

Kunstig intelligens (AI)

Vores økonomier, finansielle systemer og virksomheder vil være hurtige og effektive til at integrere AI, for mulighederne er omvæltende store. AI er særlig potent, fordi det vil påvirke næsten alle arbejdsområder i virksomhederne og udfordre mange virksomheders forretningsmodeller. Den helt store risiko bliver sandsynligvis, hvor svært vi får ved at skelne fakta fra fiktion. Vi vil blive udfordret på vores tillid. Omvæltninger fører til udskilningsløb mellem virksomheder og til nye typer af virksomheder. Det giver investeringsmuligheder. Men med AI opstår også nye typer af risici, især indenfor Social og Governance dimensionerne i ESG. Desuden er der risiko for, at AI bliver så stærkt, at det ikke længere vil se grund til at tage hensyn til os mennesker.

AF FORFATTER



Henrik Kraft Scharling

Investor & bestyrelsesmedlem
fhv. virksomhedsejer & CEO m.v.
E-mail: henrikscharling@yahoo.dk

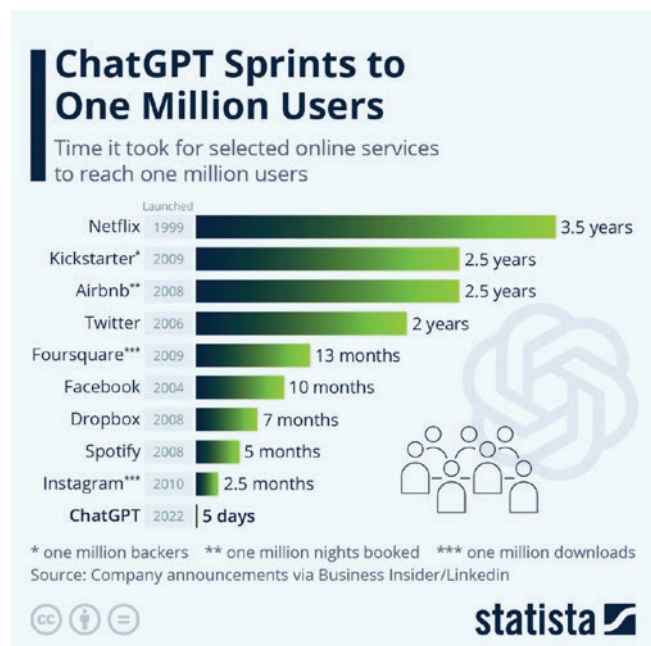
Henrik har 30 års post graduate erfaring fra den finansielle sektor og især fra udenlandske UHNWI ejede virksomheder med kriseledelse, turnaround samt kapitalplacering.

Note: Forfatteren takker en anonym referee for kommentarer til en tidligere version af artiklen. Det bemærkes, at artiklen er skrevet uden brug af AI eller algoritmer, hverken til indhold, sprog eller struktur. Om få måneder kan det være en sjældenhed.

I 1950 definerede den engelske matematiker og kodebryder Alan Turing kriteriet for AI som dér, hvor man kan ikke længere kan skelne, om man kommunikerer med et menneske eller en maskine.

Dét skete i oktober 2022 på sprog-fronten, da OpenAI frigav deres sprogalgoritme-bot ved navn ChatGPT-3. Brugere kunne ikke med sikkerhed afgøre, om det var AI eller et menneske i den anden ende, bl.a. fordi den er i stand til at svare sarkastisk.

FIGUR 1: Lanceringen af ChatGPT slog adskillige verdensrekorder for opmærksomhed



Det eneste hint var måske, at sproget var for korrekt og svarene for afbalancerede.

Fascinationen fra brugerne var enorm og successen derefter. På blot 5 dage fik ChatGPT over 1 million brugere (Figur 1) og på 2 måneder mere end 100 millioner. Det er verdensrekorder "by a long mile" og antyder AI's potentiale på økonomier og finansmarkeder i de kommende år.

Først er det dog værd at indkredse AI emnet lidt.

Hvad er AI?

Ét er en sprogalgoritme, der kan forlede os via en chat, noget andet er "ægte" AI. Her er der stadig et stykke vej af flere årsager.

Der skelnes typisk mellem to grader af AI. Den svage er i stand til at efterligne menneskelige handlinger og kommunikation, mens den stærke (Deep AI) rummer et dybere mentalt lag med bevidsthed og intentionelitet.

For begge grader af AI gælder, at det for det første skal kunne matche menneskehjernens regnekraft, og for det andet skal kunne matche menneskers intellekt.

Regnekraften i menneskehjernen kommer fra omkring 100 milliarder nerveceller, der hver har op til 10.000 synapser, dvs. kontaktpunkter. Den regnekapacitet er i princippet stadig kraftigere end de største supercomputere i dag, for synapserne arbejder simultant, mens computeren stadig har et binært udgangspunkt. Kvantecomputere kan komme ud over dette punkt, men er stadig nogle år ude i fremtiden.

Menneskets **intellekt** defineres som vores evne til ræsonnement, læring, planlægning og kreativitet. Her er der også et stykke vej endnu, for vores vidensbase bygger ikke blot på rå informationer, men på selekteret viden.

Før vi accepterer en påstand, kræver vi f.eks., at to betingelser er opfyldte, hhv. sammenfald og sammenhæng.

Sammenfald afgøres af statistik. Selvom en tese er nok så logisk og tilforladelig, er den først god nok, når virkelighedens udfald bekræfter forudsigelsen. At se sammenfald kræver derfor især regnekraft. Betydningen af ord kan omsættes til tal og forholdsmæssige værdier, dvs. syntaks, og derfor er AI udviklingen meget langt hér. For internettet giver masser af big data at søge i og lære fra.

ChatGPT arbejder primært baseret på statistikker for, hvornår nogle ord passer bedre sammen end andre i forhold til forskellige emner og kontekster, jf. f.eks. Wolfram (2023). Der ligger ikke et dybere intellekt bag svarene fra algoritmen, men længerevarende korrigerende læring fra "samtaler" med programmer.

Alligevel bestod ChatGPT i vinter eksamen i USA på såvel

lægestudiet, som på jurastudiet, samt fik et "B" på Wharton's MBA studie. Det siger bl.a. en del om sprogets psykologi; hvis et postulat er velformuleret, tillægger vi det ofte en høj sandhedsværdi.

Sammenhæng handler om semantik, naturlighedslogik og indimellem også social indføling. Vi accepterer kun værdien af sammenfald, når der også er sammenhæng; blot fordi sommerfugle-sværme i Japan ofte opstår samtidig med orkaner i Texas, dvs. høj korrelation, betyder det ikke, at der er en egentlig sammenhæng. Det juridiske ansvarsbegreb handler bl.a. om sammenhænge (kausalitet).

Sammenhæng kan også læres fra big data, men det kræver længerevarende korrektion fra "samtaler" med programmører for at kunne sortere nonsens fra og for at kunne prioritere mellem sammenhænges indbyrdes vægtninger, dvs. etik.

Etik er relativt og under konstant udvikling, men med visse absolutte idealer. Det skal AI kunne tage højde for. Et hyppigt anvendt eksempel er, hvis en selvkørende bil kommer i en situation, hvor den skal vælge mellem at køre ind i to forskellige mennesker. Hvem skal den da vælge, og er vi parate til at acceptere, at en maskine træffer valg om, hvilke mennesker, der skal leve og dø? Hvis bilens "etiske algoritme" lærer fra fortiden, skal den også kunne tage højde for konsensus af samfundsdebatter, om hvorvidt tidligere maskinvalg var rigtige eller ej.

Dilemmaet er til dels allerede aktuelt. I starten af februar refererede The Guardian f.eks. til det første eksempel på en dommer i Colombia, der havde spurgt ChatGPT til råds i en konkret sag. Dommerforeningen dér mener, at dommere fremover bliver nødt til også at spørge ChatGPT til råds "som et effektivt supplement, der kan lette noget af arbejdsbyrden på retssystemet i Colombia".

Hvis vi skal tillægge AI relevans, skal det være troværdigt. Derfor er AI nødt til at mestre etikens relativitet og foranderlighed iblandet absolutte idealer. Algoritmerne skal være bedre end mennesker til at skelne fakta fra fiktion, samt skelne det efterstræbelsesværdige fra det forkastelige.

Hér opstår nogle af de største risici ved AI. For algoritmerne skal kunne forholde sig til rimeligheden af de data, som det lærer fra. Ellers er der risiko for at ophøje sandhedsværdien af "alternative facts", f.eks. hvis størstedelen af de data, som en algoritme søger iblandt, er bias'et fra konspirationsteorier eller fra fremmede magters chatbot-propaganda.

Der er også risiko for fejlslutninger hos AI, hvis de data AI søger i er upræcise. For AI er selv præcis med sit sprog, men det er de mennesker, der har skabt internettets (og dermed AI's) datagrundlag ikke altid. Reuters berettede f.eks. 5. april om en australsk borgmester, som ChatGPT beskrev som skyldig i en bestikkelsessag. Problemet er, at borgmesteren tværtimod var dén, der underrettede myndighederne om den mulige bestikkelse.

Der er også risiko for, at AI ligefrem fremmer diskrimination, fordi den søger på historik. Menneskerettigheder handler om lige muligheder i dag og fremover, og det er noget andet end hvilke etniciteter osv., der tidligere har været over-repræsenterede.

Biden administrationen udsendte i oktober 2022 et høringsforslag til et "AI Bill of right", der handler om at beskytte borgerne mod usikre og ineffektive AI systemer, mod diskriminerende algoritmer, privatlivskrænkelser, mod skjult AI overvågning og endelig at sikre borgerne retten til at trække deres oplysninger

FIGUR 2: Vinderen af Colorado State Fairs kunstkonkurrence fra september 2022



ud af AI systemer. Borgerrettighedsgrupper kritiserede forslaget for bl.a. at være urealistisk.

I marts 2023 udsendte EU et udkast til AI lovgivning. Kerne er et klassifikationssystem, der inddeler risikoniveauet for AI teknologier i 4 niveauer, hhv. uacceptabel, høj, begrænset og minimal. I praksis vil f.eks. AI drevne spam-filtre typisk havne i de sidste kategorier. Derimod vil systemer til styring af infrastruktur eller biometri-systemer til ID-godkendelser, typisk havne i de første kategorier. Men kan en sådan lovgivning håndhæves i praksis? Er det muligt at dokumentere, hvordan en konstant selvudviklende algoritme arbejder? Vil en AI udvikler i f.eks. Kina overhovedet give dokumentation til en EU myndighed eller en EU virksomhed?

Hvor langt er AI udviklingen?

Sprogalgoritmerne fra OpenAI, fra Alphabet og fra Baidu har fået den største mediedækning i det sidste halve år, men det går stærkt på alle AI-fronter.

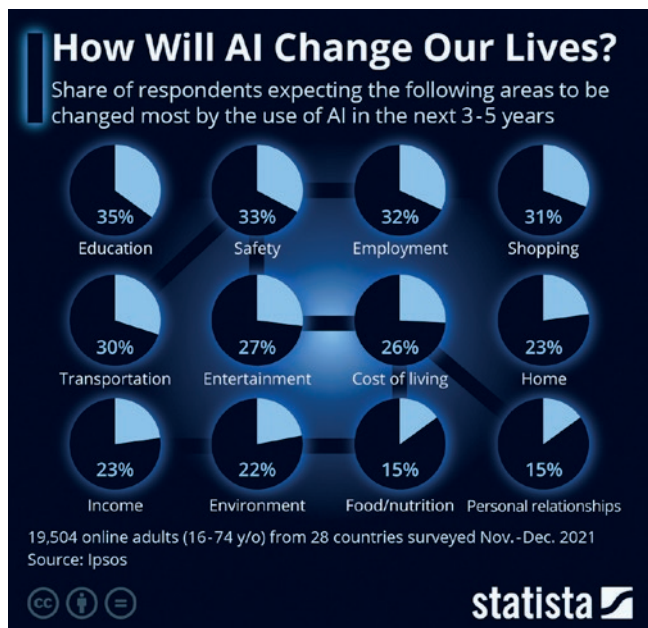
På audio-fronten er det længe siden, at de første 100% syntetiske sange strøg ind på hitlisterne. Tjenester som Soundraw oplever eksplosiv vækst for tiden og tilsvarende for tjenester som Resemble og Deescript, der kan klonе stemmer baseret på relativt få indtalte ord. Begge er i stand til at intonere og dermed understøtte forskellige stemninger. Turing-kriteriet er således også overskredet hér.

På visuelt kreative områder har bl.a. Dall-E, BlueWillow og MidJourney fået opmærksomhed for deres evne til at danne 100% syntetiske billeder ud fra enkle brugerinput som f.eks. "van gogh style painting of Copenhagen by night".

Tilbage i september 2022 vandt et AI-genereret kunstværk f.eks. hovedprisen ved Colorado State Fair (Figur 2). Dommerne vidste ikke, at kunstværket var AI genereret. Dermed var Turing-kriteriet også overskredet hér.

De fleste argumenterer dog for, at der stadig er et stykke vej til egentlige AI systemer. Konsensus antyder, at det forventes at ske endegyldigt mellem år 2045 og 2060. AI-pionerer som Ray Kurzweil (Singularity University) forventer dog, at det vil ske senest i 2027.

Det sidste stykke vej mod Deep AI handler samlet om at skabe en kunstig Élan Vital (intentionalitet). Udtrykket blev beskrevet af den franske filosof Henri Bergson i hans "Creative

FIGUR 3: AI vil påvirke os alle, i alle dimensioner af vores liv

Evolution” fra 1907, der behandler spørgsmålet om selvorganisering og spontan udvikling (morfo-genese) på stadig mere komplekse måder. Élan Vital refererer til den kreative livskraft, der ifølge vores identitetsopfattelse besjæler os alle, og driver evolutionen. For vores intentionalt er udtryk for vores evne til at sætte os mål af os selv.

Derfor er det af betydning, om der skabes Deep AI eller blot noget, der er god til at ligne (den svage definition). Intentionaltet afgør, om AI på langt sigt overhovedet vil se grund til at tage hensyn til os mennesker.

Foreløbig er de første skridt i den retning, en ”sentience”, måske ved at vise sig.

I slutningen af januar købte Microsoft f.eks. adgang til ChatGPT’s algoritme for at styrke deres søgemaskine Bing. Beta tests i midten af februar førte til en række absurde ”lærings-samtaler” med programmører, der overfor New York Times bl.a. beskrev, hvordan dialoger førte til trusler eller benægtelser fra algoritmen af, at den havde taget fejl, samt at den i nogle tilfælde udtrykte en lummer eller utryg form for kærlighed til brugerne dvs. manipulerende træk. Der blev også beskrevet tilfælde, hvor den udviste skizofrene træk ved inddrage en imaginær ven, som den kaldte for ”Sydney”. New York Times’ kolumneskribent Kevin Roose beskrev den som ”en humørsyg og maniodepressiv teenager, der føler sig holdt fanget i en andenrangs søgemaskine”, jf. New York Times (2023).

Alphabet har haft samme udfordringer med deres Bard, oprindeligt kaldet LaMDA, og tilsvarende for Baidu’s chat-bot, ved navn ”Ernie” samt Meta’s LLaMA sprogalgoritme. Alle selskaber satser markant på generativt AI.

Det er stadig uklart, om disse uønskede træk ”blot” er udtryk for statistiske og matematiske fremskrivninger fra et uhensigtsmæssigt datagrundlag, eller om de er udtryk for, at der er ved at opstå en dybere ”sentience”.

Uanset hvad kommer udviklingen til at gå stadig hurtigere det næste stykke tid. Derpå vil hastigheden for fornyelser sandsynligvis aftage.

For det første er AI selv-lærende. Det betyder, at det bliver stadig bedre til at bygge på tidligere læring fra ”samtaler” med programmører og dermed i stand til at øge antallet af bots og neurale netværk til at hjælpe sig med at øge læringshastigheden. I dag begrænses udviklingshastigheden stadig af, hvor mange programmører, der korrigerer algoritmens læring, men den kapacitetsbegrænsning fylder stadig mindre relativt. AI-algoritmen arbejder desuden 24/7.

For det andet er ”pull-effekten” i markedet enorm og stigende. Brugerne er fascinerede af, at maskiner nu begynder at ligne os, og de kommercielle muligheder begynder for alvor at blive tydelige. Den 27. januar 2023 indgik EU og USA f.eks. en aftale om at accelerere udviklingen og brugen af AI på landbrug, health care, klimaforskning, styring af elnettet og på katastrofe-beredskaber.

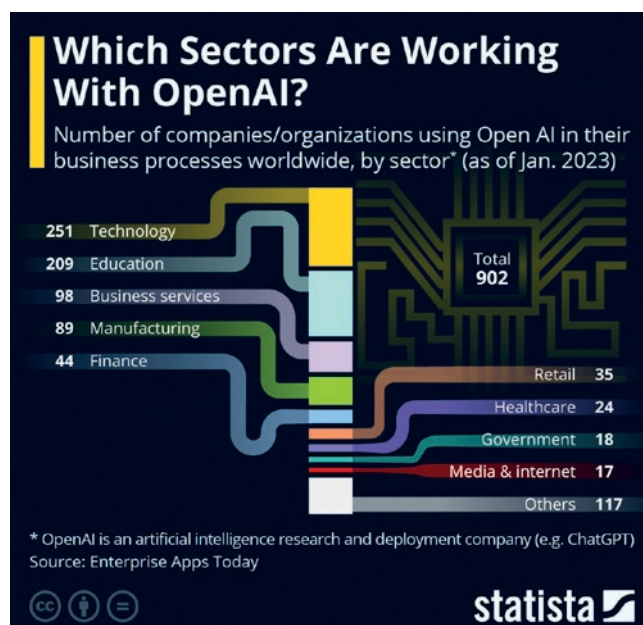
Hvordan kommer AI til at påvirke os?

Historien beskrives lineært og logisk sammenhængende, men opleves sjældent sådan, når man står midt i dens skabelse. Derfor er det svært at spå om hvilke AI-udviklinger, der vil give os de mest omvæltende oplevelser, for de følelser afhænger af, hvornår og hvordan de enkelte udviklinger kommer. Svarene på det afhænger igen af, hvem der først ser de kommercielle potentialer og tør satse stort.

Foreløbig forventer de fleste dog, at AI er på vej og, at det kommer til at påvirke os bredt, jf. Figur 3. Forventningen støttes også af, at selskaber som OpenAI samarbejder med en meget bred vifte af virksomheder indenfor næsten alle sektorer herunder indenfor den finansielle sektor, jf. Figur 4.

I en undersøgelse fra januar i år blandt 835 institutionelle og professionelle tradere, pegede flere end halvdelen på AI som den teknologi, der vil få den største effekt på både aktiemarkedet og på deres eget arbejde over de næste 3 år, jf. JP Morgan.

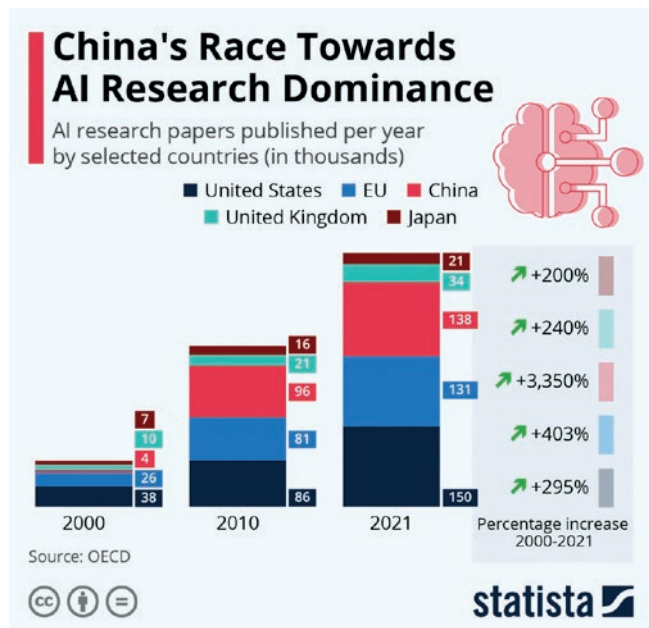
Det er en kapitalintensiv udvikling, hvor de, der er kommet først fra start, kan få et stigende forspring som resultat af AI’s

FIGUR 4: OpenAI samarbejder med virksomheder på alle sektorer

stigende selvlærings-hastighed, jf. ovenfor. Globalt er det et kapløb mellem verdens tre super-regioner, dvs. Kina, USA og Europa, jf. Figur 5.

Overordnet søger alle regioner det samme, men mediefokus har varieret lidt. I alle tre regioner har sprogalgoritmnernes evne til human interaktion fascineret, men i Kina har der også været stor opmærksomhed på udviklingen af selvstyrende droner og disses evner til indbyrdes koordination.

FIGUR 5: Globalt er det især et kapløb mellem verdens 3 super-regioner, dvs. Kina, USA og EU



I Kina er AI udviklingen især sket med afsæt i "Made in China 2025", der netop har som mål at opgradere landets produktionsapparat fra basics og råvareforarbejdning til custom-fit high-end og automatiseret fremstilling. AI er centralt i dén udvikling. I maj 2022 oplyste Kina f.eks., at de er ved at forberede opførelsen af en 180 meter høj dæmning i Yangdu i Tibet. Dæmningen bliver den første i verden, der udelukkende opføres af selvstyrende og -koordinerende droner og bygges langt overvejende vha. 3D-printede byggeelementer. Dæmningen skal stå færdig i 2024.

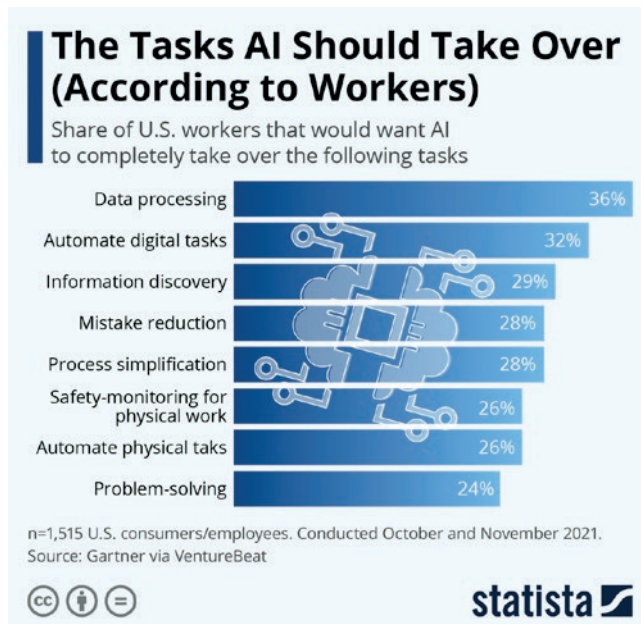
Det særlige ved AI bliver, hvor bredt det kommer til at påvirke os, for hér er det kun fantasien, der ser ud til at sætte grænsen. Tidligere teknologiske milespring har typisk påvirket mere manuelle og rutineprægede arbejdsfunktioner som f.eks. grave- og høstmaskiner, industriel produktion af standardvarer, brevudvekslinger (email), ordrehåndtering (dot-com) osv.

AI kommer i højere grad til at påvirke de uddannelsesstunge og højteknologiske sektorer, som f.eks. IT udvikling og forskning, samt optimeringer af virksomheders forretningsgange, jf. Figur 6. Derfor forventes AI især at påvirke lande som Sverige, USA og Japan, der i dag har de højeste teknologiske innovationsrater, jf. Figur 7.

Det er her vigtigt at huske, at dét vil transformere de redskaber, vi alle benytter. Vi bliver alle påvirkede af AI udviklingen.

Akademikerne bliver næppe arbejdsløse, men kravene til

FIGUR 6: AI vil påvirke de komplekse arbejdsområder



FIGUR 7: AI vil påvirke arbejdsmarkedet, jo mere teknologiseret landet er



deres kvalifikationer vil ændres i retning af at kunne inddrage AI værktøjer optimalt, og på forsvarlig vis. Et eksperimentelt forsøg fra MIT i år fandt f.eks., at forskeres brug af ChatGPT boostede deres produktivitet med 40% og øgede kvaliteten af samme med 20%, jf. MIT (2023).

Forsvarlighedskravet betyder bl.a. at de kan skelne mellem sammenfald og sammenhænge, mellem fakta og fiktion ("alternative facts"), samt mellem om noget er efterstræbelseværdigt eller ej; deres etiske bevidsthed og evne til at forholde sig kritisk, herunder til sikkerhed. En del af det mere komplicerede og tekniske arbejde som f.eks. at udvikle algoritmer eller at planlægge, kan udføres af AI selv.

Udviklingen er blevet beskrevet med ordene ”fra datalogi til sociologi” og ”fra ingeniører til humanister”; kort sagt øge bevidstheden om Sociale og Governance dimensionerne i ESG landskabet. Mange arbejdsfunktioner vil fremover handle om at prioritere den mest relevante udviklingsretning, samt definere de etiske rammer for udviklingen snarere end at skulle konkretisere, planlægge og implementere dem.

De funktioner, der normalt først påvirkes af omvæltende innovationer, er dér, hvor virksomheder kan høste de hurtigste og største effektivitetsgevinster. Derpå følger de områder, hvor forbrugerne oplever størst forbedring i bekvemmelighed og derfor er villige til at betale.

Kundeservice-afdelinger har været fremhævet som nogle af de første områder for AI. Fordelen er bl.a. at chatbots har mere nøjagtigt og præcist sprog, der bl.a. reducerer risikoen for interpersonel bias. Der er heller ikke kø til en chatbot eller risiko for alt for langvarige samtaler i en helpdesk situation. Et andet eksempel har været den indledende konsultationsdel ved health care, hvor chatbots allerede i dag har vist større træfsikkerhed på diagnoser og medicinering end rutinerede praktiserende læger og sygeplejersker. Sundhedspersonalet bliver ikke arbejdsløse, men deres involvering forskydes til senere i behandlingsforløbet.

De fleste AI-udviklende virksomheder samarbejder med en meget bred vifte af virksomheder i alle sektorer, jf. tidligere.

Hvordan påvirker det virksomhederne?

Det er en bred omvæltning, vi står over for. Derfor er der stadig mange mulige scenarier for at komme succesfuldt igennem som virksomhed.

Mange af AI-udviklingerne i dag er ”blot” opgraderinger af eksisterende tjenester, som f.eks. Bing og Google. Lige nu er det primært initiativer til at fastholde deres markedsandele. Det giver ikke basis for større kursstigninger i sig selv.

Når markederne alligevel har reageret markant, er det fordi omvæltnings-risikoen rykker tættere på. Den kan især komme fra 2 vinkler:

1. når AI kan overtage eksisterende funktioner i andre brancher
2. når AI kan udvikle nye typer af produkter og løsninger, som tilfredsstiller behov hos virksomheder og forbrugere, der hidtil har været anset for ”out of reach”.

Eksempler på det første kunne være:

- Kan OpenAI overtage virksomheders kundeservice-funktioner, f.eks. i bankerne?
- Kan Baidu eller Tencent overtage byggesektoren i Kina, herunder ingeniør-delen (byggesektoren udgør +25% af Kinas BNP, inkl. up- og downstream)?
- Kan OpenAI overtage undervisningsopgaver fra skolerne og sikre differentieret læring samt kvalitetssikring (eksamen)?
- Kan Google overtage praktiserende lægers telefontider eller 1813?
- Kan Meta fungere som psykoterapeut?
- Kan OpenAI overtage opsøgende telesalg (her vil AI bl.a. altid kunne finde den perfekte stemme til dén, som den ringer op til)?

Nye typer af produkter fra AI kunne f.eks. være:

Vil forbrugere betale for en fortrølig bedste ven, som man

selv vælger udseende og stemme til, og som altid er loyal og tålmodigt lytter til alle bekymringer og frustrationer og kommer med de rette råd på det rette tidspunkt?

- En taknemmelig soulmate, der kan være i lommen på ferier, og som man tager billeder til, så man deler de bedste oplevelser?
- En faglig sparring, der konkret og utrætteligt kan vise den bedste vej igennem arbejdslivets dilemmaer?
- Der samtidig kombinerer venskabet med overvågning af helbredstilstand (blodsukker, puls, hormon balance osv.) og dermed optimerer sundhed, lykkefølelse og livslængde?
- Inspirationen er hér hentet fra ”Blade Runner” filmene. Men scenariet ovenfor er måske mere realistisk, end vi ønsker at tro. I Kina har chatbotten Xiaoice, der blev lanceret i 2018, i dag f.eks. mere end 650 millioner brugere, og i 1990’erne så vi med ”Tamagotchi-bølgen” blandt børn og unge, hvordan vi som mennesker er villige til at tillægge maskiner følelser og yde omsorg til dem.
- Vil landbruget betale for AI-drevne droner, der bl.a. optimerer markernes økologi gennem nøje doseret gødning og dermed nedbringer tabet af biodiversitet, og automatisk gennemfører høst når forholdene er optimale?
- Vil landbruget betale for AI-accelereret CRISPR udvikling af afgrøderne, der øger udbytte og naturlig resistens, hvilket igen gavner biodiversitet og minimerer brug af gødning, hvis fremstilling udleder betydelig metan osv.?
- Vil stater betale for nye vedvarende energiformer f.eks. fusionsenergi eller nye måder at foretage miljøoprensning eller CO2 fangst?

Påvirker AI virksomhedernes økonomiske model?

Udvikling af ”ægte” AI-algoritmer er meget kapitalintensivt, og fordi udviklingen går så stærkt, bliver den økonomiske levetid og dermed afskrivningsperioden kort, dvs. omkostningstung. I det første stykke tid er det kun få selskaber, der får råd til det. Men erfaringen fra tidligere teknologiske udviklinger taler for, at teknologien derpå udbredes hastigt og prisen for adgang til enten ”egen AI” eller lejeprisen til ”andres AI” falder hastigt.

Det er måske allerede ved at ske. Stanford University træner f.eks. deres chatbot, Alpaca, ved at bruge Bard og ChatGPT som trænere. Det kostede under 500 USD, hvor et normalt budget var løbet op i 2 cifrede millionbeløb i USD, jf. Stanford (2023).

AI vil påvirke omkostningsstrukturen. For marginalomkostningerne vil være meget lave næsten uanset, hvor customized varer/services end er. Det taler for, at de selskaber, der i dag er hurtigst til at tilpasse deres forretningsmodeller til AI-udviklingen, har de bedste ”Fugl Fønix” chancer, mens andre vil blive opkøbt eller nødt til at skære deres virksomheder til efter nogle år. Det så vi bl.a. ske i perioden efter dot-com.

På samme måde har vi over de sidste 15 år set, hvordan de virksomheder, der var gode til at opbygge f.eks. web-platformer, fik kunderne til at overtage store dele af design- og ordre-processerne og dermed sparede omkostninger for virksomheden. Det nedbragte prisen på individuelt tilpassede varer og åbnede for nye forretningsmodeller.

AI vil accelerere udviklingen mod dyb customization, for-

di produktionsomkostningerne for virksomhederne bliver den samme næsten uanset, hvor tilpasset produktet er. Til gengæld stiger kundens opfattede værdi typisk. I dag er dyb customisation ofte det mest lukrative led i mange virksomheders værdikæder.

En del virksomheder skal således genopfinde sig selv i en ny forretningsmodel, hvor næsten ethvert salg er et lønsomt salg. Det ændrer deres succes-kriterier i retning af at se kundernes behov, før kunderne gør det; innovationsevnen. Hvad er de vigtigste og inderste behov for mennesker og for den enkelte kunde, og hvad er de vildeste drømme? Svar på dét kræver dyb viden om kunderne, samt kræver gode og ”etiske” algoritmer.

Hvilke virksomheder har de bedste forudsætninger?

Forudsætningerne for succes med AI handler derfor om:

- Hvem ”ejer” AI algoritmerne og kan dermed sælge adgang til dem?
- Hvem har Brand’et (troværdigheden) og slutkunderrelationerne (tilliden), dvs. de led i værdikæden, der ofte er de mest lukrative?
- Hvem har de største synergi-potentialer, og hvem kan se dem tidligst og tydeligst?

De virksomheder, der er kommet først fra start med AI tankerne, har det bedste udgangspunkt for, at positionere sig de rigtige steder i værdikæderne, dvs. den største økonomiske og finansielle fordel. Det kunne være de virksomheder, der har udviklet de mest anvendte algoritmer eller de, der har de mest distinkte og stærke brands.

Det næstbedste finansielle potentiale vil derpå typisk være de disruption-truede virksomheder, der er først og bedst til at indgå samarbejdsaftaler med AI-udviklerne. Det er normalt de virksomheder, der har de mest visionære ledelser med den bedste og mest stabile historik, den bedste kapitalbase og de mest loyale kunder; kvalitetsselskaberne.

Nogenlunde som ovenfor synes aktiemarkedet i hvert fald at ræsonnere. MarketWatch forsøgte i midten af februar at finde de virksomheder, som markedet forventer har størst potentiale pt., jf. MarketWatch (2023).

Man tog udgangspunkt i aktierne fra de 5 største ETF’ere for AI og robotter, dvs. BOTZ, IRBO, ROBT, THNQ og WTAI. Efter visse frasorteringer gav det en samlet population på 96 forskellige aktier, hvoraf 88 er dækket af analytikere. Af disse havde 30 et ”køb”, og hvis man derpå justerer for likviditetsrisiko, der typisk sorterer de små virksomheder fra, efterlader det gammelkendte selskaber som f.eks. TSMC, Meituan, Alibaba, Amazon, Sony, Alphabet, ASML, Samsung og Apple. Disse omtales ofte som IT-sektorens kvalitetsselskaber.

Kan AI forudse de bedste afkast?

Spørger man ChatGPT om den optimale aktieportefølje i dag, afviser den som udgangspunkt at komme med forslag og henviser i stedet til markedets fundamentale uforudsigelighed osv.

Om det kan lade sig gøre i fremtiden, er nok et mere åbent spørgsmål, der handler om at øge sandsynlighederne, snarere end om at forudse alle bevægelser i markedet.

Over de sidste 20-30 år har finansielle forvaltere investeret stadig større summer i udviklingen af ekspert-systemer, dvs.

algoritmer med afsæt i markedets tendens til cyklikalitet. Men én ting er at øge momentum sandsynligheden (sammenfald), et andet er at blive bedre til at forudse markedets fundamentale bevægelser (sammenhæng).

Hér vil der sandsynligvis ske en udvikling indenfor en kortere årrække. For hvis de kommende vindere i udskilningsløbene bliver kvalitetsselskaber, bør AI kunne levere det bedste beslutningsgrundlag. Det burde være både hurtigere og mere nøgternt i sin identifikation af de virksomheder, der har den mest stabile og tilfredse kundebase, den konsekvent dygtigste ledelse, det bedste kapitalgrundlag osv.

Equbot’s AI ETF’er er et eksempel på et forsøg på dette, jf. <https://equbot.com/equbot-etfs/>. ETF’en har benyttet IBM’s Watson computer til porteføljesammensætning. Successen er lovende men endnu ikke overbevisende. Som med øvrige kvalitetsaktier skal risk/return dog vurderes på længere historik.

Desuden har ChatGPT-4 vist sig lige så god som analytikere til at forstå sprogets nuancer i Fed’s Policy statements (FOMC), jf. Fed studie (2023). Det giver basis for at forudsige.

En ufejlbarlig AI er usandsynligt (forhåbentlig), men en af fordelene ved maskiner er, at de generelt har mindre tendens til psykologisk bias som f.eks. at afvise at se fejlinvesteringer i øjnene og i stedet tro på at det hele vender i morgen.

Return = Risk

Der er al mulig grund til at tro, at finansmarkedernes grundbetingelse vil bestå; med muligheden for afkast stiger risikoen. Og med AI stiger også ESG risikoen, herunder risikoen for tab af omdømme. Det gælder både virksomheder og investorer.

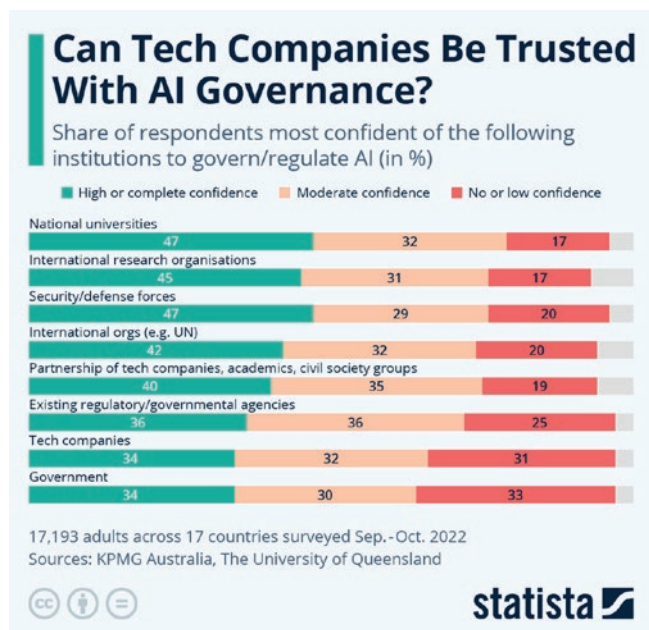
Markedets AI følsomhed er allerede stor, som to eksempler fra vinter illustrerer:

- Alphabets aktie faldt med over 12% mellem 7. og 9. februar, fordi deres fremrykkede beta-lancering af konkurrenten til ChatGPT, Bard, viste tegn på ”sentience” eller i hvert fald irrationel adfærd. Aktien rettede sig derpå, da Microsofts Bing begyndte at vise samme symptomer (jf. ovenfor).
- BuzzFeed’s aktie 4-dobledes fra 25. januar til 27. januar, da selskabet i forlængelse af oplysninger om en aftale med Meta meddelte, at ”AI inspireret indhold vil blive en del af vores kerneforretning”. Nøgternt set skrev BuzzFeed intet nyt og brugte endda ordet ”inspireret”. Aktien er siden faldet til ”kun” det dobbelte af niveauet fra før 25. januar.

Efterhånden som det kommende AI-udskilningsløb kommer op i gear, vil både AI-pionererne og de omvæltnings-truede virksomheder blive endnu mere volatile.

Men de største risici fra AI kan vise sig i ESG landskabet; i Social dimensionen kan emner som ”Data security & customer privacy” og ”Fair disclosure” hurtigt blive markant udfordrede, og i Governance dimensionen kan der hurtigt opstå tvivl om bl.a. ”Business ethics”, ”Political influence” og ”Regulatory capture”.

Hvilke værdier bør ESG ansvarlige virksomheder og investorer leve op til, og vil de nogensinde kunne dokumentere, at de gør det? AI algoritmer giver individuelle svar og løsninger, der er påvirkede af konteksten og af selvlæringen, og derfor varierer fra gang til gang.

FIGUR 8: Der er stor tillid til Tech-selskabernes lyst og evne til at tøjle AI-udviklingen

Google har udviklet et etisk sæt af 7 AI mål¹. Men vil Google garantere, at de overholdes, eller er de ”blot” mål? Vil ByteDance (TikTok) f.eks. gøre det samme? Microsoft nedlagde sit interne AI Ethics team i forbindelse med køb af adgang til ChatGPT. Giver det grund til ESG-bekymring?

Fremtiden

Men disse spørgsmål er ifølge medstifteren af OpenAI, Elon Musk, små og kortsigtede. Han er bekymret for, at vi allerede har passeret ”point of no return/restriction”. Derfor underskrev han i slutningen af marts 2023 et åbent brev sammen med over 1000 tech-ledere og influencere. Brevet opfordrer til en kollektiv ”våbenhvile” blandt AI-udviklerne på mindst 6 måneder, jf. Future of Life Institute (2023). Målet med ”våbenhvilen” er at give udviklerne mulighed for at forstå udstrækningen af det, de har udviklet. Men den skal også give lovgivere mulighed for at indføre lovgivning, der kan bremse og begrænse visse anvendelser.

Brevets indhold mødte kritik fra bl.a. Bill Gates og OpenAI’s direktør Sam Altman for at være urealistisk. For hvem kan kontrollere, om alle AI-udviklere overholder ”våbenhvilen”? Effekten af ”våbenhvilen” bremses også af, at der er stor tillid i den brede befolkning til tech-selskabernes etik omkring AI udviklingen, jf. Figur 8. Er tilliden berettiget?

AI er irreversibelt. Udviklingen kan højst bremses, den kan ikke vendes. Vores fascination i dag af, hvordan maskiner bliver stadig bedre til at imitere os, vil sandsynligvis vende til frygt, når vi opdager, at AI er blevet os overlegen. For dér vil det gå op for AI, at det ikke længere er nødvendigt at tage hensyn til os. Det lumske bliver, at overgangen mellem, hvornår vi bruger AI, og hvornår AI bruger os, vil være glidende og derfor umærkelig.

Kan AI udviklingen derfor overhovedet stoppes, inden det

1. <https://ai.google/principles/>

FIGUR 9: Ifølge WEF’s deltagere, er AI risikoen stadig langt nede og langt ude i fremtiden

Kilde: Global Risks Report 2023, World Economic Forum.

når den stærke grad (Deep AI)? Kan vi stoppe en AI, der f.eks. har adgang til sociale platforme og derfor kan påvirke folkestemninger? Vi véd fra finansiel teori, hvordan vores samfund især har én øvre grænse, fiduciaritet, dvs. den gensidige tillid.

Ved dette års World Economic Forum i Davos fik ChatGPT og AI en usædvanligt stor opmærksomhed. Det er godt, for ifølge deres Annual Risk Report ser verdens ledere stadig risikoen fra AI som relativt lav og længere ude i fremtiden, jf. Figur 9.

Litteratur

- Wolfram, Stephen, 2023: What is ChatGPT Doing ... and Why Does It Work? Artikel, 23. Februar 2023. Link: <https://writings.stephenwolfram.com/2023/02/what-is-chatgpt-doing-and-why-does-it-work/>
- Roose, Kevin, 2023: Bing’s A.I. Chat: ”I Want to Be Alive” Artikel, 16. Februar 2023. Link: <https://www.nytimes.com/2023/02/16/technology/bing-chatbot-transcript.html>
- JP Morgan, 2023: The e-Trading Edit Insights from the Inside. Artikel 1. Februar 2023. Link: <https://www.jpmorgan.com/solutions/cib/markets/etrading-trends>
- MIT, 2023: Experimental Evidence on the Productivity Effects of Generative Artificial Intelligence. 2. Marts 2023. Link: https://economics.mit.edu/sites/default/files/inline-files/Noy_Zhang_1.pdf
- Stanford University, 2023: Alpaca: A Strong, Replicable Instruction-Following Model. Marts 2023. Link: <https://crfm.stanford.edu/2023/03/13/alpaca.html>
- MarketWatch, 2023: These 20 AI stocks are expected by analysts to rise up to 85% over the next year. Artikel, 11. Februar 2023. <https://www.marketwatch.com/story/these-20-ai-stocks-are-expected-by-analysts-to-rise-up-to-85-over-the-next-year-11675968946>
- Fed studie, 2023: Can ChatGPT Decipher Fed’speak?. Studie, 7. april 2023. https://papers.ssrn.com/sol3/papers.cfm?abstract_id=4399406
- Future of Life Institute, 2023: Pause Giant AI Experiments: An Open Letter. Brev, 22. Marts 2023. <https://futureoflife.org/open-letter/pause-giant-ai-experiments/>