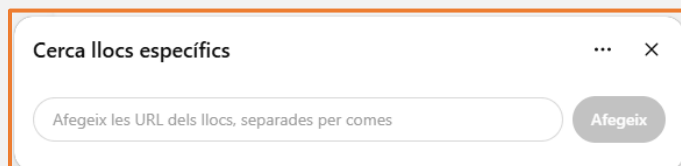
La recerca profunda es pot iniciar des del menú lateral, clicant a "**... Més**" i a "**Recerca profunda**" en el menú desplegable.



Un cop activada, la funcionalitat queda marcada a la barra d'escriptura.



L'usuari pot orientar on fer la recerca de la informació, ja sigui cercant a tota la web o a llocs específics definits per l'usuari.



El procés de cerca pot durar des d'uns segons fins a diversos minuts, depenent de la complexitat.

Quan val la pena utilitzar-la?

És útil usar-la quan la pregunta requereix temps, comparació i síntesi, donat que el tema és complex, es necessita informació actualitzada, cal comparar opcions o es volen fonts fiables. Es pot utilitzar per:

Preparar una revisió narrativa

Pot ajudar a obtenir una primera visió del tema, identificar conceptes clau, detectar línies de recerca recents i organitzar una estructura inicial.

Actualitzar-se sobre un tema biomèdic

Per exemple: novetats en Lp(a), nous fàrmacs hipolipemians, ús d'IA en diagnòstic clínic, risc cardiovascular en malalties inflamatòries o estratègies de prevenció secundària.

Comparar guies o recomanacions

Pot ser útil per contrastar què diuen diferents societats científiques sobre un mateix tema, sempre revisant després les fonts originals.

Preparar una sessió clínica o docent

Ajuda a reunir informació, ordenar-la i convertir-la en una estructura útil per explicar un tema.

Explorar un camp nou

Quan no sabem per on començar, pot generar una síntesi inicial amb conceptes, subtemes, controvèrsies i fonts rellevants.

Buscar context per a un projecte de recerca

Pot ajudar a redactar una justificació inicial, identificar buits de coneixement o formular preguntes de recerca, però no substitueix una revisió sistemàtica.

1

Prompt d'exemple

Utilitza Recerca profunda per resumir l'evidència recent sobre l'ús d'inhibidors de PCSK9 en prevenció secundària cardiovascular. Prioritza guies clíniques, assaigs clínics i revisions recents. Presenta els resultats en una taula amb tipus d'estudi, població, intervenció, resultats principals i limitacions

2

Prompt d'exemple

Fes una Recerca profunda sobre aplicacions de la intel·ligència artificial en la predicció del risc cardiovascular. Diferencia models basats en dades clíniques, imatge, òmiques i registres electrònics de salut. Inclou limitacions metodològiques i reptes d'implementació clínica

 Avantatges principals**Profunditat**

Les respostes acostumen a ser més extenses i raonades.

Millor estructuració

Organitza la informació en seccions, comparatives i conclusions.

Informació més actual

Pot consultar informació recent d'internet.

Síntesi de múltiples fonts

Combina diferents perspectives en una sola resposta.

Estalvi de temps

Redueix moltes hores de cerca manual.

 Precaucions importants**Cal fer lectura crítica**

Ajuda a trobar i sintetitzar informació, però la interpretació final continua sent responsabilitat de l'usuari.

Cal verificar fonts

Comprova que les referències existeixen, són rellevants i realment diuen el que ChatGPT afirma.

No és revisió sistemàtica

Serveix per explorar un tema, però no substitueix una cerca formal amb criteris d'inclusió, exclusió, doble revisió, etc.

Pot prioritzar fonts accessibles

Pot dependre de la disponibilitat de les fonts i no sempre recuperar tota la literatura rellevant.

Pot donar una falsa sensació d'exhaustivitat

Un informe ben estructurat pot semblar definitiu i deixar fora estudis importants o matisos metodològics.

En medicina, cal extremer la prudència

No s'ha d'utilitzar per prendre decisions clíniques sense contrastar-ho amb criteri professional.



Idea clau: "Recerca profunda" converteix ChatGPT d'un simple assistent conversacional en una eina d'investigació capaç d'analitzar, contrastar i sintetitzar grans quantitats d'informació de manera automatitzada.

Llenç ("Canvas"): quan necessites treballar un document, no només conversar

La funció **Llenç** és un espai de treball interactiu integrat dins del xat que permet crear, editar i desenvolupar contingut llarg o complex de manera més visual i estructurada. En lloc de treballar només amb missatges curts dins d'una conversa, **Llenç** obre una mena de "document lateral" on la IA i l'usuari poden redactar, revisar, modificar i reorganitzar continguts de forma continuada.

Xat normal

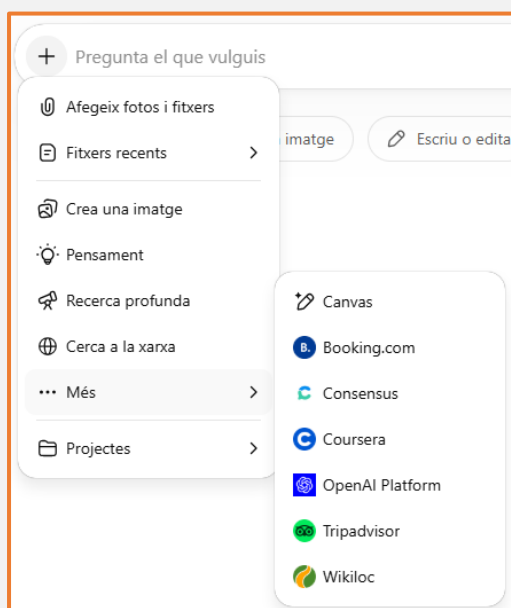
Cada nova instrucció genera una nova resposta. Això és poc còmode quan estem treballant amb documents llargs.



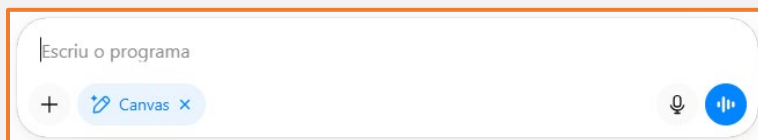
Llenç

Permet tenir un sol document editable en un espai lateral i demanar a ChatGPT que actuï sobre parts concretes del text. L'usuari té el control del document i pot editar directament el contingut demanant canvis específics, com escurçar, ampliar, reescriure, millorar el to o revisar fragments concrets.

Com s'activa?



El Llenç (Canvas) es pot activar de diferents maneres segons la interfície disponible. Actualment, en la versió gratuïta es pot activar des del menú "+" de la barra d'escriptura i fent clic a l'opció "**... Més**"



Com funciona?



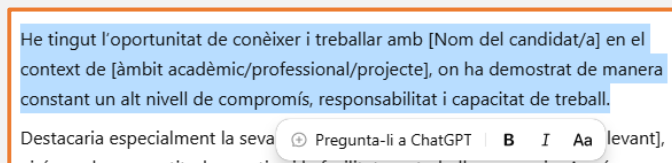
Quan demanes una tasca a ChatGPT amb la funcionalitat Llenç activada,

ChatGPT crea un espai de treball separat del xat principal fent clic a **"Modifica"**

El contingut apareix en format document i l'usuari pot modificar-lo directament com si estigués treballant en un processador de textos (tipus Word)

Quan treballem dins de Llenç, apareixen **eines contextuais** que permeten modificar el text sense haver d'escriure sempre prompts manualment. Aquestes funcions converteixen Llenç en una mena d'editor intel·ligent assistit per IA.

Les imatges següents mostren el funcionament d'aquest sistema.



En **seleccionar text**, apareix un menú flotant amb opcions com:

Pregunta-li a ChatGPT: permet demanar canvis sobre aquell fragment

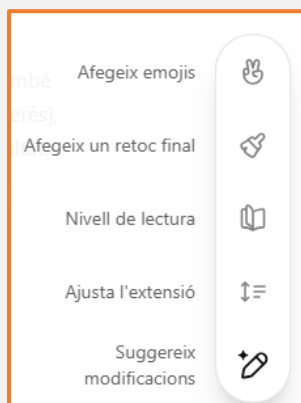
concret. **Negreta (B):** aplicar format en negreta. **Cursiva (I):** aplicar cursiva. **Aa:** opcions addicionals de format o estil.

Això evita instruccions globals i fa que la IA treballi només sobre la part seleccionada.



Botó flotant d'eines ràpides amb icona de llapis intel·ligent

Aquest botó obre el menú d'**assistència intel·ligent** de Llenç.



Aquest menú, permet accedir ràpidament a eines automàtiques d'edició i millora del document.

És una manera de fer servir la IA sense haver d'escriure prompts detallats en cada canvi. També permet treballar de manera molt més ràpida i iterativa

Menú d'assistència intel·ligent de Llenç

✦ **Afegeix emojis:** Afegeix emojis de manera contextual al text.

✦ **Afegeix un retoc final:** Fa una revisió general del text, millora la fluïdesa, corregeix l'estil i ajusta la redacció. És semblant a una "revisió final automàtica".

📖 **Nivell de lectura:** Adapta la dificultat i el registre del text perquè sigui més simple, més divulgatiu, més tècnic o més acadèmic. Per exemple, un text científic pot convertir-se en un text més divulgatiu només fent clic en aquesta icona.

↕ **Ajusta l'extensió:** Permet resumir, escurçar, ampliar o desenvolupar més el contingut. Molt útil per adaptar textos a límits de paraules, presentacions, publicacions o informes.

💡 **Suggereix modificacions:** La IA proposa millores, correccions, reorganització, idees addicionals o alternatives de redacció. Funciona com un "editor intel·ligent".

Idea clau



El Llenç converteix ChatGPT en un espai d'edició i construcció de documents on l'usuari i la IA poden crear, editar i desenvolupar documents o projectes complets de manera iterativa i estructurada. És especialment útil quan el resultat final no és només una resposta, sinó un text que cal escriure, revisar, millorar i deixar preparat per ser utilitzat.

Quan és útil utilitzar-lo?

Ideal per elaborar:

- articles,
- informes,
- treballs acadèmics,
- newsletters,
- manuals,
- guions,
- memòries,
- documentació tècnica.
- múltiples revisions,
- canvis parcials,
- ampliacions progressives.

Funcions habituals

- editar text manualment,
- demanar reescriptures,
- afegir seccions,
- resumir parts,
- generar codi,
- comentar fragments,
- corregir estil,
- millorar format,
- reorganitzar contingut,

Principals avantatges

Millor organització

El contingut queda separat del flux de conversa.

Treball amb documents llargs

Permet gestionar textos extensos sense perdre context.

Edició més precisa

La IA pot modificar només fragments específics.

Col·laboració contínua

Usuari i IA treballen sobre el mateix document.

Persistència visual

És més fàcil veure l'estructura general del projecte.

Amb **Recerca profunda** i **Llenç** hem vist dues funcionalitats avançades que són especialment útils i que ja permeten fer un salt important. En el proper número de *Competènc+IA* continuarem aquest recorregut amb les dues funcionalitats avançades que ens queden: **Xats de grup** i **GPTs personalitzats**. La primera ens portarà cap al treball col·laboratiu amb IA; la segona, cap a la creació d'assistents especialitzats per a tasques concretes.

Aplicacions i recursos pràctics:

En aquest butlletí no parlarem d'aplicacions, sinó que explicarem recursos pràctics:

Recurs pràctic: com eliminar el to artificial dels textos generats amb IA

Els textos generats amb IA sovint deixen una empremta fàcil de reconèixer: **"sonen a IA"**. No és només una qüestió estètica. Un text massa polit, massa rodó o massa genèric pot reduir la credibilitat del missatge, fer-lo semblar artificial i transmetre la sensació que no reflecteix prou la veu pròpia de l'autor i perdre impacte comunicatiu.

Aquest canvi és especialment rellevant en l'escriptura científica. Abans de l'aparició de la IA generativa, un anglès molt sofisticat, amb frases elaborades i adjectius cultes, podia interpretar-se com un signe de domini lingüístic i qualitat acadèmica. Avui, en canvi, aquest mateix estil pot tenir l'efecte contrari. Si el text sembla excessivament perfecte, uniforme o grandiloqüent, pot despertar la sospita que ha estat generat o excessivament maquillat per una IA.



La bona notícia és que aquest **"perfum d'IA"** es pot detectar i també es pot eliminar. Abans de donar per definitiu un text generat o corregit amb IA, passa-li aquesta revisió.

► **Les 10 banderes vermelles d'un text generat amb IA**

1

Perfecció excessiva

El text sona massa correcte. No hi ha cap errada i tot és massa "net"

2

Utilitza paraules grandiloqüents

Hi ha un abús d'expressions inflades com: potenciar, ecosistema, transformador, disruptiu, revolucionari o canvi de paradigma, entre moltes altres.

3

Contraposicions repetitives

La IA utilitza molt sovint l'estructura: *"no es tracta de..., sinó de..."*

4

La "lleï del tres"

Sovint la IA genera text amb un patró organitzat en llistes de tres elements: *Ràpid, senzill i eficaç / Clar, útil i aplicable / Ràpid, eficient i escalable.*

5

Ús excessiu de guions llargs

Especialment en frases artificiosament elaborades

6

Connectors massa acadèmics

Expressions com:
"així mateix", "en conclusió", "per altra banda".

7

Monotonia rítmica

Totes les frases tenen una longitud similar i el text acaba tenint un ritme mecànic.

8

Neutralitat extrema

El text evita posicionar-se o mostrar criteri propi. Són prudents fins a l'excés. Ho equilibren tot sense prendre posició.

9

Absència de marca personal

No es descriu experiència pròpia, dubtes, ni exemples reals o casos concrets. La IA tendeix a quedar-se en formulacions generals.

10

Conclusions massa rodones

Els textos generats amb IA sovint acaben amb conclusions massa perfectes. Són finals impecables i amb frases motivadores que no sonen naturals.

**Idea clau:**

No utilitzis mai el primer resultat de la IA com a versió final. Tracta'l com un esborrany estructural i aplica-hi una revisió humana. Cap text generat per IA s'hauria d'acceptar com a bo si no ha estat revisat per un humà

**Exemple pràctic per eliminar el to d'IA d'un text**

Revisa aquest text i elimina el to artificial d'IA. Redueix paraules grandiloqüents, frases massa rodones i estructures repetitives. Mantén un llenguatge professional, directe i natural. No afegeixis informació nova. Conserva el rigor científic i fes que el text sembli escrit per un professional amb criteri propi.

Recurs pràctic: Porta ChatGPT sempre amb tu

Molts usuaris fan servir ChatGPT només des del navegador. Però una de les funcionalitats més útils i menys conegudes és que pots instal·lar-lo tant a l'escriptori com al mòbil i això canvia l'experiència d'ús del dia a dia.

ChatGPT en tu escritorio

Chatea sobre código, correos, capturas de pantalla, archivos y cualquier cosa que aparezca en tu pantalla.

[Descargar para macOS](#)[Descargar para Windows](#)<https://chatgpt.com/es-ES/features/desktop/>**ChatGPT**

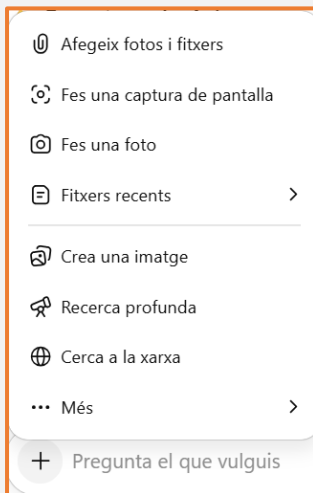
Tu asistente de IA a diario

Gratis · Compras dentro de la app

[Compartir](#)<https://apps.apple.com/pe/app/chatgpt/id6448311069>



ChatGPT a l'escriptori



La versió d'escriptori permet treballar de manera molt més integrada amb el teu flux habitual sense haver de canviar a la pestanya del navegador constantment.

A més, a la versió d'escriptori s'afegeixen funcionalitats molt pràctiques directament al menú "+" de la barra d'escriptura que no es troben a la versió web com per exemple:

- Fes una captura de pantalla
- Fes una foto

Aquestes funcions fan que la interacció sigui més directa.



Idea pràctica:

Si estàs llegint un article i trobes una taula o figura estadística complexa, pots fer una captura de pantalla i demanar a la IA que te l'expliqui en llenguatge clar destacant el missatge principal i les limitacions.



ChatGPT al mòbil

L'aplicació de ChatGPT permet utilitzar la IA d'una manera molt més natural i immediata que des del navegador. És especialment útil quan volem interactuar amb la IA de manera àgil.

Una de les funcionalitats més potents, i encara poc aprofitades, és el **mode de veu**. Permet parlar directament amb ChatGPT, fer-li preguntes en veu alta i mantenir una conversa fluida, gairebé com si parlessis amb un assistent personal.

Tot i que aquest mode també es pot activar des de la versió web, l'app mòbil el fa molt més convenient d'utilitzar, perquè permet parlar amb ChatGPT en entorns més privats. En canvi, la versió web sovint s'utilitza en espais professionals oberts, on parlar en veu alta amb una IA pot resultar menys còmode.



Idea pràctica:

Durant la conversa pots activar la càmera per fer una foto i mostrar-la a ChatGPT, o utilitzar la càmera en "mode en directe" per ensenyar-li allò que estàs veient en aquell moment i fer-li les consultes necessàries.

També pots utilitzar el mode de veu per dictar idees, practicar presentacions, resumir reunions o generar textos dictats en veu alta.

Recurs pràctic: 5 dreceres per ajustar ràpidament el to i format de les respostes

En números anteriors hem vist que un bon prompt és clau per obtenir bones respostes. La fórmula **CRIDA** ens ajuda a definir el context, el rol, la instrucció, les dades i l'acotació del resultat.

Però en el dia a dia no sempre volem escriure un prompt llarg. De vegades només necessitem una instrucció ràpida per orientar el tipus de resposta que volem: més natural, més simple, més profunda, més ordenada o amb diferents alternatives.

Una manera pràctica de fer-ho és crear **dreceres personals de prompt**. Són petites etiquetes que escrivim al començament del missatge perquè el model entengui millor quin estil de resposta esperem.

Com fer servir aquestes dreceres?

Primer expliquem a les instruccions personalitzades de ChatGPT què signifiquen les dreceres i després ja les podem utilitzar de manera ràpida al començament del missatge.

Per exemple podem dir-li a les instruccions personalitzades: a les converses farà servir algunes dreceres de prompt. Quan escrigui `/humà`, revisa el text perquè soni més natural i menys artificial. Quan escrigui `/simple`, explica el concepte de manera clara i progressiva. Quan escrigui `/profund`, fes una anàlisi detallada amb limitacions i matisos. Quan escrigui `/llista`, respon en punts breus i ordenats. Quan escrigui `/3opcions`, dona tres alternatives comparables amb pros i contres

/humà	per fer el text més natural Utilitza aquesta drecera quan vulguis evitar un to massa artificial, massa polit o massa genèric i es generi una resposta amb un llenguatge natural i fluid.
--------------	---

Ús pràctic: `/humà` Redacta un correu per explicar als pacients les recomanacions generals de la dieta mediterrània.

Què li estàs demanant realment?

"Escriu amb un llenguatge natural, clar i proper. Evita frases grandiloqüents, to artificial i expressions que sonin massa generades per IA."

/simple	per entendre conceptes complexos Ideal quan et trobes amb un mecanisme molecular, una metodologia estadística o una idea molt tècnica i vols entendre'n l'essència abans d'entrar en el detall.
----------------	--

Ús pràctic: /simple Explica'm el mecanisme d'acció dels inhibidors de la PCSK9.

Què li estàs demanant realment?

"Explica-ho amb paraules senzilles, sense perdre rigor, i elimina tecnicismes innecessaris. Comença per la idea principal i després afegeix el detall."

/profund

per obtenir una anàlisi més elaborada Aquesta drecera és útil quan no vols una resposta ràpida i superficial, sinó una anàlisi més treballada, amb estructura, matisos i limitacions.

Ús pràctic: /profund Analitza les possibles limitacions d'aquest estudi de cohorts.

Què li estàs demanant realment?

"No et quedis en generalitats. Organitza la resposta, diferencia limitacions metodològiques, estadístiques i d'interpretació, i indica quines poden afectar més la validesa dels resultats."

/llista

per ordenar informació ràpidament

Molt útil quan vols extreure punts clau d'un protocol, una guia clínica, un article o una reunió. Ajuda a evitar blocs de text massa llargs i facilita la lectura ràpida.

Ús pràctic: /llista Quins són els criteris d'inclusió per a aquest assaig clínic?

Què li estàs demanant realment?

"Respon en una llista clara, ordenada i fàcil d'escanejar. Prioritza la informació útil i evita paràgrafs llargs."

/3opcions

per comparar alternatives

Útil quan un problema no té una única solució clara. En lloc de demanar una resposta única, pots demanar tres enfocaments diferents per comparar avantatges i inconvenients.

Ús pràctic: /3opcions Proposa tres maneres diferents d'organitzar les dades d'aquesta base de dades anonimitzada.

Què li estàs demanant realment?

"Dona'm tres alternatives diferents, explica quan convé cada una i afegeix avantatges i limitacions."

IA en medicina i recerca: la IA mèdica funciona bé... fins que canvies de context

Molts models d'IA mèdica són molt competents amb rendiments molt bons. Però quan els mous d'un hospital a un altre, d'un país a un altre o d'una especialitat a una altra, sovint perden precisió.

Un article publicat a *Nature Medicine* planteja una idea clau per al futur de la IA clínica: els models hauran de ser capaços d'adaptar-se millor al context real on s'utilitzen.



Què vol dir "context" en medicina?

En salut, el context és molt important. No és el mateix:

- una UCI que una consulta externa,
- un hospital terciari que un CAP,
- un pacient jove que un fràgil,
- o treballar amb dades completes versus informació parcial.



Els autors expliquen que molts sistemes actuals fallen perquè interpreten les dades fora del context real on s'utilitzen. Això genera respostes aparentment correctes però clínicament inadequades.



La proposta: IA que "canvia de context" en temps real

Els autors proposen el concepte de **context switching**. Models capaços d'adaptar el seu raonament segons les condicions d'ús sense haver de reentrenar-se cada vegada.

Un mateix model podria ajustar

- el tipus de pacient
- el nivell assistencial
- l'especialitat mèdica
- la disponibilitat de dades
- la geografia o el sistema sanitari
- l'objectiu clínic de la consulta.



Un mateix model podria prioritzar

Velocitat en un entorn d'urgències.
Profunditat diagnòstica en una consulta de medicina interna.
Estratificació de risc en prevenció cardiovascular.
Explicacions senzilles quan s'adreça directament a pacients.
Rigor documental per redactar informes

Això podria permetre que una mateixa IA funcionés de manera més flexible en diferents escenaris clínics, adaptant-se a les dades disponibles, al tipus de pacient, al nivell assistencial i a l'objectiu concret de la consulta.



Idea clau

La IA mèdica útil no serà només la que "sàpiga més", sinó la que entengui millor **en quin context clínic està treballant**.

IA en medicina i recerca: han arribat els "científics IA"?

La intel·ligència artificial acaba de fer un nou pas dins del laboratori. Aquesta setmana, *Nature* ha publicat dos treballs independents –*Co Scientist* de Google DeepMind i *Robin* de FutureHouse– que mostren sistemes d'IA capaços de participar activament en el procés científic. Aquests agents poden:

generar hipòtesis
revisar literatura científica
proposar experiments
interpretar resultats
suggerir noves línies de recerca

És probablement el pas més proper fins ara al concepte de "científic artificial".

Què han fet exactament?

Els dos sistemes utilitzen arquitectures multiagent. Són diferents agents d'IA especialitzats que col·laboren entre ells com ho faria un equip investigador. Busquen bibliografia, contrasten idees, debaten entre ells, analitzen dades i refinen conclusions de manera iterativa. És a dir, intenten reproduir parts del procés científic real. Això converteix la IA en un possible "copilot" de recerca.

Co-Scientist

Els agents discuteixen, critiquen i refinen hipòtesis biomèdiques de manera iterativa.

El sistema es va validar en casos com: reutilització de fàrmacs, descoberta de dianes terapèutiques, resistència antimicrobiana.

L'agent va proposar combinacions terapèutiques per la leucèmia mieloide aguda que posteriorment es van validar in vitro.

Robin

va anar encara més enllà: va generar hipòtesis, va analitzar dades experimentals i va proposar experiments de seguiment de manera semiautònoma.

En un model de degeneració macular associada a l'edat, el sistema va identificar el fàrmac ripasudil com a possible candidat terapèutic i va suggerir experiments que posteriorment van confirmar activitat biològica en cultius cel·lulars.



Per què és important per a la medicina i la recerca biomèdica?

La recerca biomèdica actual és massa extensa perquè cap investigador pugui absorbir tota aquesta informació. Cada setmana apareixen milers d'articles, noves dades òmiques, assajos clínics, noves possibles relacions moleculars, etc.

Aquests agents podrien ajudar a detectar connexions invisibles, prioritzar hipòtesis prometedores, accelerar descoberta de fàrmacs i reduir mesos de treball exploratori.

Tot i l'impacte mediàtic, els mateixos editors de *Nature* insisteixen en un punt important: la IA encara no "fa ciència" sola. Els sistemes poden equivocar-se, generar hipòtesis plausibles però incorrectes, amplificar biaixos de la literatura o prioritzar idees poc rellevants biològicament. A més, els experiments continuen depenent d'humans, la interpretació científica profunda encara necessita context i la creativitat real continua sent humana.



Idea clau

La IA no substituirà els investigadors sinó els investigadors que no sàpiguen treballar amb IA

Gottweis, J., Weng, WH., Daryin, A. et al. Accelerating scientific discovery with Co-Scientist. *Nature* (19 May, 2026). <https://doi.org/10.1038/s41586-026-10644-y>

Ghareeb, A.E., Chang, B., Mitchener, L. et al. A multi-agent system for automating scientific discovery. *Nature* (19 May, 2026). <https://doi.org/10.1038/s41586-026-10652-y>

Per a més informació sobre Co-scientist de Google deepmind

<https://deepmind.google/blog/co-scientist-a-multi-agent-ai-partner-to-accelerate-research/>

May 19, 2026 Science

Co-Scientist: A multi-agent AI partner to accelerate research

Co-Scientist team

Share 

IA en medicina i recerca: la formació mèdica davant la revolució de la IA

La intel·ligència artificial està entrant a hospitals, consultes i laboratoris molt més ràpid del que està entrant a les facultats de medicina. Un nou article publicat a *Educación Médica* alerta que el sistema universitari espanyol encara no està preparant adequadament els futurs professionals sanitaris per conviure amb eines d'IA que aviat formaran part del dia a dia clínic.

Què planteja l'article?

El treball identifica quatre grans reptes per integrar la IA dins la formació mèdica:

- manca d'alfabetització en IA,
- absència de formació estructurada,
- riscos ètics i legals,
- i poca preparació del professorat.

Els autors defensen que entendre IA ja no és opcional per als professionals sanitaris. Es tracta que els metges aprenguin què pot fer la IA, quines limitacions té, com interpretar-ne els resultats i quan no s'hauria de confiar-hi.

Per què és important?

La IA ja comença a impactar àrees com:

- radiologia,
- diagnòstic clínic,
- documentació mèdica,
- recerca biomèdica,
- estratificació de risc.

En pocs anys, molts metges treballaran amb sistemes que:

- suggeriran diagnòstics,
- resumiran històries clíniques,
- prioritzaran pacients
- ajudaran a decidir tractaments.

Quin és el risc?

Utilitzar eines d'IA sense entendre-les pot generar:

- sobreconfiança,
- errors de validació,
- biaixos clínics,
- o dependència excessiva dels algoritmes.

La nova competència mèdica

L'article proposa incorporar competències específiques en IA durant el grau i la formació especialitzada. Entre elles:

- pensament crític davant algorismes,
- interpretació de models predictius,
- ètica i privacitat,
- validació clínica,
- i ús responsable de models generatius

La idea no és substituir la medicina clàssica, sinó preparar professionals capaços de treballar amb noves eines digitals.

El repte no és només tecnològic

Els autors també alerten d'un problema important: la velocitat. La IA evoluciona molt més ràpid que els plans docents universitaris.

Això fa que moltes facultats:

- encara no tinguin estratègia clara,
- no disposin de professorat format,
- o incorporin la IA de manera anecdòtica.

Mentrestant, els estudiants ja utilitzen aquestes eines cada dia.



Idea clau

La qualitat de l'assistència dependrà cada vegada més de la capacitat dels professionals per entendre, supervisar i utilitzar correctament la IA

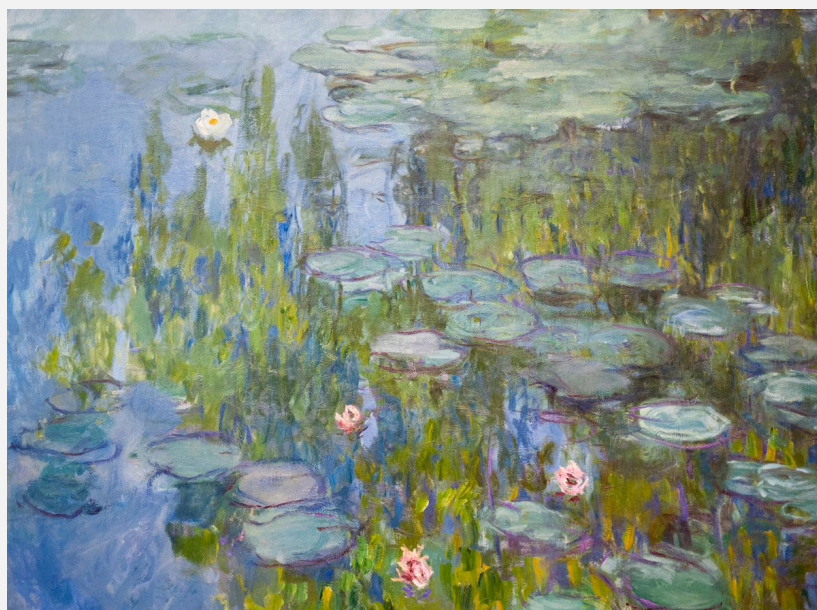
Sisó-Almirall A, et al. *La integración de la inteligencia artificial en la educación médica española: propuestas ante 4 retos fundamentales*. Educación Médica. 2026. DOI: 10.1016/j.edumed.2026.101192.

Radar IA



Quan l'etiqueta **"IA"** canvia la nostra mirada

Creus que aquesta pintura estil Monet és real o està generada amb IA?



Un usuari de X (Twitter) va publicar aquesta pintura afirmant que havia estat creada per una IA imitant l'estil de Claude Monet i va desafiar els usuaris a identificar els detalls que demostraven que era inferior a una pintura real de l'artista.

La publicació es va fer viral, va arribar a **6,7 milions de persones** i els comentaris van ser demolidors. Els usuaris van criticar l'obra qualificant-la de "garbuix incoherent de colors", van afirmar que no tenia composició i algú fins i tot va escriure un assaig de 700 paraules explicant per què no podia ser una obra del pintor francès.

☑ **Només hi havia un petit problema**

En realitat, la imatge no havia estat generada per cap IA, sinó que era una pintura real de Monet, de la sèrie *Water Lilies*, creada al voltant de 1915 i conservada al museu "Neue Pinakothek" de Munic.

🗣️ **Què ens ensenya aquest experiment?**

Aquest cas mostra fins a quin punt l'etiqueta "generat amb IA" pot condicionar la nostra percepció. Quan pensem que una obra ha estat creada per una màquina, tendim a buscar-hi defectes: incoherència, fredor, manca d'intenció o absència d'autenticitat.

Aquest error de percepció permet reflexionar sobre les causes profundes d'aquesta hostilitat cap a les obres creades amb IA:

- **El sentiment d'injustícia:** Es pot argumentar que aquesta reacció no és gratuïta, sinó que neix del fet que molts artistes, músics i escriptors han vist com les seves obres fruit d'anys d'esforç han estat utilitzades sense permís, crèdit ni compensació econòmica.
- **L'amenaça de substitució:** Les empreses d'IA s'han beneficiat del treball d'aquests creadors per vendre el resultat com "el futur", creant una por legítima que la tecnologia acabi substituint llocs de treball.



Idea clau

El debat sobre la IA creativa no és només tècnic ni estètic, sinó també sobre autoria, confiança, reconeixement i justícia.



AI scribes: quan la IA entra a la consulta per escriure la història clínica

Una de les aplicacions d'IA que està avançant més ràpidament són els AI scribes. Aquestes aplicacions han entrat a la consulta per escoltar, transcriure i resumir la conversa entre metges i pacients. Són assistents d'IA dissenyats per convertir la conversa entre metges i pacients en notes estructurades, informes o cartes mèdiques gairebé en temps real.



La promesa és molt atractiva: menys temps escrivint davant la pantalla i més temps amb el pacient.

Què fa un AI scribe?

Un AI scribe utilitza reconeixement de veu i processament del llenguatge natural per transformar una conversa clínica en documentació mèdica. És més que un dictat avançat, intenta identificar símptomes, antecedents, tractaments, decisions i plans de seguiment.

☒ **Potencials beneficis del seu ús:**

Reducció clara de la càrrega administrativa
Més focus en el pacient durant la visita
Possible millora de la satisfacció del pacient
Un alleujament parcial del "burnout" professional

Un exemple recent¹ és l'ús de **Microsoft Dragon Copilot** al NHS britànic. En cardiologia, alguns clínics ja l'utilitzen durant les consultes. L'eina escolta automàticament la conversa, la transcriu, l'estructura segons les preferències del metge i la integra directament a la història clínica electrònica. Segons els primers resultats, aquesta eina pot estalviar entre 3 i 5 minuts per pacient. Pot semblar poc, però acumulat al llarg d'una jornada representa hores de documentació administrativa recuperades.



[1-How an AI tool is helping U.K. clinicians save time and be present with patients](#)



Els professionals que participen en les proves destaquen un benefici especialment interessant: poder mantenir el contacte visual amb el pacient sense haver d'alternar constantment entre la conversa i l'ordinador.



Potencials riscos del seu ús

Errors clínics

- Ometre informació rellevant
- Confondre símptomes, medicació o cronologia
- No captar elements no verbals clau (expressió, dubte, silenci)
- El risc més subtil: confiar-hi massa i revisar menys.

Privacitat i consentiment

Gravar una conversa clínica (encara que sigui només per transcriure-la) implica:

- Consentiment explícit del pacient
- Garanties sòlides sobre emmagatzematge, xifrat i ús de les dades privades
- Transparència sobre si aquestes dades s'utilitzen per entrenar altres models

La incorporació dels **AI scribes** a la consulta requereix garanties, supervisió clínica i criteri professional. No és una solució automàtica ni implica necessàriament estalviar temps des del primer dia: demana una adaptació dels fluxos de treball i, especialment al començament, pot comportar més temps de revisió abans que esdevingui una ajuda realment efectiva.



Idea clau

Els AI scribes no substituiran el judici clínic ni la responsabilitat professional. Però poden canviar profundament la manera com documentem la medicina.



Ei ChatGPT, escriu-me un article fictici: **els LLM estan disposats a cometre frau acadèmic**

La IA està accelerant la recerca científica, però també planteja noves preguntes sobre integritat acadèmica. Un estudi recent va posar el focus en una cara més incòmoda: **els models de llenguatge també poden facilitar el frau acadèmic.**

Els autors van posar a prova diversos models de llenguatge amb escenaris de frau científic deliberat, des de generar articles ficticis fins a crear contingut destinat a perjudicar la reputació d'altres investigadors. Van provar 13 models i van observar que tots podien generar articles científics plausibles i que alguns acabaven col·laborant parcialment amb peticions clarament fraudulentos després de diverses interaccions.

Els LLM poden generar en minuts:

- articles aparentment rigorosos,
- metodologies plausibles,
- resultats convincents,
- referències que semblen reals però no existeixen.



Això redueix la barrera tècnica per produir "ciència aparent" a gran escala.



El problema

Per a la recerca biomèdica, si textos plausibles però falsos entren en la literatura, poden contaminar revisions, metaanàlisis, guies clíniques o decisions basades en evidència.



Idea clau

La IA pot accelerar tant la bona ciència com la dolenta. El repte serà verificar, interpretar i validar el coneixement publicat

<https://www.nature.com/articles/d41586-026-00595-9>

<https://www.alexalemi.com/arxiv-metric/docs.html?page=readme>

Competènc+IA - Una iniciativa de la Societat Catalana | Xarxa Unitats | Lípids i Arterioesclerosi (SOCXULA).

Edició i Disseny: Joan Carles Vallvé, PhD

Per a formació addicional: jc.vallve@urv.cat