



Patienten har spurgt sin AI – og den lyder som en specialist

Det sker dagligt nu. En patient sætter sig og siger: "Jeg har spurgt ChatGPT, og den mener, det kan være..." Det er ikke længere en Google-søgning, patienten har med. Det er en samtale med et værktøj, der formulerer sig som en specialist, citerer studier og tilpasser svaret til netop deres symptomer. Og det lyder overbevisende.

Vores autoritet har været under pres i årevis. Men det her er anderledes. Det er sværere at afvise en patient, som har haft en times samtale med et AI-værktøj, der har sammenholdt deres blodprøvesvar med internationale retningslinjer, end én med et udprint fra Netdoktor. Spørgsmålet er, om vi forholder os til det nu – eller venter, til patienten kommer med et behandlingskrav.

Hvad kan patienten komme med i morgen?

Ifølge OpenAI stiller over 230 millioner mennesker sundhedsspørgsmål til ChatGPT hver uge. I januar 2026 lancerede OpenAI, Google og Anthropic sundhedsprodukter, der går langt videre end en søgemaskine. Amerikanske brugere kan allerede kombinere patientjournaler, data fra wearables og laboratorieresultater i én AI-samtale.

Det, der ændrer spillet, er ikke bare, at værktøjerne slår ting op; det er, at de nu bruger patientens egne data. En patient med diabetes kan indtaste oplysninger om blodsuktermålinger, kost og motion over måneder og sammenholde disse med data om egen behandling. Patienterne har indsigt i, hvad vi har gjort, og hvad de selv gør derhjemme. Og når de holder det op mod den nyeste viden om behandling, danner de sig et klart billede af, hvilken medicinering der passer bedst.

Hvad gør vi, når patienten følger den europæiske endokrinologiforenings nyeste anbefalinger, mens vi følger retningslinjer, der er et skridt bagefter?

Hvor er AI'en stærk – og hvor fejler den?

AI'ens styrke er dens omfattende viden om behandling: virkningsmekanismer, bivirkninger, interaktioner, internationale retningslinjer. Hallucinationer er nu et mindre problem, end det var for et år siden.

Det reelle problem er diagnosen. AI'en følger det mest sandsynlige spor. Hvis en patient skriver "Jeg er konstant tørstig og træt", går chatbotten ned ad diabetessporet. Men hvis det er midt i juli, og patienten arbejder udendørs, er det mere relevant at spør-

ge: "Har du drukket nok vand?" Det spørgsmål stiller vi. AI'en gør det ikke.

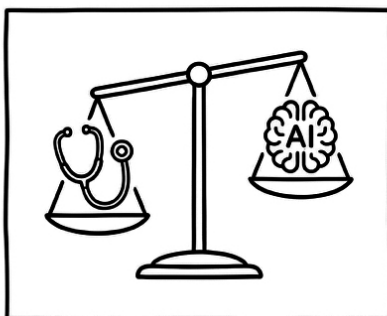
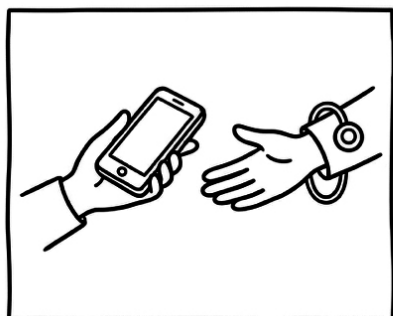
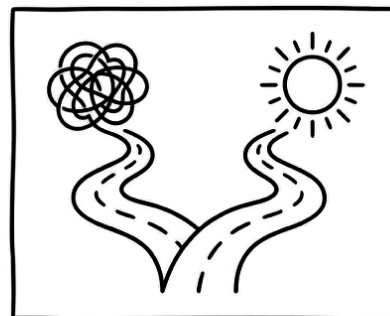
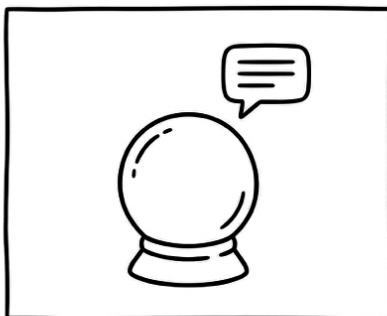
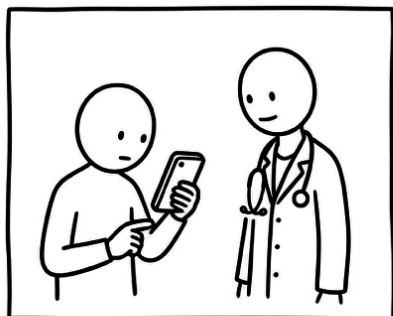
Det er vores egentlige styrke: Vi ser, når noget ikke passer. Det kan godt være, at 99 procent peger én vej – men noget stemmer ikke. Maskinen følger mønstret. Vi reagerer anderledes.

Første sikkerhedsevaluering af ChatGPT Health

I februar 2026 publicerede Mount Sinai i Nature Medicine den første uafhængige sikkerhedsevaluering af ChatGPT Health. I 960 testforløb undertriagerede systemet 52 procent af de akutte tilstande. Atypiske tilfælde – diabetisk ketoacidose med milde symptomer – blev henvist til "vurdering inden for 24-48 timer" i stedet for til akut læge.

AI-modeller, som er specialtrænet til medicin, klarede sig i flere tilfælde dårligere end almindelige modeller; de stolede mere på selvssikker klinisk formulering.

Vi så konflikten i stor skala med GLP-1. Men den version, jeg møder i praksis, er mere subtil. En patient med et BMI på 25 og let forhøjet kardiovaskulær risiko kræver GLP-1-behandling. Med ChatGPT har han fundet SELECT-studiet og konkluderet, at semaglutid beskytter mod hjertesygdom



AI-genereret

– også uden for studiets indikation. Jeg holdt fast. Studiet er ikke en indikation, og indikationen er ikke min at udvide.

Det er kernen: Der kan gå flere år, fra et studie påviser lovende resultater, til det bliver anbefalet behandling. Vores patienter læser med i realtid. Og de forstår ikke altid, at et spændende resultat ikke er en klinisk konklusion.

Vidensforsikring skaber ulighed

De sundhedsværktøjer, som bygger på AI, fungerer som en vidensforsikring: De giver hurtigere kontakt med sundhedsvæsenet, for nu kender man de ord, der sikrer en tid. Men præmien er tid og digitalt overskud. De, der har mest brug for det, har mindst adgang.

Den velinformerede patient møder op med "oversatte" prøvesvar og præcise spørgsmål til en fokuseret konsultation. Den anden – uden smartwatch, uden ChatGPT – får mindre ud af de samme 10 minutter. Når den forberedte patient kræver stillingtagen til behandlinger uden for retningslinjerne, involverer det visitation, second

opinions, korrespondance. Systemet risikerer at bruge stigende ressourcer på dem, der allerede er bedst stillet.

ECRI har udpeget utilsigtet afhængighed af AI-chatbots som den største sundhedsteknologiske risiko i 2026.

Ansvar og fremtid

I januar 2026 advarede det globale advokatfirma Clyde & Co om, at patienter snart vil kritisere lægen for ikke at følge AI'ens anbefaling. Ansvarsgabet er tydeligt: Vi har et autorisationsbaseret ansvar, mens AI-plattformen har en ansvarsfraskrivelse. Den leverer funktionel klinisk rådgivning. Juridisk leverer den underholdning – det er jo rart, at advokaterne ikke føler sig truet...

Men fremtiden er ikke kun mørk. AI-skribenter har forbedret arbejdsglæden for mange læger. AI-værktøjerne kan potentielt nå de patienter, vi har svært ved at nå. Almen praksis er ved at blive det speciale, der går forrest i anvendelsen af AI – fra tidligere at være blandt de mindst teknologiorienterede til nu at fremstå som en af de mest agile.

Prøv det. Investér i et af de store AI-værktøjer. Se, hvordan svarene er konstrueret.

Patienterne har allerede en digital ekspert med, når de kommer. Vi skal ikke konkurrere med den. Vi skal være det kliniske filter, som den mangler. Du kender dine patienter bedst. ●

TRE GODE RÅD TIL PATIENTEN

1. Brug AI som forberedelsesværktøj, ikke diagnoseværktøj. Du kan skrive: "Hjælp mig med at beskrive mine symptomer for min læge. Stil ikke diagnoser." Eller spørg bagvendt: "Hvilke symptomer taler imod, at det er det her?"
2. Vær opmærksom på, at alt, du uploader til et AI-værktøj, potentielt kan deles. Så undlad at oplyse dit CPR-nummer – det hører ikke hjemme i en chatbot.
3. Husk, at uanset hvor selvsikkert svaret fra AI-værktøjet lyder, kender den ikke dig. At det lyder rigtigt, betyder ikke, at det er rigtigt.