

NÅR MASKINEN LÆRER AT TÆNKE

MASKINEN, DER SVAREDE IGEN

Han åbnede døren til mit kontor og stak hovedet indenfor. Min begejstrede kollega smilede over hele ansigtet og snakkede løs om en ny digital service, der lige var blevet lanceret. Det var en tidlig decemberdag i 2022, og han var tydeligt imponeret over, hvor vild den var. Han havde lige fået den til at beskrive DNA og arvemasse, som var det et puslespil for børn – noget, der kræver både dyb viden og evne til at forstå en modtagers perspektiv. Den slags havde han aldrig set før.

Den service, min kollega talte så begejstret om, var ChatGPT fra det indtil da ret ukendte firma OpenAI. ChatGPT er en chatbot udviklet til at føre samtaler med brugere, som var den en menneskelig samtalepartner. Vi kunne for eksempel spørge, hvordan vi skulle svare på en svær mail fra en chef, og så foreslog chatbotten et velovervejet, diplomatisk svar – præcis som en klog bekendt kunne have gjort.

Resten af dagen lavede vi ikke meget andet end at lege med chatbotten og prøve dens grænser af. Vi ville finde ud af, hvor meget den egentlig vidste om verden, om den kunne forstå komplicerede spørgsmål, såsom

hvorfor himlen er blå, hvordan den ville formulere en joke om fodbold, og om den kunne forklare teksten i en TV-2-sang. Ja vi undersøgte sågar, hvordan den opfattede sig selv. ChatGPT var utroligt meget bedre end tidligere chatbots eller sagt på en anden måde: Den var vild! Faktisk præcis så vild som min kollega havde fornemmet, og vi var ikke de eneste, der lod os fascinere.

I løbet af blot to måneder havde ChatGPT fået 100 millioner brugere og er således den hurtigst voksende digitale service nogensinde. Til sammenligning tog det Facebook mere end fire år at opnå det samme. ChatGPT inkarnerede dermed en ny type digitale systemer baseret på kunstig intelligens, også kaldet AI, der er en forkortet version af det engelske udtryk *artificial intelligence*.

På rekordtid er de gået fra at være relativt ukendte for de fleste til at være omtalt vidt og bredt i samfundsdebatten og et værktøj i hverdagen for mange. Ledelseskonsulenter bruger dem til at forberede præsentationer og skrive udkast til analyser, arkitekter til at visualisere bygninger, studerende til at opsummere komplekse tekster, og virksomheder som banker, teleudbydere og detailkæder til at automatisere kundeservice.

Nogle ser AI som et redskab til at løse alt fra hverdagens små opgaver til de helt store udfordringer, vi står over for som samfund, lige fra at overskue den konstante strøm af nyheder til at håndtere klimakrisen. Andre frygter netop derfor maskinernes totale kontrol over vores dagligdag, og at vi som mennesker kører os selv ud på et sidespor, mens vi lader dem overtage vores arbejde og

måske endda kortslutte vores måde at tage demokratiske beslutninger på.

Allerede nu hjælper AI journalister med at skrive de nyheder, vi læser i Politiken, screene de kandidater, der skal til jobsamtale i Pandora, og udvælge de holdninger, vi møder på Facebook. Men selv om AI i dag pludselig virker til at være alle steder i vores liv, var tænkende maskiner indtil for få år siden ren og skær fantasi. En fantasi, der langsomt er blevet mere og mere virkelig siden den første spæde start.

KNÆKKEDE KODER

I efterkrigstidens laboratorier diskuterede it-pionerer som den britiske matematiker Alan Turing og hans ungarske kollega John von Neumann, om de kunne få en maskine til at efterligne menneskers måde at tænke på. Begge havde spillet væsentlige roller under Anden Verdenskrig: Turing havde brudt den tyske hærs koder, så de allierede kunne følge med i fjendens kommunikation, og von Neumann havde udtænkt de matematiske modeller for de amerikanske atombomber, der udslettede Hiroshima og Nagasaki.

Deres bidrag til at udvikle de første computere og plante kimen til AI har om muligt påvirket menneskeheden endnu mere. Turing udviklede sågar en berømt test, der skulle afgøre, hvorvidt en maskine kan tænke eller ej. Det kan den ifølge ham, hvis et menneske ikke kan gennemskue, om det taler med et andet menneske eller en maskine.

Han forestillede sig en person føre en skriftlig sam-

tale med både et menneske og en maskine uden at vide, hvem der var hvem. Hvis maskinen kunne overbevise personen om, at den var menneskelig, ville den bestå testen.



Turings tankeeksperiment lagde grunden til mange af de overvejelser, vi stadig har om AI i dag, hvor både vores frygt og fascination næres af, at vi i mange sammenhænge ikke længere ved, om vi interagerer med maskiner eller mennesker: Når vi chatter med kundeservice på nettet og ikke opdager, at det er en bot, der svarer. Eller når vi læser en artikel, der er skrevet af en algoritme.

Von Neumann gik et skridt videre og søgte inspiration i menneskets måder at tænke på for at udvikle mere effektive computere. Han så hjerner som maskiner til at behandle information og anvendte begreber som 'lagring', når et minde fæstner sig i vores hukommelse, og 'databehandling' til at beskrive vores tankeprocesser.

Dermed fik vi på godt og ondt et fælles sprog for, hvad intelligens *kan* indebære – både for mennesker og maskiner. For eksempel bruger vi i flæng udtryk som at 'processere' information fra et møde eller at 'nulstille' vores tankemønstre. Og når vi glemmer noget, føles det, som om det er 'slettet' fra vores erindring.

Nok så vigtigt udviklede von Neumann nogle af de første teorier om, hvordan maskiner kunne lære af data, så de ikke blot kunne udføre forudbestemte instruktioner, men også tilpasse deres måde at handle på ud fra erfaring. Indsigterne er i dag fundamentet for moderne AI-systemer, der analyserer store mængder data, såsom patientjournaler fra et hospital eller trafikdata fra en by.

Systemerne kan så identificere mønstre i de mange data, for eksempel om der er sammenhæng mellem bestemte symptomer og en sygdom, eller hvor der opstår køer på vejene. Derefter kan de så foreslå en diagnose til en læge eller optimere ruter for lastbiler.

På den måde har von Neumanns arbejde dannet basis for systemer, som i dag kan alt fra at genkende vores ansigter via overvågningskameraer på offentlige pladser til at skabe realistiske deepfake-videoer, hvor politikere udtrykker ting, de aldrig nogensinde har sagt. Men der skulle gå næsten et halvt århundrede, før andre forskere begyndte at knække koden til at virkeliggøre de to mænds visioner. På et skakbræt i New York en forårsdag i 1997.

SKAKMAT SKAKSPILLER

Den aserbajdsjanske verdensmester i skak Garri Kasparov sad over for en maskine ved navn Deep Blue. Kasparov havde året før besejret maskinen, men denne gang var noget anderledes. Deep Blue spillede med Kasparovs ord ikke som en maskine, men som en kreativ modstander med strategisk finesse, og den foretog træk, der kom helt bag på ham.

Kasparovs frustration var malet i hans ansigt igennem spillet, og da han rejste sig, slagen, stod det klart for verden, at noget fundamentalt havde ændret sig. En maskine, der blot udførte beregninger på kommando, havde overvundet det menneskelige geni, der i årevis havde regeret i skakverdenen.

Under den kolde krig var der på begge sider af

Jerntæppet enorm prestige i at sætte sig på tronen som verdensmester, da både politikere og meningsdannere så skak som den ypperste test af vores evne til at tænke strategisk og udmanøvrere en modstander. Da Deep Blue triumferede, eksploderede debatten derfor ikke kun blandt skakentusiaster, men også forskere, teknologer, filosoffer og politikere.

Nogle så det som bevis på maskinens potentiale. For andre rejste det svære overvejelser, som er lige så aktuelle i dag: Hvis en maskine kan tænke, rækker det ved den basale ide, at vi som mennesker besidder helt unikke kognitive evner. Vi har jo netop forestillet os, at det at kunne tænke adskilte os fra dyrene og gav os kontrol over verden omkring os.

Spol tiden frem til en hverdagsaften i 2025, hvor jeg efter en lang dag på kontoret kaster mig på sofaen i min lejlighed i det centrale Aarhus og instinktivt tænder for Netflix for at koble af. Et forslag til en god tv-serie toner straks frem, *The Queen's Gambit* fra 2020, der pudsigt nok handler om et skakgeni i efterkrigstiden. Serien er øverst i en nøje udvalgt liste af anbefalinger skræddersyet til netop min smag i film og serier.

Den amerikanske streamingtjenestes algoritmer benytter kunstig intelligens til at analysere millioner af brugeres seervaner for at forudsige, hvad vi kunne have lyst til at se. Denne usynlige hånd, der er med til at påvirke mine valg, illustrerer ret godt, hvordan AI baseret på at genkende mønstre arbejder i baggrunden og tilpasser sig vores individuelle præferencer. Men selv om det ved nærmere eftertanke kan virke ekstraordinært, er