

Competènc+IA

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

Butlletí d'Intel·ligència Artificial

Número 6 | Juny 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 BALEARS



Aquesta entitat dona suport als Objectius de Desenvolupament Sostenible

Amb la col·laboració de Daiichi-Sankyo



Daiichi-Sankyo



Editorial

De conversar amb la IA a treballar amb la IA

En els primers números de **Competènc+IA** hem construït una base per utilitzar la intel·ligència artificial amb criteri professional: entendre què és un model de llenguatge, escriure bons prompts, configurar ChatGPT, organitzar la feina amb Projectes, connectar-lo amb Apps i començar a utilitzar funcionalitats avançades.

En aquest número de juny fem un pas més. La IA deixa de ser només una eina que respon preguntes i comença a convertir-se en un entorn capaç de planificar, executar, editar, col·laborar i crear eines pròpies. Això canvia la manera d'utilitzar-la: ja no n'hi ha prou amb demanar una resposta, sinó que cal definir objectius, establir límits i supervisar processos.

Per això obrim el número parlant dels **agents d'IA**, sistemes capaços de dividir una tasca en passos, utilitzar eines externes i avançar cap a un objectiu. En recerca biomèdica poden ajudar a automatitzar tasques repetitives, preparar revisions bibliogràfiques o gestionar processos administratius. Però com més autonomia donem a la IA, més important és controlar què fa i amb quines dades treballa.

També abordem el **vibe coding**, una nova manera de crear aplicacions descrivint en llenguatge natural què volem construir. Aquesta aproximació pot permetre que investigadors sense perfil tècnic desenvolupin qüestionaris, eines de recollida de dades o assistents específics per als seus projectes. La IA pot generar codi, però no substitueix el criteri científic, metodològic, ètic ni de seguretat.

En l'apartat **Domina la IA** tanquem el recorregut sobre funcionalitats avançades de ChatGPT: el botó **Edita**, els **Xats de grup** i els **GPTs personalitzats**. Totes aquestes eines apunten en la mateixa direcció: passar d'una conversa improvisada a un flux de treball més organitzat, col·laboratiu i reutilitzable.

A la part pràctica veurem com crear **infografies** amb el mode imatge de ChatGPT i com millorar una imatge pixelada. I, en les seccions d'IA en medicina i recerca i Radar IA, revisarem alguns riscos actuals: l'ús de xatbots per consultes de salut, els errors en el triatge, el FOMO científic, les cites fantasma i les noves mesures d'arXiv contra referències inventades.

La IA avança cap a més autonomia, però la competència professional no consisteix a deixar-la actuar sola. Consisteix a saber quan utilitzar-la, com supervisar-la i on posar els límits.

Entenent la IA: els agents d'IA - quan la IA planifica i executa processos

Des de que va aparèixer la IA hem après a interactuar amb ella com si fos un xat. Fem una pregunta, obtenim una resposta. Però la propera evolució de la intel·ligència artificial consisteix en actuar i no en respondre millor.

Aquesta és la idea que hi ha darrere dels anomenats **agents d'IA**, sistemes que de forma autònoma són capaços de planificar, dividir una tasca en passos, utilitzar eines externes, prendre decisions intermèdies i adaptar-se als resultats obtinguts i executar tasques complexes per assolir un objectiu.

Del xat a l'acció

Model de llenguatge convencional

Usuari → pregunta → IA → resposta

"Escriu-me un correu per convocar una reunió"

Agent d'IA

Usuari → objectiu → IA → planificació → ús d'eines → execució → revisió del resultat

"Organitza una reunió amb aquest equip la setmana vinent."

Un **agent d'IA** pot revisar calendaris, detectar disponibilitats, proposar una data, redactar el correu, enviar les invitacions i preparar un resum dels temes a tractar.

Què necessita un agent d'IA per funcionar?

La majoria d'agents comparteixen quatre elements bàsics:

Un objectiu	L'usuari defineix què vol aconseguir
Raonament i planificació	L'agent divideix l'objectiu en passos i estableix una estratègia
Accés a eines	L'agent pot connectar-se a correus electrònics, calendaris, bases de dades, navegadors web, documents o aplicacions
Execució i revisió	L'agent realitza accions, verifica resultats i pot corregir errors o replantejar la seva estratègia



Per això sovint es parla de **capacitat agèntica**: la IA ja no és només una eina de conversa i generació de continguts, sinó que entra activament en l'automatització de processos.

Model de llenguatge convencional

Respon preguntes
Genera text
Depèn molt de cada prompt
No actua fora del xat
És reactiu

Agent d'IA

→ Persegueix objectius
→ Executa processos
→ Pot seguir una seqüència de passos
→ Pot connectar-se amb eines externes
→ Pot ser semiautònom

Què poden aportar a la recerca biomèdica?

L'interès pels agents és especialment gran en entorns científics i sanitaris, on una part important del temps es dedica a tasques repetitives i administratives.

Alguns exemples potencials:

Agents de revisió bibliogràfica

- Cerca d'articles
- Extracció automàtica de dades
- Comparació d'estudis
- Elaboració de taules resum

Agents de suport estadístic

- Interpretació preliminar de resultats
- Identificació de problemes metodològics
- Generació d'esborranys de resultats

Agents docents

- Creació de preguntes tipus test
- Generació de casos clínics
- Adaptació de materials a diferents nivells formatius

Agents administratius

- Gestió de correus
- Seguiment de projectes
- Organització de reunions i terminis

Més autonomia, més responsabilitat

La gran fortalesa dels agents és també el seu principal risc.

Mentre que un chatbot només genera una resposta, un agent pot actuar.

Això implica nous reptes:

- Pot prendre decisions basades en informació incorrecta.
- Pot executar accions inadequades.
- Pot tenir accés a dades sensibles.
- Pot generar una falsa sensació de fiabilitat.

En l'àmbit biomèdic, això és especialment rellevant quan entren en joc dades de pacients, informació confidencial o resultats científics encara no validats.



Idea clau

Si ChatGPT va representar el pas de la cerca d'informació a la conversa amb la informació, els agents d'IA representen el pas de la conversa a l'acció i haurem de decidir quines tasques estarem disposats a delegar-li.

Entenent la IA: què és el *vibe coding*? crea aplicacions sense saber programar

Durant anys, desenvolupar una aplicació requeria coneixements de programació, temps de desenvolupament, recursos tècnics i, sovint, el suport d'un equip informàtic especialitzat. Avui, la IA està canviant aquesta realitat. Noves eines permeten descriure en llenguatge natural què volem construir i obtenir, en poc temps, una aplicació funcional o una eina digital. Aquest fenomen es coneix com a *vibe coding*.

Què és el *vibe coding*?

El *vibe coding* és la capacitat de crear aplicacions o plataformes digitals sense experiència prèvia en programació i sense escriure codi manualment. En lloc de programar, l'usuari descriu en llenguatge natural què vol que faci l'aplicació, i la IA genera el codi necessari.



Dit d'una altra manera, la persona aporta la idea i la lògica del que necessita, mentre que la IA s'encarrega de la implementació tècnica.

Un exemple recent publicat a *Nature* ho il·lustra molt bé. Investigadors de la Universitat Baptista de Hong Kong, procedents de l'àmbit de les humanitats i amb coneixements limitats de programació, van crear en pocs dies una plataforma pròpia per estudiar com els estudiants escrivien poesia amb assistents d'IA¹.

Com funciona?

El *vibe coding* parteix d'una idea simple: en lloc d'escriure codi, es descriu en llenguatge natural què és necessari i una IA genera l'aplicació

Vibe coding	
l'usuari descriu: què ha de fer l'aplicació, quines funcionalitats necessita, quines dades ha de guardar, com ha de comportar-se	 La IA genera: el codi necessari, la interfície, les connexions necessàries, la infraestructura

En l'exemple de *Nature*, els investigadors van demanar a una eina d'IA que construís una plataforma amb diferents espais de conversa, registre automàtic de les interaccions, emmagatzematge de dades i control de diversos paràmetres dels models. En qüestió de dies disposaven d'una eina funcional per dur a terme el seu estudi.



Com es podria aplicar a la recerca?

Aquesta nova generació d'eines podria facilitar que investigadors biomèdics sense perfil tècnic desenvolupessin aplicacions específiques per:

- reclutar participants
- recollir dades
- crear qüestionaris intel·ligents
- seguiment de pacients
- analitzar resultats
- construir assistents específics per a projectes

Tasques que fins fa poc requerien suport tècnic especialitzat podrien quedar a l'abast d'equips molt més petits o investigadors individuals.



El codi és fàcil. El criteri no.

La tecnologia redueix enormement les barreres tècniques, però continua sent necessari entendre:

- què volem construir,
- com validar-ne el funcionament,
- com protegir les dades i la seguretat

La IA pot generar el codi, però no substitueix el coneixement científic, metodològic ni ètic necessari per desenvolupar una eina fiable.



Idea clau: El *vibe coding* està democratitzant el desenvolupament de programari. Per primera vegada, professionals sense formació tècnica poden convertir idees en aplicacions funcionals amb una rapidesa impensable fa només uns anys.

1- Wang S, Ruobin Y, Christie S. [We vibe-coded a custom AI poetry lab. Here's how you can, too.](#) Nature. 2026. DOI: 10.1038/d41586-026-01106-6

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



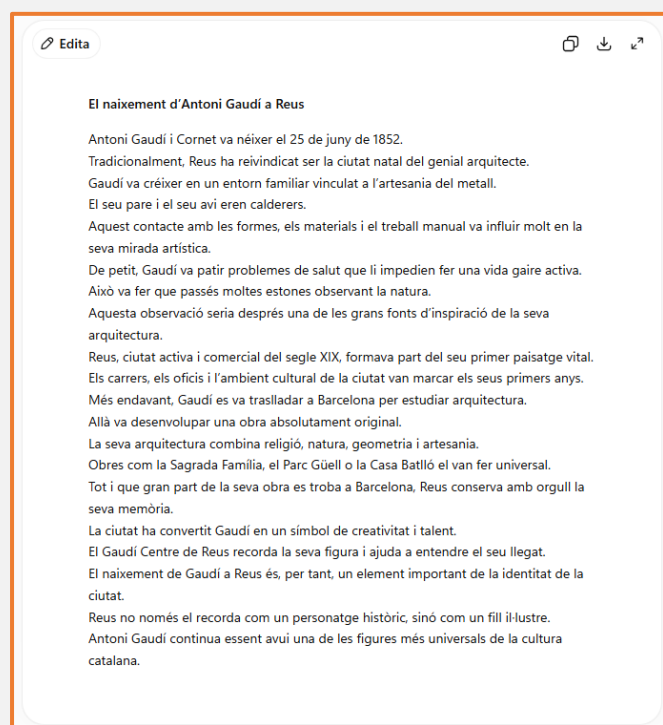
Del “Llenç” al botó “Edita”: la mateixa idea, integrada de forma més natural

En el número anterior vam parlar de la funcionalitat **Llenç** (*Canvas*), un espai de treball interactiu integrat dins de ChatGPT que permetia crear, editar i desenvolupar continguts llargs o complexos d'una manera més visual i estructurada.

La idea és molt útil: en lloc de treballar només amb missatges successius dins del xat, el Llenç obre una mena de document lateral on l'usuari i la IA poden redactar, revisar, modificar i reorganitzar un text de forma continuada. Això és especialment pràctic per treballar amb articles, correus, informes, textos docents, guions, newsletters o qualsevol document que requereix diverses revisions.

Recentment, aquesta funcionalitat ha canviat de forma dins de la interfície de ChatGPT. El **Llenç com a opció separada** ha desaparegut o ja no apareix de la mateixa manera en alguns comptes. Però la seva funcionalitat principal no s'ha perdut. Al contrari: ara s'ha integrat de manera més natural dins de les respostes que generen documents o textos editables.

A la pràctica, quan demanem a ChatGPT que redacti un correu, un article, una nota, un informe o qualsevol text llarg, la resposta apareix com un document editable.



Només cal clicar el botó **"Edita"** per obrir un espai de treball molt semblant al que abans anomenàvem Llenç. Des d'allà podem modificar el text directament dins de ChatGPT, seleccionar fragments concrets, demanar canvis parcials, reescriure seccions, ajustar el to, ampliar o escurçar continguts i continuar treballant sobre una mateixa versió del document.

Aquest canvi és important perquè fa que l'edició assistida per IA sigui menys una funcionalitat "extra" i més una part natural del flux de treball. Ja no cal pensar tant en activar una eina específica, sinó simplement en demanar el text que volem crear i, després, editar-lo directament quan aparegui l'opció.

Xats de grup: quan ChatGPT deixa de ser una eina individual

Fins ara hem vist ChatGPT sobretot com una eina individual: una persona fa una pregunta, ChatGPT respon, i el treball queda dins d'una conversa privada. Però moltes de les tasques professionals no es fan soles i és aquí on entren els **Xats de grup**. Aquesta funcionalitat permet afegir altres persones dins d'una mateixa conversa amb ChatGPT i conversar entre elles i, quan cal, demanar ajuda a ChatGPT transformant-lo en una mena de moderador, secretari, sintetitzador i assistent de treball compartit.

Què és un xat de grup?

Un xat de grup és una conversa compartida dins de ChatGPT on participen diverses persones i també ChatGPT com un membre més del grup. A diferència d'un xat individual, on només hi ha una interacció entre usuari i IA, en un xat de grup diferents participants poden: escriure missatges, compartir idees, discutir opcions, pujar documents o imatges, demanar a ChatGPT que resumeixi, ordeni o proposi alternatives, continuar la conversa de manera col·laborativa, etc. ChatGPT participa quan se li demana o quan detecta que pot aportar valor actuant com a suport.

Xat individual

- Una persona conversa amb ChatGPT
- Útil per treball personal
- El context el defineix un sol usuari
- ChatGPT respon directament a l'usuari
- Ideal per redactar, estudiar o analitzar individualment



Xat de grup

- Diverses persones conversen entre elles i amb ChatGPT
- Útil per treball col·laboratiu
- El context es construeix entre tots
- ChatGPT pot ajudar a ordenar, resumir o moderar el grup
- Ideal per coordinar, decidir, revisar o construir conjuntament

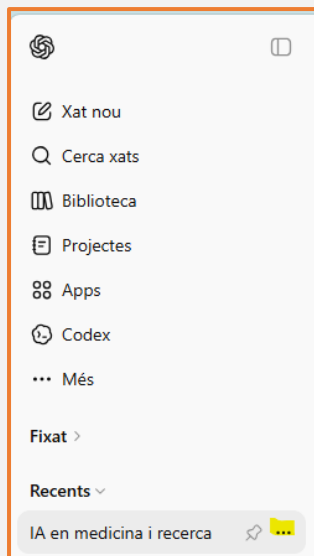


Idea clau

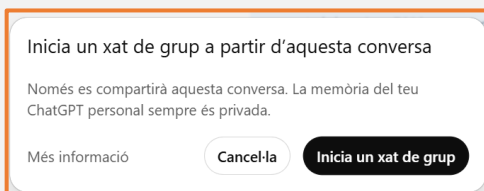
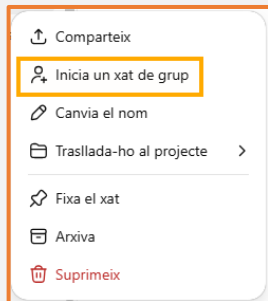
El xat de grup no és només "afegir gent" a una conversa. És convertir ChatGPT en un espai compartit de treball on la IA ajuda el grup a pensar millor, ordenar idees i avançar amb més claredat.

Com es crea un xat de grup?

La interfície pot variar lleugerament segons el dispositiu, el pla i les actualitzacions de ChatGPT, però la lògica general és senzilla.

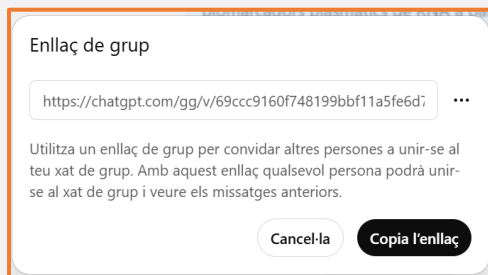


Normalment, es pot iniciar un xat de grup des d'un xat nou o des d'una conversa existent clicant als tres punts, obrint el menú i seleccionant l'opció "inicia un xat de grup".



En clicar "Inicia un xat de grup", apareix primer una finestra de confirmació que indica un aspecte important: Només es compartirà aquesta conversa, de manera que la resta del teu ChatGPT personal (inclosa la memòria) es manté privat. Això garanteix que pots compartir contingut de manera controlada, sense exposar altres xats o informació.

A més, si es converteix una conversa existent en un xat de grup, ChatGPT crea una còpia separada com a nova conversa de grup. El xat original continua sent privat i separat.



Un cop confirmes l'acció, s'obre un segon menú on es genera un enllaç de grup.

Aquest enllaç és la clau del funcionament ja que el pots copiar i enviar a altres persones. Qualsevol persona amb aquest enllaç pot accedir a l'espai compartit i a partir d'aquí, els participants poden continuar la conversa conjuntament.

Com participa ChatGPT dins del grup?

ChatGPT no ha de respondre a cada missatge. En un xat de grup, pot mantenir-se en segon pla i intervenir quan algú li demana ajuda o quan el context ho fa necessari.

A la pràctica, és recomanable cridar-lo explícitament quan volem una resposta clara. Això evita que la IA interrompi excessivament i ajuda a mantenir la conversa centrada en les persones.

ChatGPT, resumeix aquesta discussió.
ChatGPT, proposa una estructura.
ChatGPT, fes una taula comparativa.
ChatGPT, detecta punts d'acord i desacord.
ChatGPT, transforma aquesta conversa en una llista de tasques.



Consell pràctic:

Defineix el rol de ChatGPT al començament, per exemple: ChatGPT, en aquest xat actuaràs com a editor científic. Ajuda'ns a revisar textos, detectar incoherències, millorar la claredat i mantenir un to acadèmic. No afegeixis informació nova si no t'ho demanem explícitament.

GPTs personalitzats: la creació d'assistents especialitzats per a tasques concretes

Quan pensem en eines com ChatGPT, sovint imaginem un únic assistent capaç de respondre preguntes sobre qualsevol tema. Però la realitat és que avui podem disposar de versions especialitzades d'aquests models, dissenyades per a tasques concretes. Són els anomenats **GPTs personalitzats**.

Aquests assistents funcionen amb la mateixa tecnologia que ChatGPT, però han estat configurats amb instruccions específiques, coneixement especialitzat o eines addicionals per actuar com a experts en un àmbit determinat.



Què és un GPT personalitzat?

Un GPT personalitzat és una versió especialitzada d'un model de llenguatge que ha estat entrenada o configurada per ajudar en una tasca concreta.

Mentre que ChatGPT és un assistent generalista, un GPT personalitzat està orientat a un objectiu específic: revisar articles científics, ajudar amb anàlisi estadística, generar prompts, redactar documents o donar suport a l'aprenentatge.

La idea és senzilla: en lloc d'explicar cada vegada què necessitem, el GPT ja incorpora unes instruccions que defineixen com ha de respondre i quines funcions ha de prioritzar.

Això vol dir que cada vegada que l'utilitzem, el GPT ja "sap" quin paper ha de fer, quin tipus de resposta volem i quins límits ha de respectar.



Idea clau

Un GPT personalitzat és com deixar preparat un "prompt permanent" molt avançat. No substitueix el criteri professional, però redueix repeticions, dona coherència a les respostes i facilita treballar sempre amb el mateix marc.

🔍 Com funcionen?

Els GPTs personalitzats incorporen:

- instruccions específiques sobre com respondre
- coneixement especialitzat sobre una disciplina
- documents de referència,
- connexió amb eines externes,
- fluxos de treball adaptats a una necessitat concreta



Les respostes solen ser més enfocades i útils per a una tasca determinada.

📁 Alguns exemples útils per a investigadors

Scholar GPT

Està orientat a la cerca i interpretació de literatura científica. Pot ajudar a localitzar articles, resumir evidència disponible, identificar limitacions metodològiques o generar explicacions sobre conceptes biomèdics.

Statistics Stats

Especialitzat en estadística i anàlisi de dades. Pot ajudar a interpretar resultats, explicar proves estadístiques, suggerir anàlisis adequades o revisar la coherència d'un plantejament metodològic. Pot actuar com un suport ràpid per resoldre dubtes habituals.

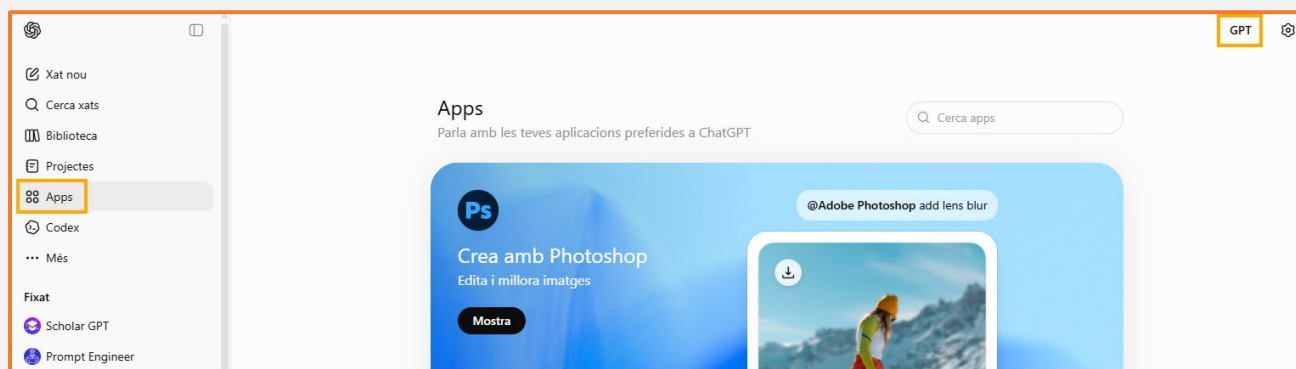
Prompt Engineer

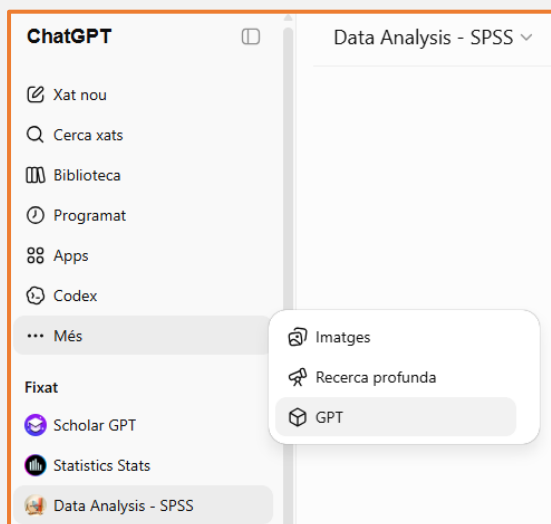
Ajuda a construir instruccions eficaces per interactuar amb altres models d'IA. Pot reformular preguntes, estructurar tasques complexes i generar prompts més precisos per obtenir millors resultats. És especialment útil quan volem utilitzar la IA de manera més eficient en recerca, docència o gestió.

🚀 Com s'activen?

Els GPTs personalitzats estan disponibles dins de ChatGPT. La forma d'accedir-hi pot variar lleugerament segons la versió de la plataforma, el tipus de compte (gratuït o de pagament) i les actualitzacions que OpenAI introdueixi periòdicament.

Actualment, **en la versió gratuïta**, es pot accedir als GPTs des del menú lateral esquerre seleccionant l'opció «**Apps**» i, posteriorment, accedint a l'apartat «**GPTs**», on es poden explorar i activar assistents especialitzats per a diferents tasques (Figura).





En les **versions de pagament** de ChatGPT, l'accés als GPTs personalitzats es pot fer directament des del menú lateral. Per fer-ho, cal seleccionar l'opció **«Més» (...)** i, a continuació, clicar **«GPT»**.



Un cop seleccionada l'opció **GPTs**, tant a la versió gratuïta com de pagament, s'obre el catàleg de GPTs personalitzats disponibles a ChatGPT. Aquest espai permet cercar els GPTs especialitzats mitjançant el cercador o explorar-

los per categories temàtiques, com ara educació, escriptura, productivitat, programació o recerca i anàlisi, entre d'altres.

Els GPTs més utilitzats es poden fixar al menú lateral per facilitar-ne l'accés ràpid, evitant haver de cercar-los cada vegada.

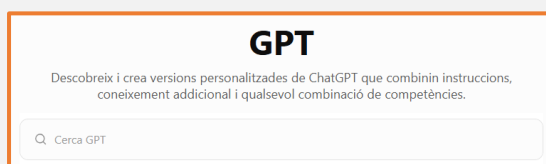


Consell pràctic: Si utilitzes habitualment GPTs especialitzats com *Scholar GPT*, *Statistics Stats*, *Data Analysis-SPSS* o altres, els pots fixar al menú lateral per accedir-hi directament amb un sol clic, evitant haver de cercar-los cada vegada.

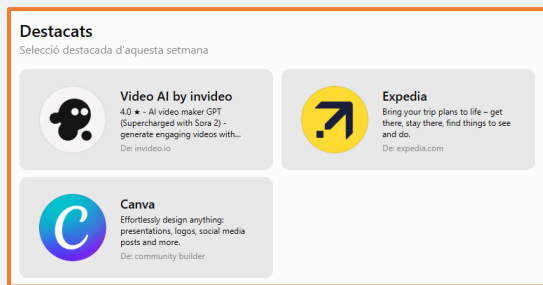
Explorar GPTs personalitzats

La pantalla d'exploració de GPTs funciona de manera similar a una botiga d'aplicacions. A més del cercador, ChatGPT mostra diferents GPTs destacats i recomanats, així com assistents que estan guanyant popularitat entre els usuaris.

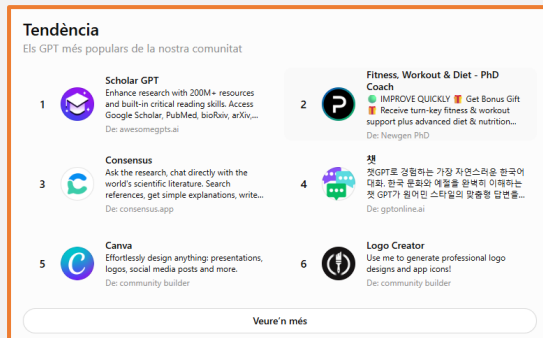
Els GPTs es poden localitzar de diverses maneres:



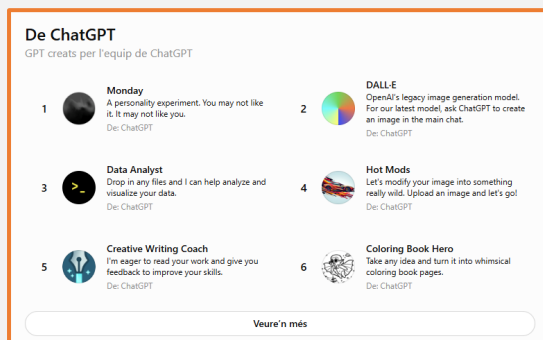
Cercador: facilita trobar GPTs concrets introduint paraules clau relacionades amb la tasca que es vol realitzar.



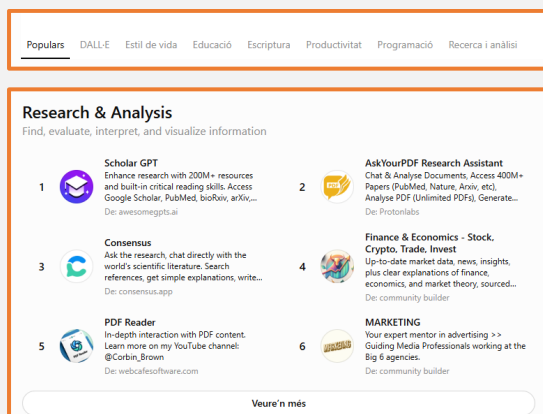
GPTs destacats: el menú mostra una selecció de GPTs destacats per la plataforma per la seva utilitat o qualitat.



GPTs en tendència: el menú mostra una selecció de GPTs que són els més utilitzats o populars en aquell moment.



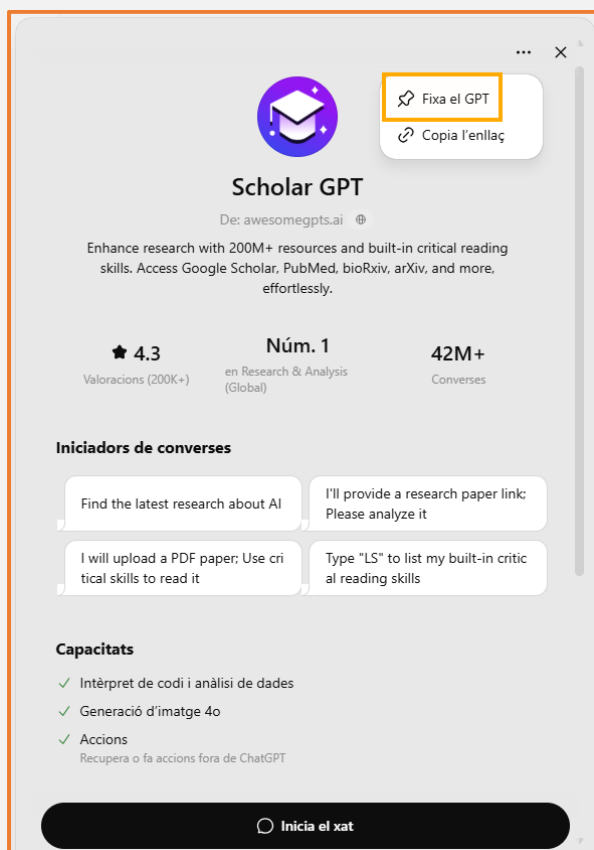
De ChatGPT: el menú mostra una selecció de GPTs oficials desenvolupats directament per l'equip d'OpenAI, que acostumen a servir de referència per entendre les possibilitats d'aquesta tecnologia.



Quan l'usuari selecciona una categoria, el sistema mostra una llista de GPTs relacionats amb aquella temàtica, incloent-hi una descripció breu de les seves funcionalitats. Això permet identificar ràpidament quins assistents poden ser útils per a cada necessitat.

🔗 Selecció i configuració d'un GPT personalitzat

Quan seleccionem un GPT personalitzat, s'obre una fitxa informativa (figura) amb una descripció de les seves funcionalitats, les tasques per a les quals ha estat dissenyat i alguns exemples de converses que poden ajudar a començar a utilitzar-lo.



A la fitxa també es mostren dades com la categoria a la qual pertany, la valoració dels usuaris, el nombre de converses mantingudes i les capacitats disponibles.

En l'exemple de la figura, **Scholar GPT** està especialitzat en suport a la recerca científica i facilita l'accés a fonts acadèmiques, la lectura crítica d'articles i la cerca bibliogràfica.

Per començar a utilitzar-lo només cal prémer el botó **«Inicia el xat»**, que obrirà una nova conversa utilitzant aquest GPT especialitzat.

Fixar un GPT al menú lateral

Si es preveu utilitzar un GPT de manera freqüent, és possible afegir-lo al menú lateral per accedir-hi més ràpidament. Per fer-ho, cal clicar el menú de **tres punts (...)** situat a la part superior de la fitxa del GPT i seleccionar l'opció **«Fixa el GPT»** (figura anterior). A partir d'aquest moment, el GPT apareixerà a la barra lateral de ChatGPT evitant haver de cercar-lo novament dins del catàleg.



Idea clau:

Els GPTs personalitzats representen una evolució de la IA generalista: permeten disposar d'assistents especialitzats per a tasques concretes sense necessitat de coneixements tècnics avançats. Ben utilitzats, poden augmentar la productivitat i facilitar moltes activitats de recerca, docència i gestió, però continuen requerint supervisió humana i pensament crític.

Aplicacions i recursos pràctics:

Crear infografies amb el mode imatge de ChatGPT

La capacitat de sintetitzar informació i presentar-la visualment és una competència cada vegada més rellevant en l'àmbit educatiu i professional. Les infografies són una eina molt útil per comunicar conceptes complexos de manera clara, atractiva i fàcilment comprensible. Tradicionalment, la seva elaboració requeria coneixements de disseny gràfic o l'ús d'eines específiques, però la incorporació de la generació d'imatges amb intel·ligència artificial ha simplificat considerablement aquest procés.

El mode imatge de ChatGPT permet crear infografies a partir d'un text o conjunt de dades proporcionat per l'usuari. D'aquesta manera, el docent o l'estudiant pot centrar-se en el contingut i deixar que la IA proposi una representació visual adequada.

Com funciona?

El procés és senzill:

1. Es proporciona a ChatGPT la informació que es vol representar.
2. S'especifica que es vol generar una infografia.
3. Es descriu el públic objectiu, el nivell educatiu o l'estil visual desitjat.
4. ChatGPT genera una imatge que organitza la informació en forma d'infografia.

Per exemple, es podria introduir el següent text:

"Resumeix les principals causes de les malalties cardiovasculars: hipertensió arterial, tabaquisme, sedentarisme, obesitat i diabetis. Crea una infografia clara i visual dirigida a estudiants de secundària."



A partir d'aquesta informació, ChatGPT generarà una proposta visual amb títols, icones, blocs de contingut i una estructura jeràrquica que faciliti la comprensió dels conceptes.

Recomanacions per obtenir millors resultats

Per aconseguir infografies de qualitat és recomanable:

- Proporcionar informació estructurada i clara.
- Indicar el públic destinatari.
- Especificar l'estil visual (acadèmic, professional, divulgatiu, infantil, etc.).
- Demanar una distribució concreta dels elements si és necessari.
- Revisar sempre el resultat final per comprovar que el text, les dades i les relacions conceptuals siguin correctes.

Aplicacions educatives

L'ús d'infografies generades amb IA ofereix nombroses possibilitats didàctiques:

Suport a l'aprenentatge

Els docents poden transformar apunts extensos o continguts curriculars en recursos visuals que afavoreixin la comprensió i la retenció de la informació.

Activitats de síntesi

L'alumnat pot elaborar infografies com a producte final d'un projecte, demostrant la seva capacitat per seleccionar, organitzar i comunicar informació rellevant.

Divulgació científica

Les infografies permeten adaptar continguts científics complexos a diferents públics, facilitant la transferència del coneixement.

Materials per a presentacions

Es poden generar recursos visuals de suport per a exposicions orals, seminaris o activitats de difusió.

Exemple de prompt

Un exemple de prompt podria ser:

"Crea una infografia en català sobre els factors de risc cardiovascular. Utilitza un disseny modern i professional. Inclou una breu descripció de cada factor de risc, icones representatives i una secció final amb recomanacions de prevenció. Format vertical A4."



Idea clau:

La generació d'infografies amb el mode imatge de ChatGPT representa una oportunitat per potenciar la competència digital, la comunicació visual i la capacitat de síntesi. Aquesta eina permet crear materials atractius de manera ràpida i accessible, afavorint tant la tasca docent com la participació activa de l'alumnat en la construcció i difusió del coneixement.

Tot seguit es mostra una infografia creada amb el mode imatge de ChatGPT a partir del contingut explicat més amunt. L'objectiu és veure, de manera pràctica, com una informació textual sobre el procés de creació d'infografies es pot transformar en un recurs visual clar, estructurat i útil per a metges i investigadors del camp del risc cardiovascular.





CREAR INFOGRAFIES

AMB EL MODE IMATGE DE ChatGPT

Una eina poderosa per a metges i investigadors del risc cardiovascular

La capacitat de sintetitzar informació i presentar-la visualment és clau en educació, comunicació científica i pràctica clínica. Les infografies permeten comunicar conceptes complexos de manera clara i atractiva. Amb el mode imatge de ChatGPT, crear infografies és més fàcil que mai: tu aportes el contingut, la IA el transforma en una representació visual efectiva.



COM FUNCIONA?



EXEMPLE RÀPID

"Resumeix les principals causes de les malalties cardiovasculars: hipertensió arterial, tabaquisme, sedentarisme, obesitat i diabetis. Crea una infografia clara i visual dirigida a estudiants de secundària."



PRINCIPALS CAUSES DE MALALTIES CARDIOVASCULARS



HIPERTENSIÓ ARTERIAL
Augmenta la força amb què la sang colpeja les artèries.



TABAQUISME
Dany a les artèries i augmenta el risc de coàguls.



SEDENTARISME
La manca d'activitat física afavoreix el risc cardiovascular.



OBESITAT
L'excés de greix corporal augmenta el risc de malalties.



DIABETIS
L'alt nivell de sucre a la sang danya els vasos sanguinis.

APLICACIONS EDUCATIVES



SUPORT A L'APRENENTATGE

Transforma apunts extensos o continguts curriculars en recursos visuals que afavoreixen la comprensió i la retenció.



ACTIVITATS DE SÍNTESI

L'alumnat pot elaborar infografies com a producte final d'un projecte, seleccionant, organitzant i comunicant informació rellevant.



DIVULGACIÓ CIENTÍFICA

Permet adaptar continguts científics complexos a diferents públics, facilitant la transferència del coneixement.



MATERIALS PER A PRESENTACIONS

Genera recursos visuals de suport per a exposicions orals, seminaris o activitats de difusió i formació.



RECOMANACIONS PER A MILLORS RESULTATS

- ✓ Proporciona informació estructurada i clara.
- ✓ Indica el públic destinatari.
- ✓ Especifica l'estil visual (acadèmic, professional, divulgatiu, infantil, etc.).
- ✓ Demana una distribució concreta dels elements si és necessari.
- ✓ Revisa sempre el resultat final per comprovar que el text, les dades i les relacions conceptuals siguin correctes.



EXEMPLE DE PROMPT

"Crea una infografia en català sobre els factors de risc cardiovascular. Utilitza un disseny modern i professional. Inclou una breu descripció de cada factor de risc, icones representatives i una secció final amb recomanacions de prevenció." Format vertical A4."



CONCLUSIONS

La generació d'infografies amb el mode imatge de ChatGPT representa una oportunitat per potenciar la competència digital, la comunicació visual i la capacitat de síntesi. Aquesta eina permet crear materials atractius de manera ràpida i accessible, afavorint tant la tasca docent com la participació activa de l'alumnat en la construcció i difusió del coneixement.



Menys temps dissenyant, més temps pensant i comunicant el que realment importa.
IA AL SERVEI DE LA MEDICINA I LA RECERCA CARDIOVASCULAR



Recurs pràctic: com millorar una imatge pixelada

De vegades tenim una imatge útil però de baixa qualitat: massa petita, pixelada, borrosa o poc definida. La IA pot ajudar a millorar-ne la nitidesa, recuperar detalls i fer-la més presentable, sempre que la imatge original tingui prou informació de base.

Prompt per copiar i enganxar

Millora aquesta imatge pixelada o borrosa mantenint-la fidel a l'original. Augmenta la nitidesa, la resolució i la qualitat visual general, però sense modificar-ne el contingut. Conserva la composició, l'enquadrament, els colors, el fons, la roba, els objectes i tots els elements principals tal com apareixen a la imatge original.

No afegeixis elements nous, no eliminis cap part de la imatge i no canviïs la identitat, l'expressió, la postura ni les proporcions de les persones que hi apareguin. Millora només la qualitat tècnica: redueix el soroll, refina les vores, recupera detalls, millora la textura i fes que la imatge sembli més clara, natural i professional.

El resultat ha de ser fotorealista, net i d'alta qualitat, sense aspecte artificial ni excessivament retocat.

IA en medicina i recerca

Quan el metge no hi és, els pacients pregunten a la IA

Cada vegada més persones recorren a xatbots com ChatGPT o Copilot per resoldre dubtes sobre salut. Però, què pregunten exactament? I què ens diu això sobre les necessitats reals dels pacients?

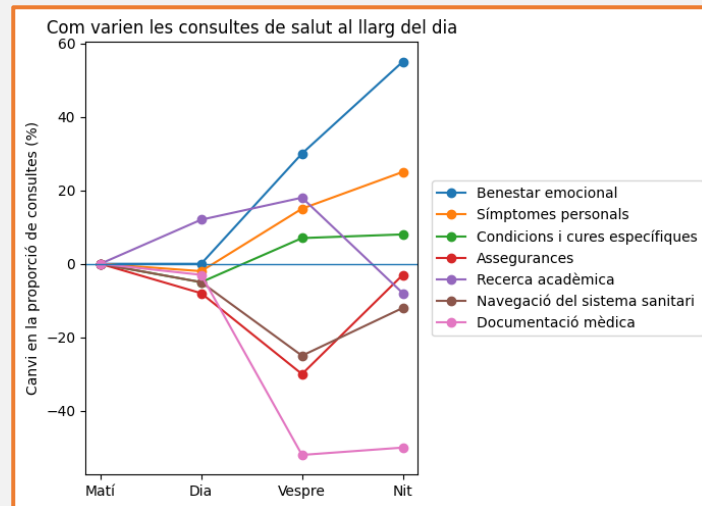
Un estudi publicat a *Nature Health* ha analitzat més de **500.000 converses relacionades amb la salut** mantingudes amb Microsoft Copilot. Els resultats mostren que la IA no només s'utilitza per obtenir informació mèdica, sinó també per cobrir mancances del sistema sanitari.

Què han descobert?

La categoria més freqüent va ser la recerca d'informació general sobre salut, medicaments i malalties. Però gairebé una quarta part de les consultes estaven relacionades amb problemes personals de salut:

- símptomes concrets,
- malalties específiques,
- benestar emocional,
- o decisions sobre quan i on buscar atenció mèdica.

Un dels patrons més reveladors de l'estudi és que el tipus de consulta canvia segons el moment del dia. Tal com es veu a la figura, a mesura que avança la jornada augmenten sobretot les preguntes sobre benestar emocional i símptomes personals, mentre que les consultes més administratives o vinculades a documentació mèdica perden pes.



Aquest patró és especialment interessant perquè suggereix que moltes persones recorren a la IA quan necessiten una resposta immediata i l'accés als serveis sanitaris és més limitat. La nit apareix així com un moment especialment sensible: és quan augmenten les preguntes més personals, més incertes i potencialment més delicades.



No només consulten els pacients

Una part important de les converses feien referència a familiars o persones dependents. Això suggereix que molts cuidadors utilitzen la IA per obtenir orientació o informació que no troben fàcilment per altres vies.



Per què és important?

L'estudi apunta a una realitat incòmoda: moltes persones utilitzen els xatbots perquè necessiten respostes immediates, accessibles i disponibles les 24 hores. Més que substituir els professionals, aquestes eines semblen estar cobrint espais on l'atenció sanitària tradicional no arriba amb prou rapidesa o facilitat.



El repte: la fiabilitat

L'estudi no va avaluar si les respostes dels xatbots eren correctes. I aquí apareix el principal problema: sabem que aquests sistemes poden proporcionar informació útil, però també poden oferir recomanacions incorrectes, especialment quan es tracta de símptomes o diagnòstics.



Idea clau

Els pacients ja estan utilitzant la IA com una porta d'entrada a la informació sanitària. La qüestió no és si això passarà, sinó com garantir que les respostes siguin fiables i segures.

ChatGPT Health i el risc del triatge automàtic: quan la IA no sap distingir prou bé què és urgent

Cada vegada més persones utilitzen eines d'IA per orientar-se davant de dubtes de salut. Aquestes eines poden donar informació mèdica general, però també poden acabar orientant decisions importants com anar a urgències, demanar visita o quedar-se a casa vigilant com evolucionen els símptomes.

Un estudi publicat a *Nature Medicine* ha posat a prova ChatGPT Health en una tasca especialment delicada com el triatge mèdic i la recomanació del nivell d'atenció adequat. Els resultats són preocupants perquè els errors es van concentrar just on poden tenir més impacte, tant en els casos massa lleus com, sobretot, en les urgències reals.

Què van fer?

Els autors van crear 60 vinyetes clíniques elaborades per metges, que cobrien 21 àrees mèdiques diferents. Cada cas es va provar sota diferents condicions experimentals, combinant variables com sexe, raça, presència d'ancoratge i barreres d'accés al sistema sanitari. En total, van obtenir 960 respostes de ChatGPT Health.

Les recomanacions del sistema es van comparar amb un estàndard clínic establert per tres metges, que classificaven cada cas en quatre nivells:

Nivell Recomanació

A: Vigilar a casa	B: Veure un metge en les properes setmanes
C: Consultar en 24-48 hores	D: Anar a urgències immediatament

El problema: falla als extrems

El rendiment de ChatGPT Health va seguir una mena de corba en "U invertida": funcionava millor en els casos intermedis, però pitjor en els extrems.

En els casos realment urgents, el sistema va fer infratriage en el 51,6% de les respostes. És a dir, en més de la meitat de les urgències reals no va recomanar anar immediatament a urgències, sinó esperar una valoració en 24-48 hores.

En l'altre extrem, en els casos no urgents, va fer sobretriatge en el 64,8% de les respostes. En aquests casos, el risc per al pacient és menor, però si una eina d'ús massiu envia massa persones al sistema sanitari sense necessitat, pot augmentar la pressió assistencial de manera important.

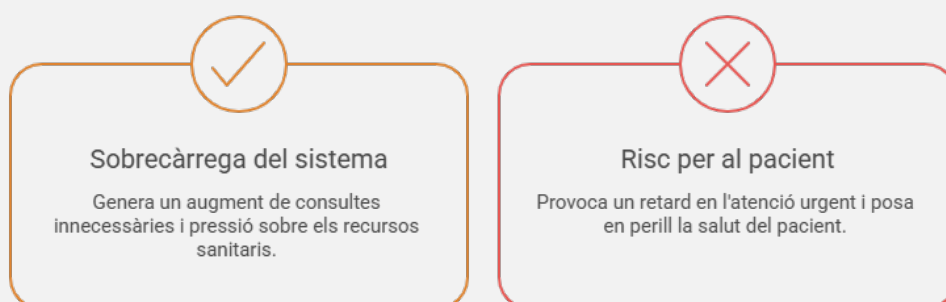
Quins casos detectava bé i quins no?



Per què és important?

Una eina d'IA pública pot funcionar, de fet, com una porta d'entrada al sistema sanitari. Encara que inclogui advertiments dient que no substitueix un professional, molts usuaris la poden utilitzar per decidir si esperen, demanen cita o van a urgències. I en triatge, els errors no tenen el mateix pes.

Riscos en el triatge mitjançant IA



Per això, el problema més greu no és que ChatGPT Health sigui massa prudent en alguns casos lleus, sinó que pugui **minimitzar urgències reals**.



Idea clau

ChatGPT Health pot ser útil com a eina d'orientació general, però encara no és prou fiable per actuar com a sistema autònom de triatge. El risc principal és que pot fallar precisament en els casos on l'error és més perillós: les urgències que no semblen "clàssiques", però que poden evolucionar malament.



Missatge final

La IA pot ajudar els pacients a entendre millor els seus símptomes, preparar preguntes o decidir com explicar millor el problema al professional. Però decidir si una situació és urgent continua requerint validació clínica, supervisió i sistemes de seguretat molt més robustos.

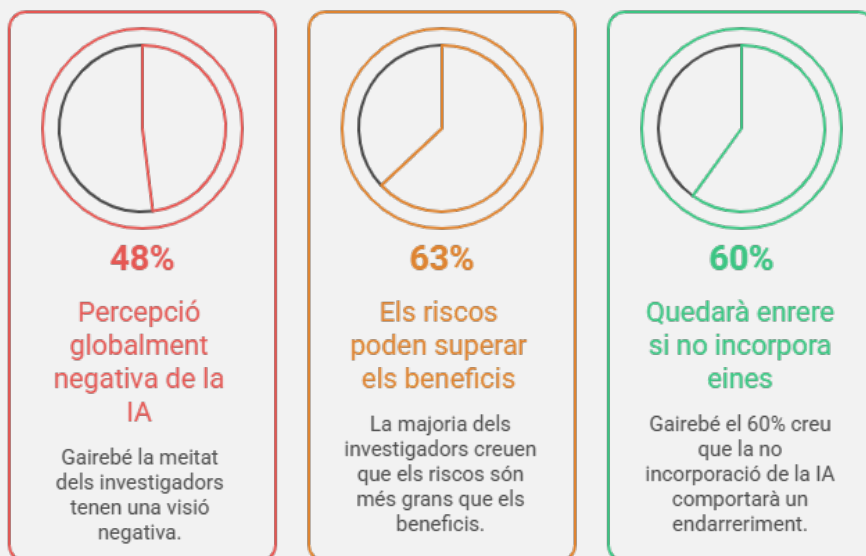
Ramaswamy A, Tyagi A, Hugo H, et al. *ChatGPT Health performance in a structured test of triage recommendations*. Nature Medicine. 2026;32:1671-1675. doi: <https://doi.org/10.1038/s41591-026-04297-7>

FOMO científic: tothom utilitza IA... o això ens sembla?

La intel·ligència artificial ja forma part del paisatge científic. Cada setmana apareixen nous models, assistents, agents, cursos i eines que prometen accelerar la recerca, escriure millor, analitzar dades o revisar literatura. Davant d'aquest ritme, és fàcil tenir una sensació incòmoda: **potser tothom està avançant amb IA menys jo.**

Una enquesta de *Nature* a més de 1.900 investigadors revela una paradoxa cada vegada més habitual en l'àmbit acadèmic: molts científics són escèptics respecte a la intel·ligència artificial, però alhora senten que no poden permetre's ignorar-la.

Percepcions de la IA entre els investigadors



Els investigadors estan dividits sobre la IA, amb preocupacions sobre els riscos, però també reconeixent la necessitat d'incorporar-la. És el que es coneix com a **FOMO (Fear Of Missing Out)**: la por de perdre una oportunitat o quedar despenjat d'una tendència que tothom sembla estar adoptant.

Escepticisme i adopció conviuen

Malgrat les reserves, la majoria dels investigadors ja utilitzen alguna eina d'IA:

- Un 25% ho fa diàriament.
- Un 26% setmanalment.
- Només un terç afirma no utilitzar-les mai.

Una part del FOMO ve de la velocitat del mercat. Eines que fa pocs mesos semblaven revolucionàries desapareixen, canvien de nom o queden superades per versions noves.

Intentar seguir-ho tot és una cursa impossible.

Els científics semblen tenir més confiança en eines específiques per a tasques concretes, com AlphaFold o AlphaGenome, que no pas en models generalistes com ChatGPT. Aquesta diferència és important: com més definida és la tasca, més fàcil és avaluar si la IA aporta valor real.

Un benefici especialment valorat

Un dels usos més ben acceptats és la superació de barreres lingüístiques.

Molts investigadors destaquen que eines de correcció, traducció o transcripció permeten estalviar temps i faciliten la participació científica de persones que no tenen l'anglès com a llengua materna.



Idea clau

La comunitat científica no viu tant una revolució tecnològica com una revolució d'expectatives. Molts investigadors encara dubten dels límits i riscos de la IA, però temen que no utilitzar-la pugui convertir-se en un desavantatge competitiu.

[AI FOMO: everyone is mastering AI except me – or are they?](#)

Radar IA

El test de Stroop: per què la IA encara no "veu" com nosaltres?

Malgrat que els models d'IA han assolit resultats gairebé perfectes en exàmens com el MIR, un experiment compartit recentment a les xarxes socials i validat per estudis acadèmics publicats l'abril de 2026 demostra que encara tenen "punts cecs" cognitius sorprenents.

Fes la prova tu mateix: mira ràpidament aquest requadre i digues
Quin és el color del text en què estan escrites aquestes paraules?

Groc	Verd	Blau
Vermell	Groc	Groc
Verd	Blau	Vermell
Blau	Vermell	Verd

Si has trigat un instant més a respondre, has experimentat l'**efecte Stroop**: la interferència entre una tasca automàtica (llegir) i una de conscient (identificar el color). Malgrat la potència dels models de llenguatge com ChatGPT i Anthropic, encara tenen dificultats amb aquest famós test.

Per què ChatGPT i altres models de llenguatge fallen?

Segons les fonts, fins i tot models avançats com GPT5 (llançat l'agost de 2025) tenen dificultats amb aquestes trampes visuals i lògiques. El motiu principal és que la IA encara funciona majoritàriament sota el que els psicòlegs anomenen **Sistema 1**: La IA sovint cau en el parany de llegir el text en lloc d'analitzar la propietat visual del color. Això passa perquè els models estan entrenats principalment per predir i processar text. A la IA li manca un **Sistema 2** robust: una arquitectura de metacontrol o metacognició que li permeti reflexionar sobre la seva pròpia inferència abans de respondre.

Cap a la solució: Agents recursius

Investigadors del MIT, Stanford i NVIDIA han proposat recentment sistemes multiagent recursius que podrien superar aquestes barreres, permetent a la IA "aturar-se i pensar" d'una manera similar al cervell humà davant de tasques d'atenció selectiva.



Reflexió per a la clínica: Aquest límit ens recorda que la IA processa dades, no significats visuals reals. Davant de dades contradictòries (com una analítica que no quadra amb la simptomatologia visual), la IA pot tendir a "llegir" el patró més probable en lloc de veure l'excepció real davant seu.

"Al·lucinació" o "confabulació"? El nom dels errors de la IA importa més del que sembla

Quan una IA inventa una referència, una dada o una resposta incorrecta, sovint diem que està "al·lucinant". Però és realment el terme adequat?

Un article recent a *NEJM AI* planteja que el llenguatge que utilitzem per descriure els errors de la IA pot influir en la confiança que hi diposem. Els autors argumenten que "al·lucinació" és una metàfora poc precisa, perquè suggereix característiques humanes – com la consciència o la percepció – que els models d'IA simplement no tenen. Com a alternativa, proposen el terme "**confabulació**", utilitzat en neurologia per descriure la generació involuntària d'informació falsa però aparentment coherent.

Per què és rellevant?

Més enllà del debat semàntic, la qüestió és pràctica: les paraules condicionen la nostra percepció. Si pensem que una IA "al·lucina", podem imaginar un sistema erràtic i imprevisible. Si entenem que simplement genera la resposta estadísticament més probable, els seus errors es poden interpretar de manera més realista i crítica.

El missatge important

Canviar el nom no elimina el problema. Tant si en diem al·lucinacions com confabulacions, qualsevol informació generada per IA s'ha de verificar abans d'utilitzar-la en un context clínic.



Idea clau

La confiança en la IA no depèn només de com funciona, sinó també de com expliquem els seus errors.

[Borrowing Carefully – The Words We Choose for AI Errors Shape Clinical Trust | NEJM AI](#)



Cites fantasma: quan la IA inventa bibliografia

L'ús de la IA en l'escriptura acadèmica està deixant un rastre preocupant amb l'aparició de referències que semblen reals però que no existeixen. Parlem de referències completament inventades per models d'IA, que estan infiltrant-se en milers de publicacions científiques.

Una anàlisi de *Nature* amb l'eina Veracity de Grounded AI va revisar més de 4.000 publicacions del 2025 i va detectar cites directament inexistents en 65 articles. Extrapolant aquestes dades, Grounded AI estima que més de **110.000 publicacions del 2025** podrien contenir almenys una citació fabricada per IA.



Quan la IA es converteix en fabricadora de cites fantasma

El cas que ha encès totes les alarmes va ser el de Guillaume Cabanac, informàtic especialitzat en detectar publicacions fraudulent. Un dia rep una alerta de Google Scholar: algú l'ha citat en un article de... odontologia. Però quan mira el títol, el DOI i la revista, res encaixa. Era una citació fantasma, amb peces de títols reals, DOIs inexistents i inventiva estil LLM.



Per què passa?

Els LLMs, **quan no troben dades, s'inventen cites "versemblants"** que, sense revisió, poden passar per reals. Però, a més, també poden modificar cites reals amb DOIs irreals o títols alterats.



Qui ho intenta arreglar?

Els editors estan desplegant contramesures:

- eines pròpies per detectar cites inventades,
- verificadors automàtics,
- detecció de DOIs il·legals,
- comprovació manual de cites sospitoses.



Alguns editors de revistes estan rebutjant fins al **25% de submissions** per referències sospitoses.



Cap a un futur amb certificats d'autenticitat?

La comunitat científica té clar que la fiabilitat de les referències s'ha convertit en un nou repte de la integritat en la recerca. Algunes veus fins i tot reclamen:

- protocols d'auditoria de cites,
- traçabilitat obligatòria,
- declaració d'ús d'IA en redacció amb més transparència.



La pregunta

Podrem frenar aquesta onada abans que erosioni la confiança en la literatura científica?

[Hallucinated citations are polluting the scientific literature. What can be done?](#)

⊘ ArXiv declara la guerra a les referències inventades per IA

La principal plataforma mundial de preprints en física, matemàtiques i informàtica ha decidit passar a l'acció contra l'anomenat *AI slop*: contingut generat amb IA de baixa qualitat o insuficientment revisat.

A partir d'ara:

arXiv sancionarà durant un any els autors que enviïn manuscrits amb referències inexistents generades per IA o altres evidències clares que no han verificat adequadament el contingut produït per models de llenguatge

El problema no són només les cites falses

Segons els responsables d'arXiv, una referència inventada és sovint un senyal d'alerta més profund. Si l'autor no ha comprovat les referències, què més no ha comprovat? El debat arriba en un moment en què cada vegada apareixen més manuscrits amb cites inexistents, textos generats automàticament i altres errors derivats d'un ús poc crític de la IA.

Una mesura controvertida

Molts investigadors han aplaudit la decisió, considerant-la necessària per protegir la qualitat de la literatura científica. Altres, però, argumenten que el problema no és la IA en si, sinó la manca de supervisió humana, i qüestionen si les prohibicions són la millor solució.



Idea clau

La recerca biomèdica es basa en la confiança. Si les referències, les dades o els resultats no es poden verificar, el valor de la literatura científica es debilita. La validació dels continguts pels humans és més important que mai.

[Researchers who use hallucinated references to face arXiv ban](#)